



PROSIDING SEMINAR NASIONAL METODE KUANTITATIF

SNMK 2017

PENGUNAAN MATEMATIKA, STATISTIKA,
DAN KOMPUTER DALAM BERBAGAI DISIPLIN ILMU
UNTUK MEWUJUDKAN KEMAKMURAN BANGSA



**SEMINAR NASIONAL
METODE KUANTITATIF
2017**

PROSIDING
Seminar Nasional
Metode Kuantitatif 2017

ISBN No. 978-602-98559-3-7

Penggunaan Matematika, Statistika, dan Komputer dalam Berbagai Disiplin Ilmu
untuk Mewujudkan Kemakmuran Bangsa

Editor :

Prof. Mustofa Usman, Ph.D

Dra. Wamiliana, M.A., Ph.D.

Layout & Design :

Shela Malinda Tampubolon

Alamat :

Fakultas Matematika dan Ilmu Pengetahuan Alam

Universitas Lampung, Bandar Lampung

Jl. Prof. Dr. Sumantri Brojonegoro No. 1 Bandar Lampung

Telp. 0721-701609/Fax. 0721-702767

KATA SAMBUTAN KETUA PELAKSANA SEMINAR NASIONAL METODE KUANTITATIF 2017

Seminar Nasional Metode Kuantitatif 2017 diselenggarakan oleh Jurusan Matematika Fakultas Matematika dan Ilmu Pengetahuan Alam (FMIPA) Universitas Lampung yang dilaksanakan pada tanggal 24 – 25 November 2017. Seminar terselenggara atas kerja sama Jurusan Matematika FMIPA, Lembaga Penelitian dan Pengabdian Masyarakat (LPPM) Unila, dan Badan Pusat Statistik (BPS).

Peserta dari Seminar dihadiri lebih dari 160 peserta dari 11 institusi di Indonesia, diantaranya : Kementerian Pendidikan dan Kebudayaan, Badan Pusat Statistik, Universitas Indonesia, Institut Teknologi Bandung, Universitas Sriwijaya, Universitas Jember, Universitas Islam Negeri Sunan Gunung Djati, Universitas Cendrawasih, Universitas Teknokrat Indonesia, Universitas Malahayati, dan Universitas Lampung. Dengan jumlah artikel yang disajikan ada sebanyak 48 artikel hal ini merefleksikan pentingnya seminar nasional metode kuantitatif dengan tema “penggunaan matematika, statistika dan computer dalam berbagai disiplin ilmu untuk mewujudkan kemakmuran bangsa”.

Kami berharap seminar ini menjadi tempat untuk para dosen dan mahasiswa untuk berbagi pengalaman dan membangun kerjasama antar ilmunan. Seminar semacam ini tentu mempunyai pengaruh yang positif pada iklim akademik khususnya di Unila.

Atas nama panitia, kami mengucapkan banyak terima kasih kepada Rektor, ketua LPPM Unila, dan Dekan FMIPA Unila serta ketua jurusan matematika FMIPA Unila dan semua panitia yang telah bekerja keras untuk suksesnya penyelenggaraan seminar ini.

Dan semoga seminar ini dapat menjadi agenda tahunan bagi jurusan matematika FMIPA Unila`

Bandar Lampung, Desember 2017

Prof. Mustofa Usman, Ph.D

Ketua Pelaksana

KEPANITIAAN

Penasehat : 1. Prof. Dr. Hasriadi Mat Akin, M.P
2. Prof. Dr. Bujang Rahman
3. Prof. Dr. Ir. Kamal, M.Sc
4. Ir. Warsono, M.Sc., Ph.D
5. Dr. Hartoyo, M.Si

Pengarah : 1. Prof. Warsito, S.Si., DEA, Ph.D
2. Prof. Dr. Sutopo Hadi, S.Si., M.Sc
3. Dian Kurniasari S.Si., M.Sc
4. Drs. Suratman Umar, M.Sc.

Penanggung Jawab : Dra. Wamiliana, M.A., Ph.D

Ketua Pelaksana : Prof. Drs. Mustofa, M.A., Ph.D

Sekretaris : Dra. Dorrah Aziz, M.Si

Bendahara : Amanto, S.Si., M.Sc

Kesekretariatan : Subian Saidi, S.Si., M.Si

Dr. Notiragayu, M.Si

- Syamsu Huda, S.I.P., M.M

- Srimati, S.Pd

- Johan, S.P

- Riendi Ferdian, S.I.P

- Siti Marbiyah, S.Si

- Rosihin Anwar, S.Kom

- Shela Malinda T

- Della Desiyana

- Nandra Adi Prayoga

- Himatika

Seksi-seksi :

Acara : Dr. Aang Nuryaman, M.Si

Dr. Khoirin Nisa, M.Si

Drs. Rudi Ruswandi, M.Si

Drs. Eri Setiawan, M.Si

Konsumsi : Widiarti S.Si., M.Si
Dr. Asmiati, M.Si

Transportasi/akomodasi : Drs. Nusyirwan, M.Si
Agus Sutrisno, S.Si., M.Si

Perlengkapan : Drs. Tiryono R., M.Sc., Ph.D
- Agus Suroso, A.Md
- Tamrinsyah
- Supriyadi
- Drajat
- Maeda Sulistiana

Reviewer : Drs. Suharsono, M.Sc., Ph.D
- Dr. La Zakaria S.Si., M.Sc
- Dr. Muslim Ansori, S.Si., M.Si
- Dr. Ir. Netti Herawati, M.Sc

DAFTAR ISI

KATA SAMBUTAN	iii
KEPANITIAAN.....	iv
DAFTAR ISI	vi
Aplikasi Metode Analisis Homotopi (HAM) pada Sistem Persamaan Diferensial Parsial Homogen (<i>Fauzia Anisatul Farida</i>).....	1
Simulasi Interaksi Angin Laut dan Bukit Barisan dalam Pembentukan Pola Cuaca di Wilayah Sumatera Barat Menggunakan Model Wrf-Arw (<i>Achmad Rafli Pahlevi</i>)	7
Penerapan Mekanisme Pertahanan Diri (Self-Defense) sebagai Upaya Strategi Pengurangan Rasa Takut Terhadap Kejahatan (Studi Pada Kabupaten/Kota di Provinsi Lampung yang Menduduki Peringkat <i>Crime Rate</i> Tertinggi) (<i>Teuku Fahmi</i>).....	18
Tingkat Ketahanan Individu Mahasiswa Unila pada Aspek Soft Skill (<i>Pitojo Budiono</i>).....	33
Metode Analisis Homotopi pada Sistem Persamaan Diferensial Parsial Linear Non Homogen Orde Satu (<i>Atika Faradilla</i>)	44
Penerapan Neural Machine Translation Untuk Eksperimen Penerjemahan Secara Otomatis pada Bahasa Lampung – Indonesia (<i>Zaenal Abidin</i>)	53
Ukuran Risiko Cre-Var (<i>Insani Putri</i>)	69
Penentuan Risiko Investasi dengan Momen Orde Tinggi $V@R-Cv@R$ (<i>Marianik</i>).....	77
Simulasi Komputasi Aliran Panas pada Model Pengering Kabinet dengan Metode Beda Hingga (<i>Vivi Nur Utami</i>).....	83
Segmentasi Wilayah Berdasarkan Derajat Kesehatan dengan Menggunakan <i>Finite Mixture Partial Least Square</i> (Fimix-Pls) (<i>Agustina Riyanti</i>).....	90
Representasi Operator Linier Dari Ruang Barisan Ke Ruang Barisan L $3/2$ (<i>Risky Aulia Ulfa</i>).....	99
Analisis Rangkaian Resistor, Induktor dan Kapasitor (RLC) dengan Metode Runge-Kutta Dan Adams Bashforth Moulton (<i>Yudandi Kuputra Aji</i>)	110
Representasi Operator Linier dari Ruang Barisan Ke Ruang Barisan L $13/12$ (<i>Amanda Yona Ningtyas</i>)	116
Desain Kontrol Model Suhu Ruangan (<i>Zulfikar Fakhri Bismar</i>)	126

Penerapan Logika Fuzzy pada Suara Tv Sebagai Alternative Menghemat Daya Listrik (Agus Wantoro)	135
Clustering Wilayah Lampung Berdasarkan Tingkat Kesejahteraan (Henida Widyatama).....	149
Pemanfaatan Sistem Informasi Geografis Untuk Valuasi Jasa Lingkungan Mangrove dalam Penyakit Malaria di Provinsi Lampung (Imawan Abdul Qohar)	156
Analisis Pengendalian Persediaan Dalam Mencapai Tingkat Produksi <i>Crude Palm Oil</i> (CPO) yang Optimal di PT. Kresna Duta Agroindo Langling Merangin-Jambi (Marcelly Widya W.)	171
Analisis <i>Cluster Data Longitudinal</i> pada Pengelompokan Daerah Berdasarkan Indikator IPM di Jawa Barat (A.S Awalluddin).....	187
Indek Pembangunan Manusia dan Faktor Yang Mempengaruhinya di Daerah Perkotaan Provinsi Lampung (Ahmad Rifa'i).....	195
<i>Parameter Estimation Of Bernoulli Distribution Using Maximum Likelihood and Bayesian Methods</i> (Nurmaita Hamsyiah).....	214
Proses Pengamanan Data Menggunakan Kombinasi Metode Kriptografi <i>Data Encryption Standard</i> dan <i>Steganografi End Of File</i> (Dedi Darwis).....	228
<i>Bayesian Inference of Poisson Distribution Using Conjugate A and Non-Informative Prior</i> (Misgiyati).....	241
Analisis Klasifikasi Menggunakan Metode Regresi Logistik Ordinal dan Klasifikasi Naïve Bayes pada Data Alumni Unila Tahun 2016 (Shintia Faramudhita).....	251
Analisis Model <i>Markov Switching Autoregressive</i> (MSAR) pada Data <i>Time Series</i> (Aulianda Prasyanti)	263
Perbandingan Metode Adams Bashforth-Moulton dan Metode Milne-Simpson dalam Penyelesaian Persamaan Diferensial Euler Orde-8 (Faranika Latif).....	278
Pengembangan Ekowisata dengan Memanfaatkan Media Sosial untuk Mengukur Selera Calon Konsumen (Gustafika Maulana).....	293
Diagonalisasi Secara Uniter Matriks Hermite dan Aplikasinya pada Pengamanan Pesan Rahasia (Abdurrois)	308
Pembandingan Metode Runge-Kutta Orde 4 dan Metode Adam-Bashfort Moulton dalam Penyelesaian Model Pertumbuhan Uang yang Diinvestasikan (Intan Puspitasari)	328
Menyelesaikan Persamaan Diferensial Linear Orde-N Non Homogen dengan Fungsi Green (Fathurrohman Al Ayubi).....	341
Penyelesaian Kata Ambigu pada Proses Pos Tagging Menggunakan Algoritma <i>Hidden Markov Model</i> (HMM) (Agus Mulyanto).....	347

Sistem Temu Kembali Citra Daun Tumbuhan Menggunakan Metode Eigenface (Supiyanto)	359
Efektivitas Model <i>Problem Solving</i> dalam Meningkatkan Kemampuan Berfikir Lancar Mahasiswa pada Materi Ph Larutan (Ratu Betta Rudibyani)	368
<i>The Optimum Bandwidth for Kernel Density Estimation of Skewed Distribution: A Case Study on Survival Data of Cancer Patients</i> (Netti Herawati)	380
Karakteristik Larutan Kimia Di Dalam Air Dengan Menggunakan Sistem Persamaan Linear (Titik Suparwati)	389
Bentuk Solusi Gelombang Berjalan Persamaan $\Delta\Delta$ mKdV Yang Diperumum (Notiragayu)	398
Pendugaan Blup Dan Eblup(Suatu Pendekatan Simulasi) (Nusyirwan)	403

APLIKASI METODE ANALISIS HOMOTOPI (HAM) PADA SISTEM PERSAMAAN DIFERENSIAL PARSIAL HOMOGEN

Fauzia Anisatul Farida¹, Suharsono S.¹ dan DorrahAzis¹

¹ Jurusan Matematika FMIPA UNILA, Jl. Soemantri Brojonegoro No.1, Gd. Meneng, B.Lampung 35145
Email: anisafauzhia@gmail.com, suharsono.1962@fmipa.unila.ac.id, dorrah.aziz@fmipa.unila.ac.id

ABSTRAK

Masalah non linear biasanya sulit diselesaikan baik secara analitik maupun secara numerik. Oleh sebab itu, maka dibutuhkan metode-metode yang dapat menyelesaikan persamaan diferensial parsial ini. Metode analisis homotopi berhasil diterapkan dalam menyelesaikan berbagai tipe persamaan dan sistem persamaan tak linier, homogeny atau tak homogen. Penelitian ini bertujuan untuk menyelesaikan system persamaan diferensial parsial dengan metode baru yaitu metode analisis homotopi (HAM) yang memiliki beberapa keunggulan dibandingkan dengan metode sebelumnya. Untuk memperlihatkan bahwa solusi dari metode homotopi mendekati solusi eksak dalam penelitian ini juga dilakukan investigasi numerik. Dari penelitian ini di peroleh kesimpulan bahwa solusi homotopi mendekati solusi eksak untuk $h = 1$.

Kata kunci: HAM, solusieksak, sistem PDE

1. PENDAHULUAN

Matematika merupakan ratu bagi ilmu pengetahuan lain, dimana perkembangan ilmu matematika memiliki peran yang sangat penting dan bermanfaat bagi kehidupan sehari-hari serta merupakan penunjang bagi perkembangan ilmu pengetahuan dan teknologi.

Dalam kehidupan sehari-hari banyak ditemui permasalahan yang berhubungan dengan matematika. Permasalahan ini biasanya berhubungan dengan persamaan diferensial, khususnya persamaan diferensial parsial, baik diferensial parsial linear maupun nonlinear. Masalah nonlinear ini biasanya sulit diselesaikan baik secara analitik maupun secara numerik. Oleh sebab itu, maka dibutuhkan metode-metode yang dapat menyelesaikan persamaan diferensial parsial ini.

Beberapa penelitian difokuskan pada penemuan metode untuk memperoleh solusi dari masalah yang dimodelkan dalam persamaan nonlinear. Beberapa metode yang digunakan antara lain, metode homotopi, persamaan aljabar, persamaan diferensial biasa dan persamaan diferensial parsial. Selain itu, metode analisis homotopi adalah metode yang bebas, artinya tidak memperhatikan kecil atau besarnya suatu parameter. Metode analisis homotopi berhasil diterapkan dalam menyelesaikan berbagai tipe persamaan dan sistem persamaan tak linier, homogen atau tak homogen. Oleh karena itu, penulis menggunakan metode homotopi dalam menyelesaikan permasalahan sistem persamaan diferensial parsial ini.

2. LANDASAN TEORI

Metode analisis homotopi (HAM) pertama kali dirancang pada tahun 1992 oleh Shijun Liao dari Shanghai Jiaotong University dalam disertasi PhD-nya dan dimodifikasi lebih lanjut pada tahun 1997 untuk

memperkenalkan parameter tambahan nol-nol atau disebut sebagai parameter konvergensi-kontrol yang dilambangkan dengan c_0 untuk membangun homotopi pada sistem diferensial dalam bentuk umum. Metode analisis homotopi (HAM) adalah teknik semi analitis untuk memecahkan masalah tak linear biasa atau persamaan diferensial parsial. Misalkan terdapat persamaan diferensial berikut,

$$N_i[z_i(x, t)] = 0 \quad ; i = 1, 2, \dots, n \quad (1)$$

Dimana N_i adalah operator nonlinear yang mewakili seluruh persamaan, x dan t adalah variable bebas dan $z_i(x, t)$ adalah fungsi yang tidak diketahui. Dengan cara mengeneralisasi metode homotopi sederhana, [1] menyusun persamaan deformasi orde nol.

$$(1 - q)L[\phi_i(x, t; q) - z_{i,0}(x, t)] = qh_iN_i[\phi_i(x, t; q)] \quad (2)$$

Dimana $q \in [0, 1]$ adalah parameter *benamanhi* adalah fungsi bantu tak nol tambahan, L adalah operator linear tambahan, $z_{i,0}(x, t)$ adalah tebakan awal dari $z_i(x, t)$ dan $\phi(x, t; q)$ adalah fungsi yang tidak diketahui. Penting untuk diingat bahwa, bebas untuk memilih objek tambahan seperti h_i dan L pada HAM. Terlihat jelas bahwa saat $q = 0$ dan $q = 1$, keduanya menghasilkan

$$\phi_i(x, t; 0) = z_{i,0}(x, t) \text{ dan } \phi_i(x, t; 1) = z_i(x, t). \quad (3)$$

Maka, sejalan dengan meningkatnya q dari 0 ke 1, solusi $\phi_i(x, t; q)$ berubah dari tebakan awal $z_{i,0}(x, t)$ ke solusi $z_i(x, t)$. Memperluas $\phi_i(x, t; q)$ dalam deret Taylor terhadap q , akan menghasilkan

$$\phi_i(x, t; q) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t)q^m \quad (4)$$

Dimana

$$z_{i,m} = \frac{1}{m!} \frac{\delta^m \phi_i(x, t; q)}{\delta q^m} \Big|_{q=0} \quad (5)$$

Jika operator linear tambahan, tebakan awal, parameter tambahan h_i , dan fungsi tambahan dipilih dengan benar, maka deret pada persamaan (3) konvergen ke $q = 1$ dan

$$\phi_i(x, t; 1) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t) \quad (6)$$

Yang merupakan satu dari solusi – solusi persamaan – persamaan nonlinear aslinya, seperti yang telah dibuktikan oleh Liao. Jika $h_i = -1$, persamaan (2) menjadi

$$(1 - q)L[\phi_i(x, t; q) - z_{i,0}(x, t)] + qN_i[\phi_i(x, t; q)] = 0 \quad (7)$$

Berdasarkan (4), persamaan yang dibentuk dapat di deduksi dari persamaan deformasi orde nol (2). Kita definisikan vektor

$$\vec{z}_{i,n} = \{ z_{i,0}(x, t), z_{i,1}(x, t), \dots, z_{i,n}(x, t) \} \quad (8)$$

Mendiferensialkan (2) sebanyak m kali terhadap parameter *benaman* q dan lalu masukkan $q = 0$ dan akhirnya membaginya dengan $m!$, kita dapatkan yang disebut dengan persamaan deformasi orde ke- m .

$$L[z_{i,m}(x, t) - X_m z_{i,m-1}(x, t)] = h_i R_{i,m}(\vec{z}_{i,m-1}) \quad (9)$$

Dimana

$$R_{i,m}(\vec{z}_{i,m-1}) = \frac{1}{(m-1)!} \frac{\delta^{m-1} N_i[\phi_i(x, t; q)]}{\delta q^{m-1}} \Big|_{q=0} \quad (10)$$

Dan

$$X_m = \begin{cases} 0, & m \leq 1 \\ 1, & m > 1 \end{cases} \quad (11)$$

Harus ditekankan bahwa $z_{i,m}(x, t) (m \geq 1)$ dibentuk oleh persamaan linear (9) dengan syarat batas linear yang didapat dari masalah awalnya. Selanjutnya akan dicari solusi dari sistem persamaan diferensial parsial [2]

$$u_t - v_x + (u + v) = 0; \quad v_t - u_x + (u + v) = 0 \quad (12)$$

3. METODOLOGI PENELITIAN

Metode yang digunakan dalam menyelesaikan permasalahan diferensial parsial dengan menggunakan metode analisis homotopi (HAM) adalah sebagai berikut:

1. Mengumpulkan bahan literatur serta studi pustaka seperti buku dan jurnal dari perpustakaan dan internet.

2. Merumuskan masalah dalam bentuk persamaan diferensial parsial tak linear.

Adapun langkah-langkah dalam menyelesaikan permasalahan persamaan diferensial parsial $u_t - v_x + (u + v) = 0$ dan $v_t - u_x + (u + v) = 0$ dengan metode analisis homotopi adalah sebagai berikut :

1. Misalkan diberikan suatu persamaan nilai awalnya

$$u(x, t) = \sinh x, \quad v(x, t) = \cosh x$$

2. Menyelesaikan persamaan (3.1) menggunakan metode analisis homotopi dengan operator bantu linear L.

- a. Menjelaskan operator linear L.
- b. Menentukan persamaan deformasi orde nol.
- c. Menentukan rangkaian solusi dari komponen yang telah diperoleh melalui perkiraan awal jika $q = 0$, $q = 1$.
- d. Mengkonstruksikan persamaan deformasi orde-m.
- e. Menentukan solusi persamaan deformasi pada persamaan yang diperoleh dari langkah ke (d) untuk setiap $m = 1, 2, 3, \dots, n$.
- f. Mensubstitusikan hasil ini kedalam deret homotopi, sehingga diperoleh solusi homotopi.

4. PEMBAHASAN

Dengan syarat awal

$$u(x, t) = \sinh x, \quad v(x, t) = \cosh x \quad (13)$$

Untuk memecahkan persamaan 12 dengan menggunakan metode analisis homotopi (HAM).

Tulis operator linear sebagai berikut:

$$L[\varphi(x, t; q)] = \frac{\delta \varphi(x, t; q)}{\delta t} \quad (14)$$

Selanjutnya definisikan operator nonlinear sebagai berikut

$$N_1[\varphi_i(x, t; q)] = \frac{\delta \varphi_1(x, t; q)}{\delta t} - \frac{\delta \varphi_2(x, t; q)}{\delta x} + \varphi_1(x, t; q) + \varphi_2(x, t; q)$$

$$N_2[\varphi_i(x, t; q)] = \frac{\delta \varphi_2(x, t; q)}{\delta t} - \frac{\delta \varphi_1(x, t; q)}{\delta x} + \varphi_1(x, t; q) + \varphi_2(x, t; q)$$

Dengan menggunakan definisi tersebut, diperoleh persamaan deformasi rangka orde-nol

$$(1-q)L[\varphi_i(x, t; q) - z_{i,0}(x, t)] = qh_i N_i[\varphi_i(x, t; q)], \quad i = 1, 2 \quad (15)$$

Ketika $q=0$ dan $q=1$ maka

$$\varphi_1(x, t; q) = z_{1,0}(x, t) = u_0(x, t), \quad \varphi_2(x, t; q) = z_{2,0}(x, t) = v_0(x, t)$$

$$\varphi_1(x, t; q) = u(x, t), \quad \varphi_2(x, t; q) = v(x, t)$$

Oleh karena itu, karena parameter q meningkat dari 0 ke 1, $\varphi_i(x, t; q)$ bervariasi dari perkiraan awal $z_{i,0}(x, t)$ ke solusi $z_{i,m}(x, t)$ untuk $i = 1, 2$. Memperluas $\varphi_i(x, t; q)$ menggunakan deret Taylor sehubungan dengan $q=1$, sehingga diperoleh solusi:

$$\varphi_i(x, t; q) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t)q^m$$

Dimana

$$z_{i,m}(x, t) = \frac{1}{m!} \left. \frac{\partial^m \varphi_i(x, t; q)}{\partial q^m} \right|_{q=0}$$

Jika operator linier tambahan, tebakan awal dan parameter pelengkap h_i dipilih dengan benar, rangkaian di atas dalam konvergen pada $q = 1$, didapat solusi:

$$u(x, t) = z_{1,0}(x, t) + \sum_{m=1}^{+\infty} z_{1,m}(x, t) \quad (16)$$

$$v(x, t) = z_{2,0}(x, t) + \sum_{m=1}^{+\infty} z_{2,m}(x, t) \quad (17)$$

Yang harus menjadi salah satu solusi persamaan nonlinear asli, seperti yang dibuktikan oleh Liao. Sekarang didefinisikan vektornya

$$\overrightarrow{z_{i,n}} = \{ z_{i,0}(x, t), z_{i,1}(x, t), \dots, z_{i,n}(x, t) \}$$

Jadi, persamaan deformasi orde- m adalah:

$$L[z_{i,m}(x, t) - X_m z_{i,m-1}(x, t)] = h_i R_{i,m}(\overrightarrow{z_{i,m-1}}) \quad (18)$$

Dengan

$$z_{i,m}(x, 0) = 0$$

Dimana

$$R_{1,m}(\overrightarrow{z_{1,m-1}}) = (z_{1,m-1})_t - (z_{2,m-1})_x + z_{1,m-1} + z_{2,m-1}$$

$$R_{2,m}(\overrightarrow{z_{2,m-1}}) = (z_{2,m-1})_t - (z_{1,m-1})_x + z_{1,m-1} + z_{2,m-1}$$

Sekarang, solusi dari persamaan deformasi orde- m untuk $m \geq 1$ adalah

$$z_{i,m}(x, t) = X_m z_{i,m-1}(x, t) + h_i \int_0^t R_{i,m}(\overrightarrow{z_{i,m-1}}) \partial \tau + C_i \quad (19)$$

Berikut ini merupakan hasil dari persamaan (19) dengan mencari dahulu nilai $R_{i,m}(\overrightarrow{z_{i,m-1}})$, sehingga diperoleh untuk:

$$z_{1,0}(x, t) = u_0(x, t) = \sinh x \quad (20)$$

$$z_{1,1}(x, t) = ht \cosh x \quad (21)$$

$$z_{1,2}(x, t) = \frac{ht}{2} [2(1 + h) \cosh x + ht \sinh x] \quad (22)$$

$$z_{1,3}(x, t) = \frac{ht}{6} [(6 + 12h + 6h^2 + h^2t^2) \cosh x + 6ht(1 + h) \sinh x] \quad (23)$$

$$z_{1,4}(x, t) = \frac{ht}{24} [(24 + 72h + 72h^2 + 24h^3 + 12h^2t^2 + 12h^3t^2) \cosh x (36ht + 72h^2t + 36h^3t + h^3t^3) \sinh x] \quad (24)$$

$$z_{1,5}(x, t) = \frac{ht}{120} [(120 + 480h + 720h^2 + 480h^3 + 120h^4 + 120h^2t^2 + 240h^3t^2 + 120h^4t^2 + h^4t^4) \cosh x + (240ht + 720h^2t + 720h^3t + 20h^3t^3 + 240h^4t + 20h^4t^3) \sinh x] \quad (25)$$

$$z_{2,0}(x, t) = v_0(x, t) = \cosh x \quad (26)$$

$$z_{2,1}(x, t) = ht \sinh x \quad (27)$$

$$z_{2,2}(x, t) = \frac{ht}{2} [2(1 + h) \sinh x + ht \cosh x] \quad (28)$$

$$z_{2,3}(x, t) = \frac{ht}{6} [(6 + 12h + 6h^2 + h^2t^2) \sinh x + 6ht(1 + h) \cosh x] \quad (29)$$

$$z_{2,4}(x, t) = \frac{ht}{24} [(24 + 72h + 72h^2 + 24h^3 + 12h^2t^2 + 12h^3t^2) \sinh x (36ht + 72h^2t + 36h^3t + h^3t^3) \cosh x] \quad (30)$$

$$z_{2,5}(x, t) = \frac{ht}{120} [(120 + 480h + 720h^2 + 480h^3 + 120h^4 + 120h^2t^2 + 240h^3t^2 + 120h^4t^2 + h^4t^4) \sinh x + (240ht + 720h^2t + 720h^3t + 20h^3t^3 + 240h^4t + 20h^4t^3) \cosh x] \quad (31)$$

Selanjutnya ketika $h = -1$ hasil yang diperoleh dari persamaan (20) yang disubstitusikan ke dalam persamaan (12) menjadi:

$$u(x, t) = \sinh x \left(1 + \frac{t^2}{2!} + \frac{t^4}{4!} + \dots \right) - \cosh x \left(t + \frac{t^3}{3!} + \frac{t^5}{5!} + \dots \right) \quad (32)$$

$$v(x, t) = \cosh x \left(1 + \frac{t^2}{2!} + \frac{t^4}{4!} + \dots \right) - \sinh x \left(t + \frac{t^3}{3!} + \frac{t^5}{5!} + \dots \right) \quad (33)$$

Sehingga diperoleh solusi homotopi:

$$u(x, t) = \sinh(x - t), \quad v(x, t) = \cosh(x - t) \quad (34)$$

Berdasarkan perhitungan dengan mengikuti langkah-langkah dalam metode homotopi, diperoleh kesimpulan bahwa untuk $h = -1$, solusi eksak dapat dicari dengan solusi homotopi.

Selanjutnya akan diperiksa bahawa $u(x, t)$ dan $v(x, t)$ memenuhi persamaan (12)

$$u(x, t) = \sinh(x - t) \quad v(x, t) = \cosh(x - t)$$

$$u_t = -\cosh(x - t) \quad v_x = \sinh(x - t)$$

$$u_x = \cosh(x - t) \quad v_t = -\sinh(x - t)$$

Maka:

Untuk persamaan (12)

$$u_t - v_x + (u + v) = -\cosh(x - t) - \sinh(x - t) + (\sinh(x - t) + \cosh(x - t)) = 0$$

$$v_t - u_x + (u + v) = -\sinh(x - t) - \cosh(x - t) + (\sinh(x - t) + \cosh(x - t)) = 0$$

Jadi berdasarkan beberapa uraian di atas persamaan (34) adalah solusi analitik untuk persamaan (12).

4. KESIMPULAN

Berdasarkan penguraian langkah-langkah yang telah dikerjakan dalam metode homotopi ini, diperoleh kesimpulan bahwa solusi dari persamaan diferensial parsial $u_t - v_x + (u + v) = 0$ dan $v_t - u_x + (u + v) = 0$ yang diperoleh berdasarkan metode analisis homotopi dengan hasil $u(x, t) = \sinh(x - t)$, $v(x, t) = \cosh(x - t)$ sama dengan solusi eksaknya jika $h = -1$.

DAFTAR PUSTAKA

- [1] Liao, S. (2012). Homotopy Analysis Method in Nonlinear Differential Equation. Beijing: Higher Education Press.
- [2] Bataineh, Noorani, dan Hashim. (2007). Approximate Analytical Solutions of System of PDEs by Homotopy Analysis Method. *Journal of Computers and Mathematics with Applications*, Malaysia, 55, 2913-2923.

SIMULASI INTERAKSI ANGIN LAUT DAN BUKIT BARISAN DALAM PEMBENTUKAN POLA CUACA DI WILAYAH SUMATERA BARAT MENGGUNAKAN MODEL WRF-ARW

Achmad Rafli Pahlevi

Magister Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam
Universitas Lampung
Stasiun Meteorologi Maritim Lampung
Badan Meteorologi Klimatologi dan Geofisika
Email: raflithium@hotmail.com

ABSTRAK

Wilayah Padang terletak di wilayah equator, tepat pada lintang 0. Seperti daerah lainnya di Indonesia, iklim Padang secara umum bersifat tropis dengan suhu udara yang cukup tinggi, yaitu antara 22,6° C sampai 31,5° C. Padang memiliki tipe iklim Ekuatorial yaitu daerah yang antara musim hujan dan musim kemarau tidak jauh berbeda sehingga memiliki dua puncak musim hujan atau bisa dikatakan tidak memiliki musim kering. Untuk itu, faktor lokal menjadi penting dalam pembentukan pola cuaca di wilayah Padang. Angin laut yang membawa uap air dari wilayah Samudera Hindia menjadi sumber utama proses cuaca yang terjadi di wilayah Padang. Adanya Bukit Barisan di Timur Padang yang berjajar sepanjang pulau Sumatera memberikan pengaruh sendiri terhadap pembentukan pola cuaca di wilayah ini. Penggunaan model WRF-ARW (Weather Research and Forecast – Advanced Research) dapat mensimulasikan seberapa besar pengaruh interaksi antara angin laut dan Bukit Barisan terhadap pembentukan pola cuaca di wilayah pesisir Padang.

KataKunci: WRF-ARW, Iklim Ekuatorial, Angin Laut

1. PENDAHULUAN

Hujan merupakan salah satu bentuk presipitasi berupa tetes-tetes air yang jatuh dari dasar awan. Diameter dan konsentrasi tetes sangat bervariasi tergantung intensitas presipitasi terutama jenis dan asalnya, yakni hujan kontiniu, hujan shower, dan lain-lain[1]. Curah hujan merupakan salah satu unsur cuaca dan iklim yang memiliki peranan sangat penting terhadap kehidupan di bumi [2], dalam hal ini disamping merupakan sumber daya alam yang dibutuhkan namun juga dapat menjadi sumber bencana [3].

Sumatera Barat terletak di pesisir barat bagian tengah pulau Sumatera yang terdiri dari dataran rendah di pantai barat dan dataran tinggi vulkanik yang dibentuk oleh Bukit Barisan.. Seperti daerah lainnya di Indonesia, iklim Sumatera Barat secara umum bersifat tropis dengan suhu udara yang cukup tinggi, yaitu antara 22,6° C sampai 31,5° C. Sumatera Barat memiliki tipe iklim Ekuatorial yaitu daerah yang antara musim hujan dan musim kemarau tidak jauh berbeda sehingga memiliki dua puncak musim hujan atau bisa dikatakan tidak memiliki musim kering [4].

Selain dipengaruhi oleh faktor regional, cuaca di Padang juga dipengaruhi oleh faktor lokal, yaitu adanya angin laut. Wilayah Padang yang terletak di dekat Samudera Hindia, menjadikan cuaca di Padang, mendapat pengaruh besar dari Samudera Hindia. Angin laut yang membawa uap air dari wilayah Samudera Hindia menjadi sumber utama proses cuaca yang terjadi di wilayah padang. Selain itu, adanya Bukit Barisan yang memanjang di sepanjang pulau Sumatera dapat berpengaruh dalam pembentukan hujan di wilayah Padang.

2. LANDASAN TEORI

2.1. Interaksi Angin Laut

Angin laut merupakan respon dari variasi pemanasan permukaan atmosfer skala meso. Angin laut terjadi terbatas pada lapisan 1 – 2 km, karena itu angin laut, besar dipengaruhi oleh proses viskositas dan konduksi lapisan bidang batas [5]. Angin laut biasanya lebih kuat dibandingkan angin darat, kecepatannya mencapai $4 - 8 \text{ ms}^{-1}$ dan ketebalan lapisan udara mencakup ketinggian 1000 m. Angin laut di tropis dapat masuk ke darat sejauh 100 km. Pada beberapa lokasi, angin laut mungkin dapat mendorong rintangan (*barrier*) topografis pantai dan menembus kedarat. Membedakan angin laut pada jarak lebih 50 km dari pantai akan sulit karena angin ini berinteraksi dengan sirkulasi lokal lain. Pada beberapa kilometer di darat, udara naik dalam akibat dari konveksi angin laut dan kembang menuju laut pada ketinggian 1500 – 3000 m [6].

Jika angin laut konvergen dengan angin dari arah berbeda, maka sering terbentuk “*front angin laut*” yang dapat menyebabkan pembentukan awan lokal dan hujan [7]. Awan tumbuh dalam zona konvergensi antara sistem skala sinoptik dan lokal yang berlawanan ini. Diatas pulau dan semenanjung sistem angin laut yang konvergen dari pantai yang berhadapan dapat menyebabkan curah hujan maksimum pada sore hari [7].

Angin laut yang bergerak terus ke daratan biasanya akan terus naik melewati gunung, tetapi ada beberapa kasus dimana angin tersebut tidak dapat melewati gunung ataupun sela-sela gunung. Jika suatu penghalang dengan bentuk memanjang ke samping, digerakkan secara horizontal dengan kecepatan kecil maka angin tersebut akan terdorong berbalik arah, efek ini disebut ‘*blocking*’. Blocking terjadi saat parsel fluida yang memiliki energi kinetik yang tidak cukup untuk mengatasi gaya apung yang akan dialami saat melewati pegunungan [8]. Panchal (1992) dalam penelitian perubahan kelembaban besar dipengaruhi oleh angin laut dan kelembaban di laut itu sendiri [9]. Pada penelitian Steele tahun 2012, penggunaan WRF-ARW dapat menampilkan untuk prakirawan indikasi dari ketergantungan angin laut pada variabel prognostik dan model fisik [10].

2.2. Model WRF-ARW (*Weather Research Forecast – Advanced Research*)

Model *Weather Research and Forecasting* (WRF) adalah sebuah prediksi cuaca numerik atau *numeric weather prediction* (NWP) dan sistem simulasi atmosfer yang khusus didesain untuk penelitian dan pengaplikasian secara operasional [11]. WRF merupakan hasil kolaborasi dari *National Center for Atmospheric Research's* (NCAR), *Mesoscale and Microscale Meteorology* (MMM) *Division*, *National Oceanic and Atmospheric Administration's* (NOAA), *National Centers for Environmental Prediction* (NCEP), dan *Earth System Research Laboratory* (ESRL).

Model WRF-ARW menggunakan persamaan pengatur (*Governing Equation*) dalam mensimulasikan atmosfer. Persamaan pengatur tersebut terdiri dari beberapa persamaan, yaitu [12]:

1. Koordinat dan Variabel Vertikal

Persamaan model WRF-ARW diformulasikan menggunakan kordinat tekanan-hidrostatik vertikal yang mengikuti bentuk topografi dinotasikan sebagai η dan didefinisikan sebagai:

$$\eta = (p_h - p_{ht})/\mu \text{ dimana } \mu = p_{hs} - p_{ht} \quad (1)$$

p_h adalah komponen hidrostatik dari tekanan, p_{hs} dan p_{ht} merepresentasikan nilai tekanan pada lapisan permukaan dan lapisan atas

2. Persamaan momentum horizonal

Persamaan momentum merupakan persamaan utama yang mengatur pergerakan di atmosfer (pergerakan angin) baik dalam arah horizontal maupun vertikal. Prinsip yang diambil dari persamaan momentum adalah hukum kedua newton yang menyatakan bahwa $F = m \cdot a$, dimana gaya adalah massa dikali percepatan. Persamaan momentum di atmosfer bumi terbentuk dari gabungan beberapa gaya yang bekerja di atmosfer. Persamaan momentum atmosfer didefinisikan dengan:

$$\frac{\partial u}{\partial t} = -u \frac{\partial u}{\partial x} - v \frac{\partial u}{\partial y} - \omega \frac{\partial u}{\partial z} + f v - g \frac{\partial z}{\partial x} + F_x \quad (2.a)$$

$$\frac{\partial v}{\partial t} = -u \frac{\partial v}{\partial x} - v \frac{\partial v}{\partial y} - \omega \frac{\partial v}{\partial z} - f u - g \frac{\partial z}{\partial y} + F_y \quad (2.b)$$

dimana u dan v adalah kecepatan angin dalam arah zonal dan meridional (satuan m/s), x dan y adalah arah zonal dan meridional (dalam m), z adalah arah vertical (dalam m), ω adalah kecepatan vertical (dalam mb/s), f adalah parameter koriolis (dalam s^{-2}), g adalah gravitasi (m/s), serta F_x dan F_y adalah gaya gesek pada atmosfer.

3. Persamaan Kekekalan Massa (Kontinyuitas)

Persamaan kekekalan massa menyatakan bahwa massa udara di atmosfer selalu tetap dan tidak akan pernah berubah berdasarkan waktu. Persamaan kekekalan massa juga disebut dengan persamaan kontinyuitas, yang menyatakan bahwa dalam suatu volume yang tetap, massa udara yang masuk sama dengan jumlah massa udara yang keluar. Persamaan kekekalan massa dinotasikan dengan:

$$\frac{\partial u}{\partial x} + \frac{\partial v}{\partial y} + \frac{\partial \omega}{\partial z} = 0 \quad (3)$$

4. Persamaan Perubahan Uap Air

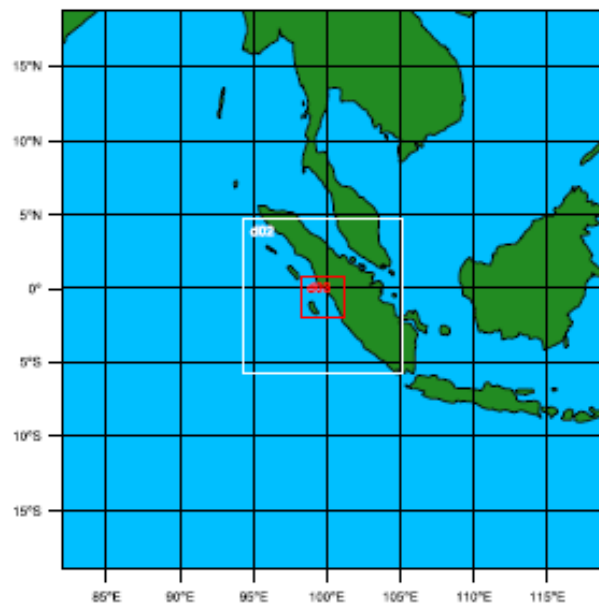
Persamaan ini menjelaskan bagaimana jumlah total uap air pada sebuah parsel udara di atmosfer bergerak, kecuali jika adanya penambahan (dari evaporasi E) atau pengurangan (presipitasi P).

$$\frac{\partial q}{\partial t} = -u \frac{\partial q}{\partial x} - v \frac{\partial q}{\partial y} - \omega \frac{\partial q}{\partial p} + E - P \quad (4)$$

dimana q adalah mixing ratio.

3. METODOLOGI PENELITIAN

Metode yang digunakan dalam penelitian ini adalah simulasi menggunakan WRF-ARW, yang akan mensimulasikan proses angin laut yang mempengaruhi proses terjadinya hujan ekstrem. Parameter cuaca yang ditampilkan adalah vektor angin, RH dan angin meridional sampai lapisan 100 mb. Pengolahan dilakukan dalam 3 domain seperti pada gambar 1.



Gambar 1.Domain Penelitian

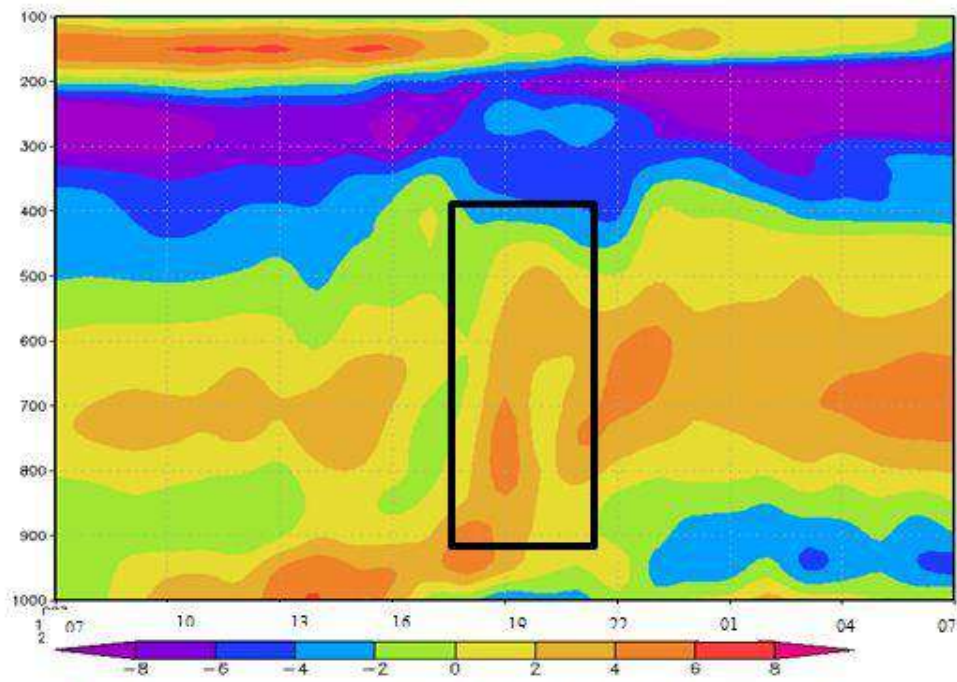
4. HASIL DAN PEMBAHASAN

Pada subbab ini penulis akan memulai dengan melakukan analisis dari beberapa kasus hujan lebat yang terjadi di Padang serta mengaitkannya dengan angin laut yang sering terjadi di pesisir Barat Sumatera. Setelah dilakukan analisis, penulis akan melakukan pembahasan berapa besar pengaruh dari angin laut terhadap pembentukan hujan di pesisir barat Padang.

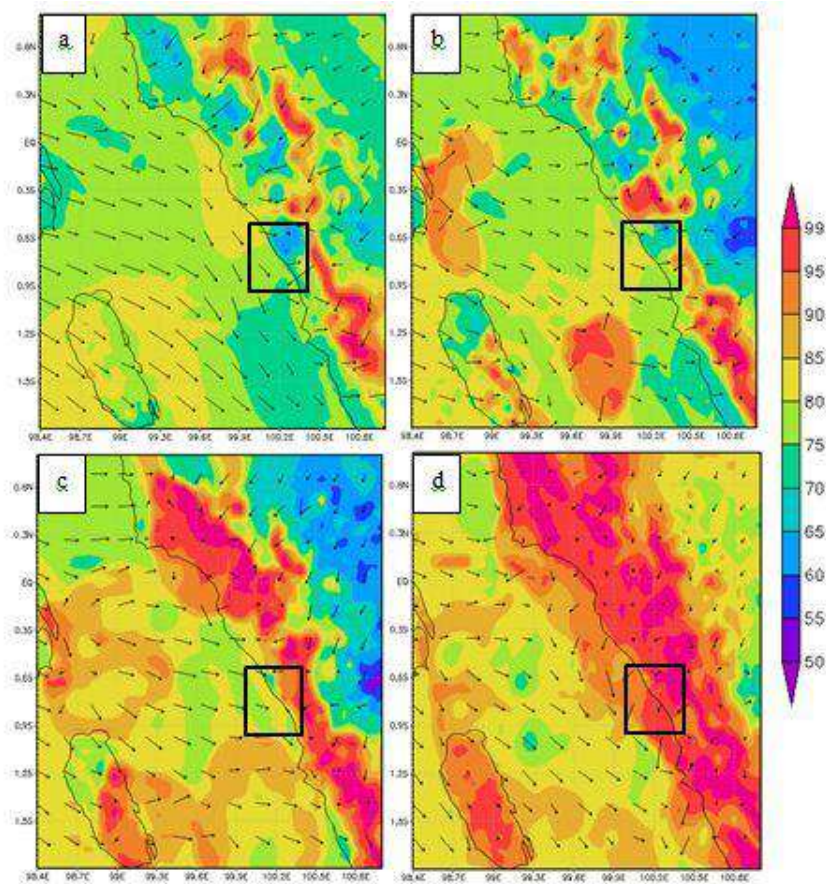
4.1. Kejadian Hujan Tanggal 17 Januari 2015

Tabel 1. Curah Hujan Tanggal 17 Januari 2015 (dalam mm)

Jam (WIB)							
07-10	10-13	13-16	16-19	19-22	22-01	01-04	04-07
0	0	0	11.2	36.5	0.5	0	0



Gambar 2. Diagram Hovmoler Angin Zonal di Wilayah Padang pada tanggal 17 Januari 2015 (dalam satuan m/s) (waktu dalam WIB)



Gambar 3. Angin dan RH pada lapisan permukaan a) 10.00 WIB, b) 13.00 WIB, c) 16.00 WIB dan d) 19.00 WIB

Berdasarkan Tabel 1 dan laporan synop tanggal 17 Januari 2015, dapat dilihat hujan dimulai pada jam 16.00 dengan intensitas 11.2 mm, lalu hujan terus berlanjut hingga pukul 22.00 dengan intensitas 36.5 mm dan berhenti pada pukul 22.00 dengan intensitas 0.5 mm. Pada kasus ini hujan dimulai pada sore hari dan berakhir pada malam hari.

Pada gambar 2 dapat dilihat angin laut mulai aktif di wilayah Padang jam 09.00. Angin laut terus berhembus menuju wilayah Padang dari Samudera Hindia hingga jam 22.00. Angin laut ini mencapai intensitas maksimumnya pada jam 13.00 hingga jam 16.00, kecepatan angin zonal mencapai 4 – 6 m/s, lalu mulai melemah hingga jam 19.00.

Angin laut di wilayah Padang dapat mencapai ketinggian 900 mb seperti yang terlihat pada gambar 2. Ketika hujan mulai terjadi pada pukul 19.00, angin laut yang awalnya hanya terjadi hingga lapisan 900 mb kemudian terangkat hingga mencapai lapisan 500 mb. Kemudian pada pukul 22.00 angin di permukaan berubah menjadi angin darat sedangkan angin pada ketinggian 800 hingga 500 mb berubah menjadi angin laut dengan kecepatan yang cukup tinggi.

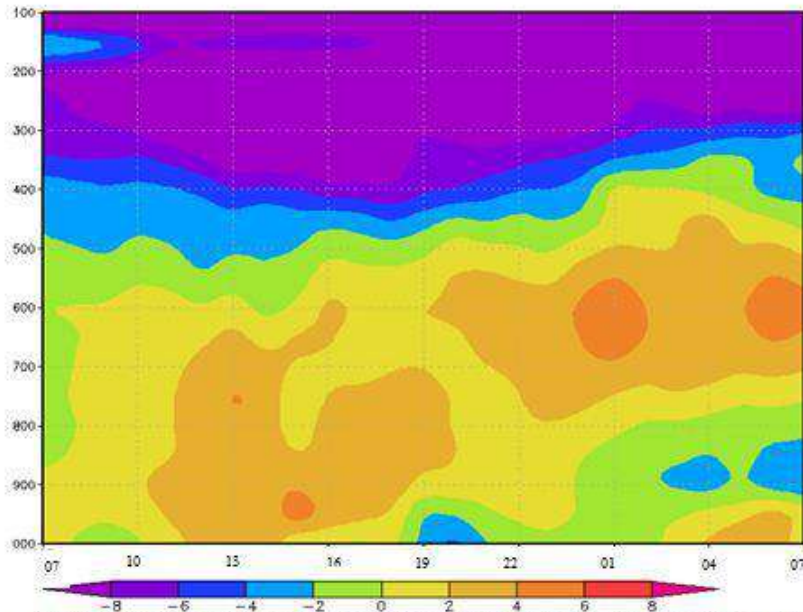
Berdasarkan gambar 3.a, dapat dilihat pada jam 09.00, ketika angin laut mulai bertiup ke darat, angin relatif lemah dan angin yang menuju darat tidak tegak lurus dengan daratan melainkan menyamping atau bersinggungan. Hal ini berbeda dengan gambar 3.b, ketika angin laut mencapai intensitas maksimumnya. Pada gambar ini, angin laut mulai menguat dan angin yang bergerak dari laut mulai tegak lurus dengan daratan. Pada gambar ini juga, mulai terlihat adanya wilayah pertemuan angin di wilayah Padang (sebelah barat Bukit Barisan). Angin yang bergerak dari barat terhalang oleh adanya Bukit Barisan (*Blocking*) dan terjadinya wilayah pertemuan angin di sebelah barat Bukit Barisan. Fenomena *blocking* yang terjadi dapat dijelaskan pada gambar 2, dimana angin laut hanya terjadi pada lapisan bawah, sehingga massa udara tidak dapat naik hingga melewati Bukit Barisan.

Pada gambar 2.c, dimana angin laut telah berlangsung selama 6 jam, terlihat bahwa angin laut membawa uap air menuju daratan. Hal ini terlihat dari penambahan RH yang terus berlangsung sejak dimulai angin laut aktif. RH yang tinggi hanya terpusat di wilayah Padang (tepatnya disebelah barat Bukit Barisan). Sedangkan pada gambar 2.d, ketika hujan mulai terjadi, angin telah berubah menjadi angin darat, meskipun relatif lemah dan RH tinggi hingga mencapai 100% mulai menyebar di sekitar wilayah Padang.

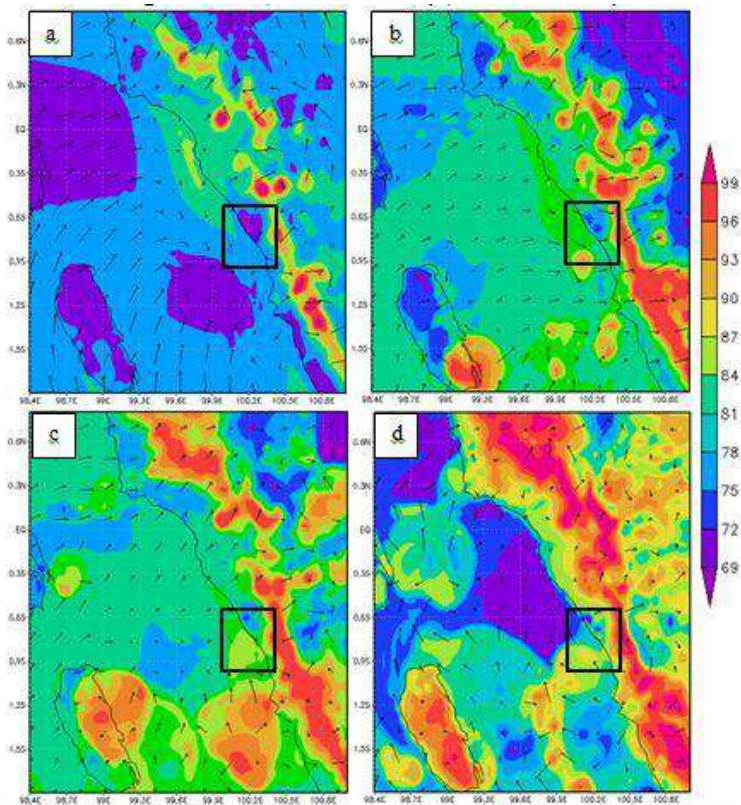
4.2. Kejadian Hujan 2 Agustus 2015

Tabel 2. Curah Hujan Tanggal 02 Agustus 2015 (dalam mm)

Jam (WIB)							
07-10	10-13	13-16	16-19	19-22	22-01	01-04	04-07
0	0	0	33	131	2	0	0



Gambar 4. Diagram Hovmoler Angin Zonal di Wilayah Padang pada tanggal 2 Agustus 2015 (dalam satuan m/s) (waktu dalam WIB)



Gambar 5. Angin dan RH pada lapisan permukaan a) 10.00 WIB, b) 19.00 WIB, c) 16.00 WIB dan d) 19.00 WIB

Berdasarkan Tabel 2, hujan mulai terjadi pada pukul 16.00 WIB, pada jam 19.00-22.00 hujan yang terjadi tergolong ekstrem dengan intensitas mencapai 131 mm, hujan mulai reda pada pukul 22.00 dengan intensitas hujan 2 mm. Pada kasus ini hujan dimulai pada sore hari dan berakhir sebelum malam hari.

Dari gambar 4. dapat dilihat bahwa angin laut mulai aktif dari pukul 10.00. Angin laut terus berhembus hingga pukul 19.00, lalu berubah menjadi angin darat dan pada jam 02.00 angin berubah kembali menjadi angin laut. Angin laut mencapai intensitas maksimumnya pada pukul 11 hingga 15 dengan kecepatan angin mencapai 4 m/s, lebih rendah dibandingkan kasus sebelumnya yang mencapai 6 m/s.

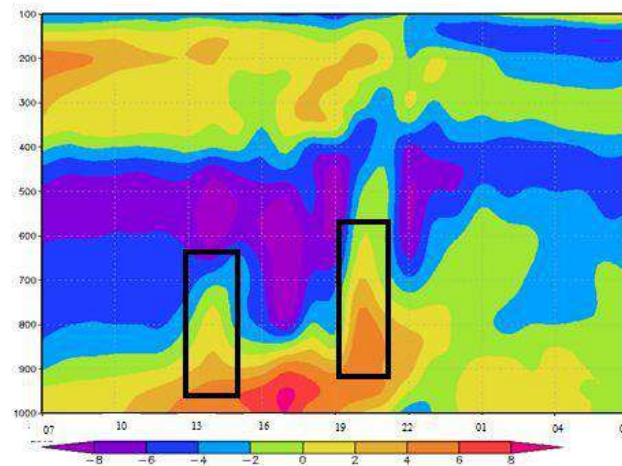
Pada kasus ini angin laut mencapai ketinggian 500 mb, cukup tinggi berbeda dengan kasus sebelumnya yang relatif lebih rendah. Pada pukul 12.00 – 14.00 angin laut di permukaan cukup tinggi dengan kecepatan mencapai 4 m/s. Kecepatan angin ini hampir sama dimulai dari permukaan hingga ketinggian 650 mb. Pada pukul 15.00 hingga 18.00, sebelum terjadinya hujan angin laut mulai melemah di permukaan tetapi di lapisan atas angin laut tetap kuat hingga lapisan 700 mb. Sesaat saat terjadinya hujan pada pukul 18.00 WIB angin mulai berubah menjadi angin darat di permukaan sedangkan di lapisan atas angin tetap angin laut. Berdasarkan gambar 5.a, dapat dilihat pada jam 10, ketika angin laut mulai bertiup ke darat, angin relatif lemah dan angin yang menuju darat tidak semuanya tegak lurus menuju daratan melainkan ada yang bersinggungan ataupun sejajar dengan daratan. Hal ini berbeda dengan gambar 5.b, ketika angin laut mencapai intensitas maksimumnya. Pada gambar ini, angin laut mulai menguat dan angin yang bergerak dari laut mulai tegak lurus dengan daratan. Pada gambar ini juga tidak terlihat adanya halangan (*blocking*) oleh bukit barisan seperti yang ditunjukkan pada kasus sebelumnya. Pada kasus ini angin terus bergerak hingga melewati Bukit Barisan..

Pada gambar 5.c, dimana angin laut telah berlangsung selama 6 jam, terlihat bahwa angin laut membawa uap air menuju daratan. Hal ini terlihat dari penambahan RH yang terus berlangsung sejak dimulai angin laut aktif. RH yang tinggi hanya terpusat di wilayah Padang (tepatnya disebelah barat Bukit Barisan). Sedangkan pada gambar 5.d, ketika hujan mulai terjadi, angin telah berubah menjadi angin darat, meskipun relatif lemah dan RH tinggi hingga mencapai 100% mulai menyebar di sekitar wilayah Padang. Tidak adanya *blocking* pada kasus ini bisa disebabkan karena angin zonal yang berasal dari laut tidak hanya terjadi pada lapisan bawah tetapi hingga mencapai lapisan udara atas (lihat gambar 4). Angin yang mencapai lapisan udara atas, menyebabkan angin dapat terus bergerak hingga melewati Bukit Barisan. Selain itu, hal inilah yang menyebabkan angin cukup besar dan energi kinetik yang cukup untuk mengatasi gaya apung dari Bukit Barisan.

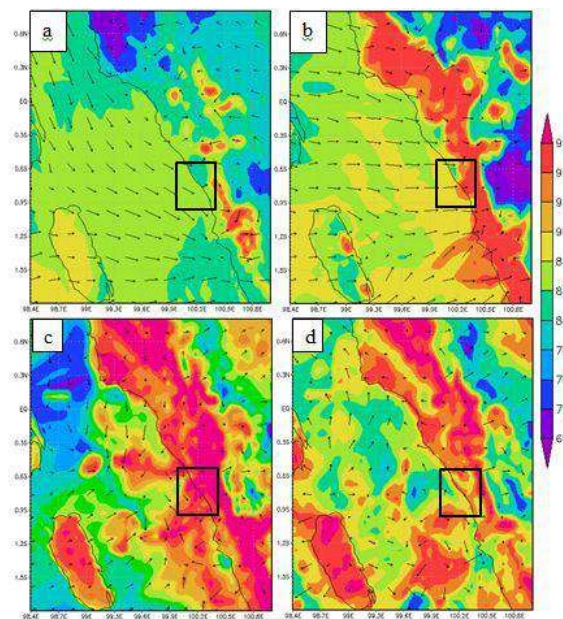
4.3. Kejadian Hujan Tanggal 15 November 2015

Tabel 3. Curah hujan pada 15 November 2015 (dalam mm)

Jam (WIB)							
07-10	10-13	13-16	16-19	19-22	22-01	01-04	04-07
0	0	42	30	0.2	0.2	68	12



Gambar 6. Diagram Hovmöller Angin Zonal di Wilayah Padang pada tanggal 15 November 2015 (dalam satuan m/s) (waktu dalam UTC)



Gambar 7. Angin dan RH pada lapisan permukaan a) 10.00 WIB, b) 13.00 WIB, c) 22.00 WIB dan d) 04.00 WIB

Pada kasus ini hujan yang terjadi berbeda dengan hujan-hujan yang terjadi pada kasus sebelumnya. Pada kasus sebelumnya hujan yang terjadi pada sore hingga malam hari pada kasus ini hujan yang terjadi dimulai dari siang hingga berakhir pada pagi hari. Hujan yang terjadi pada kasus ini berlangsung lama dengan intensitas yang lebat, hanya terjadi jeda di antara jam 19.00 hingga 01.00. Intensitas hujan yang terjadi pada hari ini mencapai 150 mm. Curah hujan pada 15 November 2015 dapat dilihat pada tabel 3.

Berdasarkan gambar 6, angin laut mulai aktif dari pagi hari pukul 07.00 dan terus bertahan hingga malam hari pukul 21.00 berubah menjadi angin darat. Kecepatan angin laut mulai meningkat saat mendekati tengah hari pada pukul 10 secara perlahan, lalu mencapai intensitas maksimumnya dari pukul 16 hingga 18. Kecepatan angin zonal maksimum dari kasus ini mencapai lebih dari 8 m/s. Pada saat terjadinya hujan pukul 13 kecepatan angin zonal cukup tinggi hingga mencapai 6 m/s. Kecepatan angin ini lebih besar jika dibandingkan dengan kasus-kasus sebelumnya.

Pada kasus ini ketinggian angin laut cukup rendah, paling rendah dibandingkan dua kasus lainnya yaitu 900 mb. Angin laut yang rendah ini lah yang menjadikan angin laut pada kasus ini memiliki kecepatan lebih tinggi jika dibandingkan dengan dua kasus lainnya. Selain itu, pada kasus ini terdapat batasan yang jelas antara angin zonal positif dan angin zonal negatif. Angin bergerak dari laut ke darat pada lapisan permukaan hingga lapisan 900 mb dan angin kembali bergerak ke laut dari lapisan 800 mb hingga lapisan 450 mb. Pada kasus ini juga terdapat tonjolan (kenaikan) angin laut pada jam 13.00 sebelum terjadinya hujan hingga kelapisan 800 mb dan tonjolan (kenaikan) angin laut yang terjadi pada jam 21.00 hingga mencapai lapisan 600 mb.

Berdasarkan gambar 7. a. dapat dilihat bahwa dimulai dari pagi hari angin laut sudah bertiup cukup kencang dan tegak lurus dengan daratan. Pada pagi hari ini, kelembaban di sekitar wilayah Padang relatif rendah. Pada gambar 7.b, saat hujan mulai terjadi, angin laut bertiup lebih cepat melewati pulau Mentawai dan langsung menuju Padang. Pada gambar ini, terlihat kelembaban yang bertambah cukup besar jika dibandingkan dengan kelembaban pada jam 10.00 WIB. Hal ini memperlihatkan bahwa adanya uap air yang terbawa saat adanya angin laut pada kasus tersebut.

Pada gambar 7.c, ketika hujan berhenti, meskipun kelembaban relatif tinggi hal ini tidak membuat hujan terus berlanjut. Hal ini disebabkan karena angin yang bergerak dari darat ke laut membuat tidak adanya tambahan pasokan uap air ke wilayah Padang. Berbeda dengan gambar 7.d, pada gambar ini terlihat adanya angin laut yang kembali mengarah ke daratan menjadi pembawa uap air yang menyebabkan hujan pada jam tersebut.

5. KESIMPULAN

Angin laut menjadi faktor utama dalam pembentukan hujan di wilayah Padang Sumatra Barat. Dalam beberapa kasus hujan lebat yang telah diamati oleh penulis, angin laut menjadi pemicu utamanya. Angin laut yang terjadi mulai dari siang hari pukul 10 WIB hingga sore hari pukul 19 WIB membawa uap air dari laut menuju daratan.

Dalam prosesnya uap air akan berkembang menjadi awan dan akan menjadi hujan dalam beberapa jam kemudian. Adanya angin yang bergerak dari timur melewati Bukit Barisan, menjadikan wilayah Padang sebagai tempat pertemuan angin (konvergensi), yang mempercepat proses terbentuknya awan. Fenomena *blocking* yang terjadi di Bukit Barisan bergantung pada angin laut yang terjadi pada level permukaan. Angin laut yang terjadi hingga level 600 mb terhindar dari fenomena *blocking* dan dapat melewati Bukit Barisan, sedangkan angin yang terjadi hanya pada level 800 mb akan terkena fenomena ini dan angin akan kembali menuju laut atau wilayah Padang.

6. KEPUSTAKAAN

- [1] Zakir, A., Sulistya, W., dan Khotimah, M., (2009), Perspektif Operasional Cuaca Tropis, BMKG, Jakarta.
- [2] Tjasyono, B., (2004), Klimatologi, ITB , Bandung

- [3] Basuki, Winarsih, I., dan Adhyani, N. L., (2009), *Analisis Periode Ulang Hujan Maksimum dengan Berbagai Metode*. BMKG, Jakarta
- [4] Buletin Stasiun Meteorologi Kelas II Tabing Padang Edisi.III/Maret 2012
- [5] Walsh, J.E. (1974). *Sea Breeze Theory and Application*. Journal of Atmospheric Science Vol.31.
- [6] Tjasyono, B.H.K., Harijono, S.W.B. (2013). *Atmosfer Equatorial*. BMKG, Jakarta
- [7] Tjasyono, B.H.K. (2008). *Sains Atmosfer*. BMKG, Jakarta
- [8] Ratag, M.A. (2006). *Dinamia Atmosfer*. BMKG, Jakarta
- [9] Panchal, N.S. (1992). *Onset Characteristics of Land/Sea Breeze Circulation and Its Effect on Meteorological Parameter at Coastal Site*. Atmosfera Vol.6
- [10] Steele, C.J., Dorling, S.R., Von Glasow, R., Bacon, J. (2013). *Idealized WRF model sensitivity simulations of sea breeze types and their effects on offshore windfields*. Atmospheric Chemistry and Physics vol. 13
- [11] Skamrock, W.C., Klemp, J.B., Dudhia, J., Gill, D.O., Barker, D.M., Duda, M.G., Huang, X.Y., Wang, W., dan Powers, J.G. (2008). *A Description of the Advanced Research WRF Version 3*. NCAR TECHNICAL NOTE: NCAR/TN-475+STR
- [12] Holton, J.S. (2004). *An Introduction to Dynamic Meteorologi*. Elsevier Academic Press: California

**PENERAPAN MEKANISME PERTAHANAN DIRI
(SELF-DEFENSE) SEBAGAI UPAYA STRATEGI
PENGURANGAN RASA TAKUT TERHADAP KEJAHATAN
(Studi Pada Kabupaten/Kota di Provinsi Lampung yang
Menduduki Peringkat Crime Rate Tertinggi)**

Teuku Fahmi

Dosen Jurusan Sosiologi FISIP Universitas Lampung
Jl. Prof. Soemantri Brojonegoro No. 1 Bandar Lampung
E-mail: teuku.fahmi@fisip.unila.ac.id Mobile +62 815 1676 721

ABSTRAK

Penelitian ini bertujuan untuk menelusuri beragam strategi pertahanan diri yang diterapkan masyarakat guna meminimalisir menjadi korban kejahatan (mengurangi risiko viktimisasi). Penelitian ini menggunakan pendekatan kuantitatif dengan tipe deskriptif dan eksplanatif. Jumlah sampel dalam penelitian ini berjumlah 384 orang yang merupakan masyarakat tinggal dan beraktivitas di kabupaten/kota di Provinsi Lampung yang memiliki crime rate tertinggi. Penelitian ini menggunakan dua teknik analisis: pertama, analisis data sekunder dengan menggunakan data statistik kriminal yang bersumber BPS Provinsi Lampung. Hasil analisis data sekunder menunjukkan bahwa Kota Bandar Lampung memiliki crime rate tertinggi jika dibandingkan dengan kabupaten/kota lainnya di Provinsi Lampung. Hasil penelitian menunjukkan dari enam strategi yang diterapkan guna mengurangi potensi kejahatan (mulai dari penggunaan jasa keamanan (satpam) hingga pemasangan peralatan antimaling/pencurian (CCTV)) secara keseluruhan berada pada tingkat "cukup efektif". Hipotesis nol diuji dengan analisis Kruskal-Wallis menyatakan bahwa enam komponen yang digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (fear reduction strategies) memiliki efek yang sama dalam upaya mencipatakan rasa aman di tengah masyarakat. Hasil perhitungan statistik Kruskal-Wallis terbukti signifikan dimana $\chi^2 (df = 5) = 53,301$; $p < 0,01$, maka dengan demikian dapat disimpulkan bahwa keenam komponen tersebut efektif digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (fear reduction strategies).

Kata kunci: mekanisme pertahanan diri, rasa takut terhadap kejahatan, crime rate

1. PENDAHULUAN

Pemberitaan mengenai maraknya tindak kejahatan yang terjadi di Provinsi Lampung, secara sadar atau tidak, mempengaruhi pola kehidupan masyarakat [1]. Implikasi dari hal tersebut diantaranya banyak aspek (kegiatan) produktif yang terhambat akibat rasa takut (*fear*) yang timbulkan. Secara psikologis hal ini dapat dipahami sebagai bagian dari reaksi tiap individu terhadap rasa takut menjadi korban kejahatan (*fear of crime*).

Situasi kekinian menunjukkan bahwa *fear of crime* menjadi isu penting yang harus disikapi oleh para pemangku kepentingan guna dicarikan solusi yang komprehensif perihal kondisi tidak teratur tersebut. Saat ini, sebagian masyarakat kerap merasakan situasi yang tidak aman ditengah aktivitas mereka. Kondisi tersebut tidak hanya terjadi pada lingkup lokal tertentu, namun secara masif, secara perlahan dan pasti, keadaan risau tersebut sudah mencapai level nasional. Beragamnya tindak kejahatan, mulai dari kejahatan jalanan seperti, pemerasan dan pencurian hingga kejahatan yang dikategorikan *extraordinary*, semisal terorisme, semakin menjadikan situasi ketidakamanan (*insecurity*) di tengah masyarakat.

Untuk konteks global, isu tentang strategi pengurangan rasa takut akan terjadinya kejahatan - atau dalam istilah Kriminologi biasa disebut "*fear of crime and fear reduction strategies*" -, telah banyak dikaji oleh para pakar. [2] banyak mengkaji dan meneliti situasi tersebut di Australia. Dalam hal ini, [2] mengidentifikasi dua faktor

dominan yang berkontribusi terhadap munculnya *fear of crime* di tengah masyarakat diantaranya, faktor gender dan keeratan sosial masyarakat.

Jauh sebelum Grabosky, di tahun 1986, The Police Foundation juga melakukan kajian perihal pengurangan rasa takut akan kejahatan di wilayah Houston dan Newark, US [3]. Dalam hal ini, The Police Foundation menerapkan beberapa strategi/program dalam upaya mengurangi *fear of crime* yang dirasakan masyarakat di Kota Houston dan Newark. Adapun program yang dijalankan diantaranya: (1) *Police-Community Newsletter*, (2) *Community Organizing Response Team*, dan (3) *Reducing the "Signs of Crime"*.

Perspektif berbeda dikemukakan [4] yang memandang bahwa terdapat keterkaitan antara berkurangnya angka kejahatan (*reduce crime*) dan ketidakteraturan (*disorder*) dengan rasa takut (*fear*) yang dialami oleh masyarakat. Dalam hal ini, pandangan [4] tersebut dikaitkan dengan efektivitas kerja kepolisian, khususnya Kepolisian Amerika Serikat.

Merujuk pada beberapa kajian mengenai *fear of crime and fear reduction strategies* tersebut, jelas terlihat bahwa rasa takut yang dialami oleh masyarakat terkait dengan tindak kejahatan, akan berkorelasi dengan pola kehidupan yang mereka jalani keseharian. [1] mengungkapkan masyarakat umumnya mengubah perilaku mereka guna menghindari menjadi korban kejahatan. Pada akhirnya, hal tersebut membatasi pilihan masyarakat dan dapat mengurangi kebebasan mereka dalam segala bentuk aktivitas.

Pada kajian terdahulu, [1] menjadikan dua lokasi di Provinsi Lampung (Kota Bandar Lampung dan Lampung Utara) untuk menelusuri tingkat risiko dan juga rasa takut (*fear*) menjadi korban kejahatan. Pada kajian tersebut, para responden hanya dimintakan mengungkapkan tingkat risiko dan *fear* yang mereka rasakan/alami, tanpa mengusut lebih lanjut strategi pengurangan rasa takut terhadap kejahatan yang mereka terapkan.

Terkait dengan hal tersebut, penelitian ini merupakan lanjutan dari kajian terdahulu dengan tema yang serupa namun berbeda pada aspek fokus kajian yang hendak didalami. Eksplorasi lebih mendalam difokuskan pada pengungkapan beragam upaya yang diterapkan oleh masyarakat sebagai strategi pengurangan rasa takut terhadap kejahatan. Premis awal yang digunakan dalam penelitian ini ialah masyarakat memiliki strategi tersendiri dalam merespon rasa takut menjadi korban kejahatan (*fear of crime*) ketika mereka melakukan aktivitas kesehariannya di ruang publik. Lebih lanjut, penelitian ini juga mengungkap beragam usaha (strategi) yang dilakukan masyarakat dalam mengurangi risiko viktimisasi. Dalam hal ini, langkah yang diambil dalam mengurangi risiko viktimisasi tersebut dapat dipahami sebagai suatu usaha guna mencegah munculnya kejahatan.

2. LANDASAN TEORI

Tinjauan tentang *Fear of Crime*

Penelitian terkait *fear of crime* banyak dikaji oleh para kriminolog ataupun para peneliti sosial lainnya yang menganggap bahwa rasa aman merupakan hak dasar bagi setiap orang. Perkembangan penelitian dengan tema tersebut pada awalnya meningkat pesat pada dekade 1960-an, terutama di Amerika Serikat dan Inggris. Hal ini

didukung oleh penggambaran yang dikemukakan bahwa pada dekade tersebut ketegangan rasial, kerusuhan, dan kekerasan perkotaan meningkat [5].

Bila ditilik lebih lanjut, penelitian kriminologi pada awalnya difokuskan pada perilaku pelaku kejahatan, bukan pada perilaku korban atau pencegahan kejahatan. Namun pada gilirannya, seiring sejalan pada situasi tertentu timbul kesadaran dan kehati-hatian khalayak ramai tentang perilaku kejahatan. Beberapa kalangan berpendapat bahwa timbulnya kesadaran tersebut merupakan sesuatu yang wajar, namun demikian bila muncul kesadaran yang berlebihan akan memiliki efek negatif bahkan sampai dengan menghasilkan kerugian yang signifikan dalam kesejahteraan pribadi. Dalam hal ini [5] mengungkapkan gambaran atau temuan yang dilakukan oleh Moore dan Trojanowicz pada tahun 1988 yakni banyak orang yang menginvestasikan waktu dan uang dalam tindakan defensif untuk mengurangi kerentanan mereka menjadi korban kejahatan [5]. Pada umumnya mereka tinggal di dalam rumah dan menghindari kegiatan lebih dari yang mereka ingin lakukan.

Hasil kajian literatur yang dilakukan [6] mengetengahkan tampilan penelitian terdahulu mengenai efek negatif dari *fear of crime*, yakni: persoalan yang menjauhkan dari kualitas hidup sehingga berimplikasi negatif terhadap kehidupan sosial dan kesejahteraan ekonomi; bahwa *fear of crime* dilihat sebagai permasalahan dalam diri individu, oleh karenanya perspektif tersebut dapat menentukan gaya hidup masyarakat, dan mengurangi penggunaan ruang publik dan fasilitas umum. Sementara kasus anak-anak, terdapat kecenderungan bila orang tua menjadi terlalu protektif, hal ini berpotensi untuk mengurangi kemampuan mereka untuk menjadi orang dewasa yang kompeten.

Tinjauan tentang *Fear Reduction Strategies*

[7] menyatakan bahwa banyak dari masyarakat mencoba untuk mengurangi risiko viktimisasi. Mereka telah melakukan sesuatu untuk merespon atau menanggapi kejahatan atau *fear of crime*. Pernyataan tersebut berdasar pada hasil analisis penelitian sebelumnya yang mengungkapkan bahwa:

"the proportions of respondents who had "limited or changed" their activities in some way because of crime ranged from 27 to 56% among 13 cities in the National Crime Survey." (Terjemahan bebas: proporsi responden yang telah "dibatasi atau diubah" kegiatan mereka dalam beberapa cara karena kejahatan berkisar pada 27-56% di antara 13 kota di National Crime Survey)

Penelitian lainnya juga telah membahas berbagai tanggapan khusus yang dilakukan orang sebagai reaksi atas kejahatan. Dalam hal ini [7] menggambarkan hasil kajian yang dilakukan DuBow dan rekannya tentang reaksi perilaku individu terhadap kejahatan, yakni digambarkan kedalam lima kategori berikut:

1. Penghindaran (*avoidance*): "tindakan yang diambil untuk mengurangi paparan kejahatan dengan menjauhkan diri sendiri atau meningkatkan jarak dari situasi dimana risiko viktimisasi atas suatu kejahatan diyakini tinggi."
2. Perilaku yang bersifat melindungi (*protective behavior*): perilaku yang "berusaha untuk meningkatkan ketahanan terhadap korban." Dua jenis diidentifikasi: (1) proteksi pada rumah "tindakan yang bertujuan untuk membuat rumah lebih terlindungi apakah itu melibatkan pembelian perangkat atau hanya

menggunakan perangkat yang ada"; (2) proteksi pribadi "tindakan yang diambil di luar rumah, selain menghindari, untuk mengurangi kerentanan ketika menghadapi situasi yang mengancam."

3. Asuransi perilaku (*insurance behavior*): perilaku yang "berusaha untuk meminimalkan biaya korban atau mengubah konsekuensi dari viktimisasi"
4. Perilaku yang bersifat komunikatif (*communicative behavior*): "berbagi informasi dan emosi yang terkait dengan kejahatan kepada orang lain."
5. Perilaku yang bersifat partisipasi (*participation behavior*): "tindakan secara bersama dengan orang lain yang termotivasi oleh kejahatan tertentu atau dengan kejahatan pada umumnya."

Sejalan dengan bahasan di atas, [8] mengungkapkan beberapa teknik lain sebagai upaya adaptasi dalam menanggapi *fear of crime*, diantaranya yakni mempersenjatai diri (senjata api), belajar teknik pertahanan diri (*self-defence*), hingga menginstalasi perlengkapan anti-maling/pencurian atau menggunakan anjing pengawas.

3. METODOLOGI PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif, adapun metode penelitian yang digunakan adalah deskriptif dan eksplanatif. Populasi dari penelitian ini adalah masyarakat pada satu kabupaten/kota di Provinsi Lampung yang memiliki *crime rate* tertinggi. Namun demikian, ketidaktersediaan kerangka sampel (*sampling frame*) dari populasi yang dimaksud, menjadikan desain sampel yang digunakan berdasar pada prinsip non probabilita (*non probability sampling*). Jumlah sampel dalam penelitian ini sebanyak 384 orang. Penentuan ukuran sampel dalam penelitian ini dilakukan dengan memperhatikan tiga faktor, yakni: variasi dalam populasi, tingkat kesalahan yang ditoleransi, dan tingkat kepercayaan [9]. Ketiga faktor yang menentukan besar sampel itu dapat dirangkum dalam rumus sebagai berikut:

$$N = (p \times q) \cdot \frac{Z^2}{E^2} \quad (1)$$

Dimana $(p \times q)$ adalah variasi populasi, Z adalah ukuran tingkat kepercayaan dan E adalah *sampling error*/kesalahan yang dapat ditoleransi. Dalam penelitian ini, taksiran tentang keragaman populasi yang digunakan adalah proporsi sebesar 50% : 50%. Ini artinya peneliti mengasumsikan populasi pada kategori yang ditetapkan adalah heterogen, keragaman populasi yang diteliti yakni, terdapat perbedaan pendapat antara jenis kelamin perempuan dan laki-laki dalam mengenali situasi *fear of crime* & *fear reduction strategies*. Ukuran tingkat kepercayaan yang digunakan sebesar 95%, pada tingkat kepercayaan 95% skor z adalah 1,96. Adapun ukuran kesalahan yang dapat ditoleransi tidak lebih $\pm 5\%$, dengan demikian $E = 0,05$. Jadi, sampel dalam penelitian ini adalah:

$$N = (0,50 \times 0,50) \cdot \frac{1,96^2}{0,10^2} \quad (2)$$

$$N = 384,16$$

Terkait dengan itu, jumlah responden berdasarkan jenis kelamin dilakukan secara berimbang. Perlu diketahui, terkait dengan sampel dalam penelitian ini kemungkinan besar tidak mewakili populasi, sehingga generalisasi yang dapat dilakukan oleh peneliti akan terbatas. Cara ini juga cenderung memiliki bias yang tinggi karena peneliti menentukan sendiri responden yang terpilih secara acak, biasanya penentuan tersebut dilakukan dengan

subjektif. Lebih lanjut, bila dilihat berdasarkan dimensi waktu, merujuk pendapat [10], pendekatan yang digunakan dalam penelitian yakni *cross-sectional research*, dimana pengumpulan data yang dilakukan hanya pada kurun satu waktu tertentu (*one-time snapshot approach*). Uji instrumen (kuesioner) dilakukan terhadap 40 responden. Merujuk hasil pengujian validitas yang dilakukan pada tiap variabel, secara keseluruhan, item pernyataan yang disodorkan dikategorikan valid, di mana tingkat signifikansi (α) tiap itemnya lebih kecil dari 0,05. Adapun hasil uji reliabilitas pada tiap variabel dalam instrumen penelitian ini menunjukkan bahwa nilai cronbach alpha pada variabel yang diteliti lebih dari 0,6 (hasil perhitungan nilai cronbach alpha sebesar 0,740).

4. HASIL DAN PEMBAHASAN

Pengukuran Kriminalitas berdasarkan Statistik Kriminal Resmi

Pada analisis ini, trend kejahatan diarahkan pada kecenderungan pertumbuhan dan penurunan angka kejahatan yang didasari pada data statistik kriminal resmi Kepolisian Daerah Lampung Tahun 2014 [11]. Untuk mengukur trend kejahatan digunakan rumusan yang dikemukakan oleh Larry Siegel, yaitu dengan mengetahui angka perimbangan kejahatan atau *Crime rate*, yakni jumlah kejahatan dibandingkan dengan jumlah penduduk, atau nilai rata-rata kejahatan per 10.000 penduduk [12].

$$CrimeRate = \frac{Angkakejahatanyangdilaporkan}{Jumlahtotalpenduduk} \times 10.000 \quad (3)$$

Secara keseluruhan, untuk mengetahui *Total Crime Rate* dilakukan perhitungan sebagai berikut, yakni *Crime Total* dalam satu tahun dibagi dengan jumlah penduduk pada tahun yang sama dan dikalikan per 10.000 penduduk. pada Tabel 1 disajikan sebaran jenis kasus laporan kejahatan di tahun 2014 pada 11 wilayah hukum kepolisian resort di Provinsi Lampung.

Perlu diketahui bahwa data statistik kriminal menurut kepolisian tidak dapat mewakili jumlah kejahatan yang ada secara keseluruhan. Tidak semua peristiwa kejahatan dicatat oleh polisi. Peristiwa kejahatan yang tidak diketahui oleh polisi yang diperkirakan jumlahnya sangat banyak tidak pernah tercatat dalam statistik kriminal polisi. Data kriminalitas yang tidak diketahui oleh polisi ini disebut sebagai angka gelap (*dark number*) kejahatan [13].

Tabel 1. Jenis Kasus Laporan Kejahatan Tahun 2014 pada 10 Wilayah Hukum Kepolisian Resort di Provinsi Lampung

No	Wilayah Hukum Kepolisian Resort	Jenis Kasus Laporan Kejahatan							Jumlah
		Pem- bunu- han	Ani- rat	Cu-ras	Cu-rat	Cu- ran- mor	Per- ko- saan	Pe- me-ra- san	
1	Lampung Barat ^{*)}	0	0	5	63	17	0	1	86
2	Tanggamus ^{**)}	2	4	42	87	40	8	3	186
3	Lampung Selatan	5	5	115	176	71	6	6	384

4	Lampung Timur	23	27	519	1014	223	46	25	1877
5	Lampung Tengah	7	40	68	105	277	7	10	514
6	Lampung Utara	37	99	308	927	1142	75	40	2628
7	Way Kanan	2	8	19	101	72	6	0	208
8	Tulang Bawang ^{***})	3	1	88	78	29	7	9	215
9	Mesuji	5	7	25	54	7	1	2	101
10	Bandar Lampung	26	491	711	3361	2556	78	67	7290
11	Metro	1	0	12	33	100	5	3	154
Jumlah		111	682	1912	5999	4534	239	166	

Keterangan:

^{*)} Mencakup Kabupaten Lampung Barat dan Pesisir Barat

^{**}) Mencakup Kabupaten Tanggamus dan Pringsewu

^{***}) Mencakup Kabupaten Lampung Selatan dan Pesawaran

^{****}) Mencakup Kabupaten Tulang Bawang dan Tulang Bawang Barat

Sumber: Lampung dalam Angka 2015

Untuk konteks ini, ada beberapa kriteria yang perlu diperhatikan dalam penggunaan statistik kriminal agar tidak menyesatkan, diantaranya: (1) menghindari pernyataan total kejahatan sebagai tolok ukur tingkat kriminalitas; (2) dalam mengukur kriminalitas akan lebih baik dikelompokkan menurut klasifikasi kejahatan yang masing-masing klasifikasi mempunyai kesamaan ciri; (3) fluktuasi kejahatan harus diperhitungkan dengan fluktuasi populasi penduduk (*crime rate*), dan; (4) dalam mengukur fluktuasi kejahatan “polisi” sering mempergunakan “angka indeks kejahatan” dan angka indeks kejahatan inilah yang digunakan sebagai tolok ukur fluktuasi kejahatan [13].

Berangkat dari beberapa kriteria penggunaan statistik kriminal tersebut, maka analisis *crime rate* pada 11 wilayah hukum kepolisian resort di Provinsi Lampung akan dilakukan pada tujuh jenis kasus dengan mempertimbangkan seriusitas kejahatan yang dilaporkan. Secara rinci *Crime Rate* data tahun 2014 pada 11 wilayah hukum kepolisian resort di Provinsi Lampung dapat diamati pada Tabel 2.

Tabel 2 menunjukkan bahwa secara umum *crime rate* data tahun 2014 pada 11 wilayah hukum kepolisian resort di Provinsi Lampung menunjukkan variasi nilai yang cukup berbeda. Bila ditelusuri berdasarkan besaran angka *crime rate*, terdapat tiga wilayah dengan nilai tertinggi yakni Kota Bandar Lampung (62,46), Kabupaten Lampung Utara (29,90), dan Kabupaten Lampung Timur (16,97). Dapat dinyatakan bahwa tiga wilayah tersebut merupakan daerah dengan angka kriminalitas tertinggi di Provinsi Lampung.

Hasil analisis data di atas menunjukkan bahwa wilayah Bandar Lampung menempati posisi tertinggi dalam nilai *crime rate* di Provinsi Lampung. Hal tersebut mengindikasikan bahwa wilayah ini (untuk konteks Provinsi Lampung) memiliki angka kriminalitas yang terjadi tergolong tinggi. Pada tahap selanjutnya, wilayah Kota Bandar Lampung ditetapkan menjadi lokasi pengumpulan data kuesioner guna memperoleh gambaran pola

penerapan *fear reduction strategies* pada wilayah tersebut.

Tabel 2. *Crime Rate (per 10.000 penduduk) Tahun 2014 pada 11 Wilayah Hukum Kepolisian Resort di Provinsi Lampung*

No.	Wilayah Hukum Kepolisian Resort	Jumlah Penduduk	Jumlah Laporan Kasus Kejahatan ^{*)}	Crime Rate
1	Lampung Barat ^{**)}	450991	86	1.91
2	Tanggamus ^{***)}	1088611	186	1.71
3	Lampung Selatan ^{****)}	1803479	384	2.13
4	Lampung Timur	1105990	1877	16.97
5	Lampung Tengah	1449851	514	3.55
6	Lampung Utara	878874	2628	29.90
7	Way Kanan	472815	208	4.40
8	Tulang Bawang ^{*****)}	666807	215	3.22
9	Mesuji	302730	101	3.34
10	Bandar Lampung	1167101	7290	62.46
11	Metro	161830	154	9.52

Keterangan:

^{*)} Mencakup tujuh jenis kejahatan yang dilaporkan yakni: pembunuhan, penganiayaan berat, pencurian dengan kekerasan, pencurian dengan pemberatan, pencurian kendaraan bermotor, perkosaan, dan pemerasan.

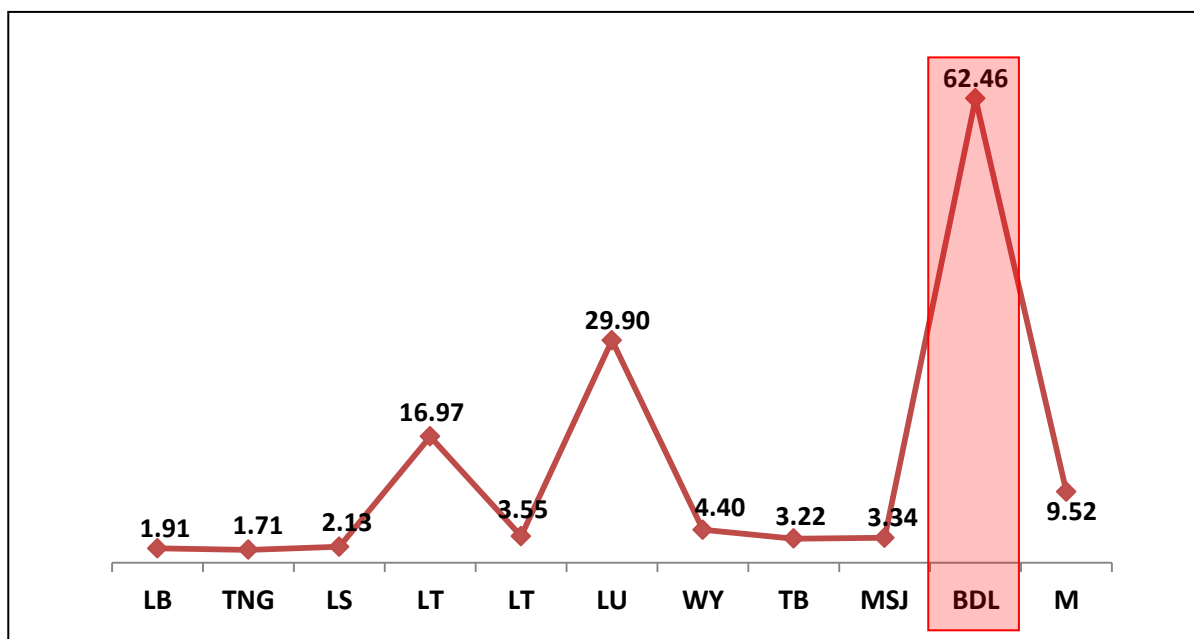
^{**)} Mencakup Kabupaten Lampung Barat dan Pesisir Barat

^{***)} Mencakup Kabupaten Tanggamus dan Pringsewu

^{****)} Mencakup Kabupaten Lampung Selatan dan Pesawaran

^{*****)} Mencakup Kabupaten Tulang Bawang dan Tulang Bawang Barat

Sumber: Olahan data sekunder, 2016



Sumber: Olahan data sekunder, 2016

Gambar 1. *Crime Rate Tahun 2014 pada 10 Wilayah Hukum Kepolisian Resort di Provinsi Lampung*

Menarik untuk dicermati angka *crime rate* Kota Bandar Lampung di tahun 2016 dengan merujuk data statistik kriminal resmi Kepolisian Daerah Lampung Tahun 2014. Bila dibandingkan dengan analisis serupa ditahun 2014 [1], terjadi kenaikan yang signifikan antara *crime rate* pada 2014 (data statistik kriminal resmi Kepolisian Daerah Lampung Tahun 2012) dengan *crime rate* pada 2016 (data statistik kriminal resmi Kepolisian Daerah Lampung Tahun 2014). Pada *crime rate* Kota Bandar Lampung 2014 hanya berkisar 11,42 sedangkan *crime rate* pada 2016 menyentuh angka 62,46.

Kenaikan signifikan pada angka *crime rate* tersebut dapat dipahami sebagai dinamika sosial masyarakat Kota Bandar Lampung. Meskipun demikian, lonjakan angka *crime rate* ini mesti menjadi perhatian serius para pemangku kepentingan (*stakeholders*) guna dijadikan sebagai penentuan kebijakan hal yang terkait dengan usaha penciptaan keamanan di wilayah Kota Bandar Lampung. Khusus untuk pihak kepolisian, tingginya angka *crime rate* memiliki kecenderungan dengan peningkatan rasa takut (*fear*) yang dialami oleh masyarakat sebagaimana yang dikemukakan oleh [14]. Oleh karenanya dibutuhkan beragam strategi yang perlu dilakukan oleh pihak kepolisian dalam mengurangi tindak kejahatan di tengah masyarakat.

Gambaran Umum Subjek Penelitian

Hasil analisis data sekunder pada subbab sebelumnya menunjukkan bahwa *crime rate* tertinggi berada di wilayah Kota Bandar Lampung. Peringkat *crime rate* tertinggi inilah yang dijadikan sebagai acuan dalam penentuan lokasi wilayah kabupaten/kota untuk dilakukan pengumpulan data primer guna menjawab permasalahan penelitian ini. Lokasi pengumpulan mencakup keseluruhan wilayah (kecamatan) di Kota Bandar Lampung. Hal ini dilakukan untuk mendapatkan gambaran yang nyata yang merepresentasikan realitas dari tingkat keamanan

yang dirasakan oleh warga Kota Bandar Lampung. Dapat dinyatakan, setiap wilayah (kecamatan) di Kota Bandar Lampung memiliki kekhasan tersendiri tentang apa yang disebut “rasa aman”, ada kalanya satu wilayah cenderung lebih aman bila dibandingkan dengan wilayah lainnya, begitu pula sebaliknya.

Pada penyajian karakteristik responden dipaparkan beberapa cakupan diantaranya jenis kelamin, kelompok umur, dan tingkatan pendidikan. Lebih lanjut, disajikan pula pengalaman responden perihal apakah mereka pernah menjadi korban tindakan kejahatan beserta jenis kejahatan yang dialaminya. Secara umum, proporsi responden perempuan yang terlibat dalam penelitian sedikit lebih banyak bila dibandingkan laki-laki, yakni 50,5 persen berbanding 49,5. Jika dilihat berdasarkan kelompok usia, sebagian besar responden berada pada rentang usia kurang dari 20 hingga 34 tahun (56,3 persen). Namun demikian bila diamati berdasarkan rentang usia, persentase terbesar berada pada kelompok usia kurang dari 20 tahun (sebesar 41,5 persen), lalu diikuti kelompok usia 20 - 24 tahun (sebesar 14,8 persen) dan lebih dari 49 tahun (sebanyak 9,2 persen). Persentase terkecil berada pada kelompok usia 35 - 39 tahun (5 persen).

Bila melihat karakteristik responden berdasarkan pada tingkatan pendidikan, terdapat kecenderungan yang seragam diantara kedua kategori responden. Sebagai gambaran, persentase terbesar dalam penelitian ini bila dilihat dari tingkat pendidikan terakhir berada pada tingkatan Tamat SMA, yakni untuk responden laki-laki sebesar 37 persen dan perempuan sebesar 34,9 persen. Terkait dengan hal itu, dalam penelitian ini juga terungkap responden yang hanya tamat SD. Secara keseluruhan, terdapat 1,8 persen dari total keseluruhan responden yang hanya tamat SD. Adapun rinciannya yakni laki-laki sebesar 1 persen dan perempuan sebesar 0,8 persen.

Penelitian ini juga mengungkapkan pengalaman responden menjadi korban kejahatan (viktimsasi). Hasil data lapangan cukup mengejutkan, dari jumlah keseluruhan responden yang berjumlah 384 orang, sebanyak 31,8 persen (122 orang) pernah menjadi korban kejahatan di Kota Bandar Lampung. Beragam jenis kejahatan yang pernah dialami para korban, mulai dari pencurian kendaraan bermotor, perampasan, pengrusakan barang, pencopetan, dan beragam jenis kejahatan lainnya.

Tabel 3. *Karakteristik Sosiodemografi Responden*

Karakteristik	Jenis Kelamin				Total	
	Laki-laki		Perempuan			
	<i>f</i>	%	<i>f</i>	%	<i>f</i>	%
Kelompok umur (N=358)						
<20 tahun	76	21,2	73	20,4	149	41,6
20 – 24 tahun	25	7	28	7,8	53	14,8
25 – 29 tahun	9	8,8	18	5	27	7,5
30 – 34 tahun	11	3,1	9	2,5	20	5,6
35 – 39 tahun	9	2,5	9	2,5	18	5
40 – 44 tahun	17	4,7	11	3,1	28	7,8
45 – 49 tahun	15	4,2	15	4,2	30	8,4

>49 tahun	15	4,2	18	5	33	9,2
Total	177	49,4	181	50,6	358	100
Tingkatan pendidikan (N=384)						
Tidak tamat SD	0	0	1	0,3	1	0,3
Tamat SD	4	1,0	3	0,8	7	1,8
Tamat SMP	6	1,6	9	2,3	15	3,9
Tamat SMA	134	34,9	142	37	276	71,9
Diploma/S1	44	11,5	36	9,4	80	20,8
S2/S3	2	0,5	2	0,4	4	1
Total	190	49,5	194	50,5	384	100
Pengalaman menjadi korban kejahatan di wilayah Kota Bandar Lampung (N=384)						
Ya, pernah	52	13,5	70	18,2	122	31,8
Tidak pernah	138	35,9	124	32,3	262	68,2
Total	190	49,5	194	50,5	384	100

Sumber: Olahan data sekunder, 2016

Persepsi Responden mengenai Rasa Takut terhadap Kejahatan (*Fear of Crime*) di Wilayah Kota Bandar Lampung

Subbab ini menyajikan penilaian responden mengenai rasa takut terhadap kejahatan (*fear of crime*) di wilayah Kota Bandar Lampung. Cakupan penilaian mengenai *fear fo crime* merujuk pada beberapa hal berikut:

1. Situasi kemanan Kota Bandar Lampung saat ini bila dibandingkan tahun lalu,
2. Kondisi/kategori tingkat kerawanan kejahatan Kota Bandar Lampung,
3. Pengalaman responden apakah pernah atau tidak pernah menjadi korban kejahatan,
4. Wilayah rawan kejahatan di Kota Bandar Lampung, dan
5. Durasi/rentang waktu yang dinilai bahwa Kota Bandar Lampung rawan kejahatan.

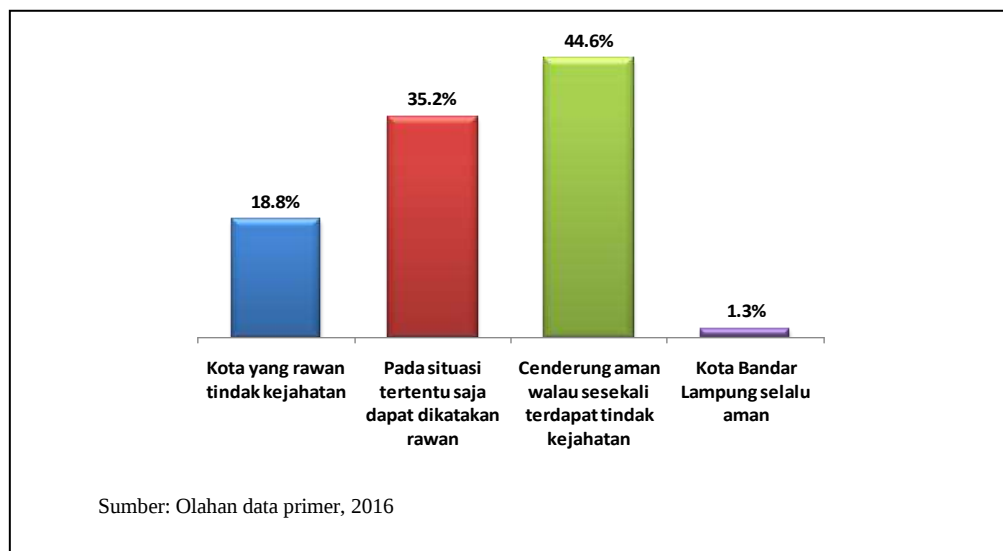
Dari penilaian telah responden berikan, untuk selanjutnya dilakukan interpretasi deskriptif pada tiap item pertanyaan sebagai indikator dari *fear of crime*. Interpretasi yang sama juga akan dilakukan untuk menelusuri lebih lanjut pada strategi yang digunakan untuk menghindarkan diri menjadi korban kejahatan atau istilah yang digunakan dalam penelitian ini *fear reduction strategies*. Dalam hal ini, digunakan enam parameter dalam mengukur *fear reduction strategies* tersebut, diantaranya yakni:

1. Penggunaan jasa keamanan (satpam)
2. Optimalisasi peran SISKAMLING
3. Optimalisasi peran Bhabinkamtibmas
4. Optimalisasi peran Babinsa
5. Menyimpan alat perlindungan sebagai alat proteksi
6. Pemasangan peralatan antimaling/pencurian (CCTV)

Secara teknis, responden dimintakan untuk memberikan penilaian efektivitas pada keenam parameter tersebut sebagai strategi yang dapat digunakan dalam mengurangi potensi menjadi korban kejahatan. Skala penilaian yang diberikan berupa skala likert mulai dari “Sangat Efektif” hingga “Tidak Efektif”. Untuk selanjutnya, hasil penilaian yang diberikan dianalisis secara inferensial guna menjawab permasalahan penelitian yang diangkat dalam kajian ini.

Terkait dengan penilaian responden tentang situasi keamanan Kota Bandar Lampung, data hasil lapangan menunjukkan bahwa responden menilai situasi keamanan saat ini “sedikit ada perbaikan” dibandingkan tahun lalu (44 persen). Persentase terbesar kedua merujuk pada kategori “tidak ada perubahan” (28,1 persen). Terdapat 17,7 persen responden yang menilai bahwa situasi keamanan saat ini “lebih buruk” bila dibandingkan tahun lalu.

Hasil data lapangan juga menunjukkan bahwa responden mengklasifikasikan Kota Bandar Lampung sebagai kota yang “cenderung aman walau sesekali terdapat tindak kejahatan” (44,8 persen). Persentase terbesar kedua merujuk pada kategori “pada situasi tertentu saja dapat dikatakan rawan” (35,2 persen). Terdapat 18,8 persen responden yang menilai bahwa Kota Bandar Lampung sebagai “kota yang rawan tindak kejahatan”. Secara rinci penilaian responden dapat diamati pada Gambar 2.



Gambar 2. Persepsi Responden tentang Klasifikasi Kondisi Kota Bandar Lampung (n = 384)

Analisis lebih lanjut dilakukan dengan mensilangkan antara klasifikasi (penggolongan) kondisi Kota Bandar Lampung berdasarkan jenis kelamin. Hal ini untuk mendapatkan gambaran nyata tentang persepsi “rasa takut” juga berkorelasi dengan jenis kelamin. Banyak literatur yang mengungkapkan bahwa perempuan cenderung memiliki penilaian jauh lebih tinggi tentang “fear” bila dibandingkan laki-laki. Hasil penelitian ini juga mendukung pernyataan sebelumnya dimana perempuan memiliki penilaian yang lebih tinggi mengenai situasi/klasifikasi kondisi rawan di Kota Bandar Lampung. Secara rinci hasil analisis silang tersebut dilihat pada Tabel 4.

Tabel 4. Persepsi Responden tentang Klasifikasi (penggolongan) Kondisi Kota Bandar Lampung berdasarkan Jenis Kelamin (N = 383)

Klasifikasi Secara Umum Kondisi Kota Bandar Lampung	Jenis Kelamin				Total	
	Laki-laki		Perempuan			
	<i>f</i>	%	<i>f</i>	%	<i>f</i>	%
Kota yang rawan tindak kejahatan	40	10,4	32	8,4	72	18,8
Pada situasi tertentu saja dapat dikatakan rawan	67	17,5	68	17,8	135	35,2
Cenderung aman walau sesekali terdapat tindak kejahatan	82	21,4	89	23,2	171	44,6
Kota Bandar Lampung selalu aman	1	0,3	4	1,0	5	1,3
Total	190	49,6	193	50,4	383	100

Sumber: Olahan data primer, 2016

Adapun untuk pengukuran enam parameter *fear reduction strategies* mengacu pada penilaian/jawaban yang diberikan oleh responden dengan menggunakan skala likert (mulai dari “Sangat Efektif” bernilai bobot 4, “Cukup Efektif” bernilai bobot 3, “Kurang Efektif” bernilai bobot 2, dan “Tidak Efektif” bernilai bobot 1). Pada Tabel 5, terlihat distribusi penilaian responden terhadap enam item parameter yang diukur.

Tabel 5. Efektifitas Penerapan Strategi Pengurangan Rasa Takut Terhadap Kejahatan (*Fear Reduction Strategies*)

Kategori Responden	Tingkat Efektivitas			
	Tidak Efektif	Kurang Efektif	Cukup Efektif	Sangat Efektif
Penggunaan jasa keamanan (satpam)	28 (7,3%)	90 (23,4%)	197 (51,3)	69 (18%)
Optimalisasi peran SISKAMLING	24 (6,3%)	88 (22,9%)	160 (41,7%)	112 (29,2%)
Optimalisasi peran Bhabinkamtibmas	27 (7,0)	78 (20,3%)	187 (48,7%)	92 (24%)
Optimalisasi peran Babinsa	28 (7,3%)	75 (19,5%)	184 (47,9%)	97 (25,3%)
Menyimpan alat perlindungan sebagai alat proteksi	35 (9,1%)	85 (22,1%)	148 (38,5%)	116 (30,2%)
Pemasangan peralatan antimaling/ pencurian (CCTV)	31 (8,1%)	55 (14,3%)	112 (29,2%)	186 (48,4%)

Sumber: Olahan data primer, 2016

Untuk selanjutnya, analisis inferensial digunakan dengan perhitungan Kruskal–Wallis Test for Several Independent Samples. Merujuk pendapat [15] Tes Kruskal-Wallis merupakan tes nonparametrik yang digunakan

pada kelompok independen yang terdiri dari lebih dari dua kelompok. Dalam penelitian ini terdapat enam kelompok independen yakni enam parameter dalam mengukur *fear reduction strategies*.

Tabel 6. Hasil Uji Statistika Nonparametrik Kruskal–Wallis terhadap Efektifitas Penerapan Strategi Pengurangan Rasa Takut Terhadap Kejahatan (*Fear Reduction Strategies*)

Kruskal-Wallis Test			
Ranks			
Komponen		N	Mean Rank
Skor penilaian	Satpam	384	1039.99
	Siskamling	384	1147.69
	Bhabinkamtibmas	384	1115.24
	Babinsa	384	1129.01
	Alat proteksi	384	1131.91
	CCTV	384	1351.16
	Total	2304	

Test Statistics^{a,b}

	Skor penilaian
Chi-Square	53.301
df	5
Asymp. Sig.	.000

a. Kruskal Wallis Test

b. Grouping Variable: komponen

Sumber: Olahan data primer, 2016

Hipotesis nol diuji oleh analisis Kruskal-Wallis adalah bahwa enam komponen yang digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (*fear reduction strategies*) memiliki tidak efek yang sama dalam upaya menciptakan rasa aman di tengah masyarakat. Hasil perhitungan statistik Kruskal-Wallis terbukti signifikan dimana χ^2 (df = 5) = 53,301; $p < 0,01$, maka dengan demikian dapat disimpulkan bahwa keenam komponen tersebut efektif digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (*fear reduction strategies*).

5.SIMPULAN

Berdasarkan hasil temuan penelitian, beberapa kesimpulan yang dapat ditarik guna menjawab permasalahan dalam penelitian ini adalah:

1. Terkait penilaian responden tentang situasi keamanan Kota Bandar Lampung menunjukkan bahwa sebanyak 44 persen menilai situasi keamanan saat ini “sedikit ada perbaikan” dibandingkan tahun lalu (44 persen). Selain itu, responden cenderung mengklasifikasikan Kota Bandar Lampung sebagai kota yang “cenderung aman walau sesekali terdapat tindak kejahatan” (44,8 persen).
2. Hasil data lapangan lainnya yakni dari jumlah keseluruhan responden yang berjumlah 384 orang, sebanyak 31, 8 persen (122 orang) pernah menjadi korban kejahatan di Kota Bandar Lampung. Beragam jenis kejahatan yang pernah dialami para korban, mulai dari pencurian kendaraan bermotor, perampasan, pengrusakan barang, pencopetan, dan beragam jenis kejahatan lainnya.
3. Hipotesis nol diuji oleh analisis Kruskal-Wallis adalah bahwa enam komponen yang digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (*fear reduction strategies*) memiliki tidak efek yang sama dalam upaya menciptakan rasa aman di tengah masyarakat. Hasil perhitungan statistik Kruskal-Wallis terbukti signifikan dimana χ^2 ($df = 5$) = 53,301; $p < 0,01$, maka dengan demikian dapat disimpulkan bahwa keenam komponen tersebut efektif digunakan sebagai strategi pengurangan rasa takut terhadap kejahatan (*fear reduction strategies*).

KEPUSTAKAAN

- [1] Fahmi, T. (2014). Perbedaan tingkat perceived risk, fear of crime, dan mekanisme coping pada masyarakat yang beraktivitas di wilayah rawan tindak kejahatan. *Sosiologi: Jurnal Ilmiah Kajian Ilmu Sosial dan Budaya*, 16(2).
- [2] Grabosky, P. N. (1995). *Fear of crime and fear reduction strategies*. Australian Institute of Criminology.
- [3] Pate, A.M., Wycoff, M.A., Skogan, W.G., & Sherman, L.W. (1986). *Reducing fear of crime in Houston and Newark, a summary report*. Washington: Police Foundation.
- [4] Weisburd, D. & Eck, J.E. (2004). What can police do to reduce crime, disorder, and fear?. *Annals of the American Academy of Political and Social Science*, Vol. 593, To Better Serve and Protect: Improving Police Practices (May, 2004), pp. 42-65.
- [5] Grohe, B. R. (2007). *Perceptions of crime, fear of crime, and defensible space in Fort Worth neighborhoods*.
- [6] National Campaign Against Violence and Crime (1998). *Fear of crime - audit of the literature and community programs*. Criminal Research Council [accessed March 5, 2016]. Available from www.criminologyresearchcouncil.gov.au/reports/1998-foc1.pdf
- [7] Garofalo, James. (1981). The fear of crime: causes and consequences. *The Journal of Criminal Law and Criminology* (1973), Vol. 72, No. 2. pp.839-857.
- [8] Doran, B.J. & Burgess M.B. (2011). *Putting fear of crime on the map investigating perceptions of crime using geographic information systems*. Springer New York.

- [9] Eriyanto. (1999). *Metodologi polling memberdayakan suara rakyat*. Bandung: PT Remaja Rosdakarya.
- [10] Neuman, W.L. (2007). *Basic of social research, qualitative and quantitative approaches, second edition*. Boston: Pearson Education, Inc.
- [11] BPS Provinsi Lampung. 2016. *Lampung dalam angka 2015*. Bandar Lampung: Badan Pusat Statistik Kota Bandar Lampung.
- [12] Siegel, Larry J. (2005). *Criminology (9th edition)*. California: Wadsworth Publishing.
- [13] Mustofa, M. (2007). *Kriminologi kajian sosiologi terhadap kriminalitas perilaku menyimpang dan pelanggaran hukum*. Depok: FISIP UI Press.
- [14] Cordner, G., & Melekian, B. K. (2010). *Reducing fear of crime strategies for police*. U.S. Department of Justice Office of Community Oriented Policing Services Cover.
- [15] Ho, R. (2006). *Handbook of univariate and multivariate data analysis and interpretation with SPSS*. CRC Press.

TINGKAT KETAHANAN INDIVIDU MAHASISWA UNILA PADA ASPEK SOFT SKILL

(Studi pada mahasiswa Fakultas Ilmu Sosial dan Politik, dan Fakultas Kedokteran, Unila)

Pitojo Budiono, Feni Rosalia dan Lilih Muflihah

Jurusan Ilmu Pemerintahan Fisip Unila

ptjbudiono@gmail.com

ABSTRAK

Ketahanan Nasional akan tangguh jika di topang dengan Ketahanan Daerah yang tangguh, dan Ketahanan Daerah akan tangguh jika di dukung dengan Ketahanan Individu. Ketahanan individu pada mahasiswa sangat dibutuhkan dan mutlak, karena ketahanan individu pada mahasiswa identik dengan soft skill. Kemampuan akademik (hard skills) yang tinggi tetap harus memperhatikan kecakapan dalam hal nilai-nilai yang melekat pada seseorang atau sering dikenal dengan aspek soft skills. Kemampuan ini dapat disebut juga dengan kemampuan non teknis yang tentunya memiliki peran tidak kalah pentingnya dengan kemampuan akademik. Terdapat empat variabel yang diujicobakan yakni a) integritas; b) Kemampuan menghadapi perubahan; c) Kepemimpinan; d) kemampuan berkomunikasi. Sedangkan metode yang digunakan adalah deskriptif kuantitatif dengan menggunakan angket tertutup dan menggunakan skala likert untuk menemukan sikap persetujuan yang kemudian di sajikan secara tabulasi. Hasil yang di dapat bahwa pada variabel integritas dan kepemimpinan terpahami dengan baik khususnya yang mahasiswa eksak di banding sosial.

Kata Kunci: Ketahanan Individu, Soft skill.

1. PENDAHULUAN

Ketahanan Individu merupan basic dari ketahanan Daerah dan Ketahanan Nasional. Pengertian Ketahanan nasional adalah kondisi dinamik bangsa Indonesia yang meliputi segenap aspek kehidupan nasional yang terintegrasi, berisi keuletan dan ketangguhan yang mengandung kemampuan mengembangkan kekuatan nasional, dalam menggapai dan mengatasi segala tantangan, ancaman, hambatan, dan gangguan baik yang dating dari luar dan dari dalam untuk menjamin identitas, integrasi, kelangsungan hidup bangsa dan negara.

Ketahanan individu terfokus pada kemampuan pribadi dalam menghadapi segala masalah dan tantangan kehidupan, sehingga faktor emosi, keimanan dan kesabaran menjadi penopang yang paling utama. Mahasiswa sebagai generasi penerus bangsa, akan berhadapan dengan hal tersebut, kemampuan dalam mengatasi masalah menjadi ukuran ketahanan dalam individu. Sifat dari ketahanan yakni lentur dan lenteng, sehingga kemampuan menghadapi segala situasi tetap mampu beradaptasi.

Ketahanan individu yang sering juga disebut *emotion question* atau bisa di setarakan dengan *soft skill*, yang menjadi syarat untuk dapat menjalankan dengan baik *hard skill*. *Soft skills* merupakan keterampilan dan kecakapan hidup, baik untuk sendiri, berkelompok, atau bermasyarakat, serta dengan Sang Pencipta. Dengan mempunyai *soft skills* membuat keberadaan seseorang akan semakin terasa tengah masyarakat. Keterampilan akan berkomunikasi, keterampilan emosional, keterampilan berbahasa, keterampilan berkelompok, memiliki etika dan moral, santun dan keterampilan spiritual menjadi kualitas yang seharusnya muncul di kalangan mahasiswa.

Mahasiswa sebagai sosok pembaharuan yang disibukkan pada agenda rutin belajar sehingga ada kecenderungan mahasiswa bersikap apatis terhadap permasalahan bangsa. Untuk membahas masalah nasionalisme dan kepedulian terhadap bangsa yang kondisi saat ini dinyatakan keprihatinan, tidak muncul suatu opini atau gerakan yang mampu memberikan sumbangan yang berarti. Seperti halnya diungkapkan oleh Paay presiden BEM dari UNSRAT (Universitas Sam Ratulangi – Menando).

“Sangat ironi kondisi di dunia kampus terlebih khusus mahasiswa. Jika dahulu semangat idealisme dalam memperjuangkan kepentingan bangsa dan Negara begitu menyala – nyala karena nasionalisme terhadap Negara Kesatuan Republik Indonesia - Presiden **BEM FKM Unsrat Jimmy Paays** mengungkapkan bahwa kondisi sekarang ini sudah berbeda dengan generasi yang terdahulu dimana sudah terjadi pergeseran budaya dalam memperjuangkan aspirasi mahasiswa dan juga aspirasi masyarakat. <http://www.seputarsulut.com/mahasiswa-sekarang-cenderung-apatis-jika-membahas-masalah-bangsa> diakses pd tgl 16/11/2017.

Ketahanan individu mahasiswa dapat dijadikan indikasi kepedulian terhadap masalah bangsa. Soft skill yang merupakan prasyarat kemampuan dalam menjalankan hard skill tidak bisa di tinggalkan, karena kemampuan hardskill tanpa soft skill akan sia sia. Masalah integritas menjadi masalah yang cukup serius karena kurang terpahaminya dengan baik maknanya. Apa yang disampaikan oleh ketua BEM FKM Unsrat, merupakan salah satu indikasi rendahnya integritas mahasiswa jaman sekarang, selain itu kemampuan dalam menghadapi perubahan merupakan tantangan tersendiri. Faktor kepemimpinan dan kemampuan berkomunikasi menjadi permasalahan sendiri.

Keempat masalah inilah yang dipandang menjadi bagian menarik untuk dicermati dalam penelitian ini khususnya dalam melihat ketahanan individu dalam perpektif *soft skill* khususnya di fakultas ilmu sosial dan ilmu politik , dan kedokteran Universitas Lampung.

2. LANDASAN TEORI

2.1. Ketahanan Nasional dan Ketahanan Individu.

Pengertian ketahanan nasional adalah kondisi dinamika, yaitu suatu bangsa yang berisi keuletan dan ketangguhan yang mampu mengembangkan ketahanan, Kekuatan nasional dalam menghadapi dan mengatasi segala tantangan, hambatan dan ancaman baik yang datang dari dalam maupun dari luar. Juga secara langsung

ataupun tidak langsung yang dapat membahayakan integritas, identitas serta kelangsungan hidup bangsa dan negara.

Dalam perjuangan mencapai cita-cita/tujuan nasionalnya bangsa Indonesia tidak terhindar dari berbagai ancaman-ancaman yang kadang-kadang membahayakan keselamatannya. Cara agar dapat menghadapi ancaman-ancaman tersebut, bangsa Indonesia harus memiliki kemampuan, keuletan, dan daya tahan yang dinamakan ketahanan nasional.

Kondisi atau situasi dan juga bisa dikatakan sikon bangsa kita ini selalu berubah-ubah tidak statik. Ancaman yang dihadapi juga tidak sama, baik jenisnya maupun besarnya. Karena itu ketahanan nasional harus selalu dibina dan ditingkatkan, sesuai dengan kondisi serta ancaman yang akan dihadapi dan konsep inilah yang disebut dengan sifat dinamika pada ketahanan nasional.

Untuk mengetahui ketahanan nasional, sebelumnya kita sudah tau arti dari wawasan nusantara. Ketahanan nasional merupakan kondisi dinamik yang dimiliki suatu bangsa, yang didalamnya terkandung keuletan dan ketangguhan yang mampu mengembangkan kekuatan nasional. Kekuatan ini diperlukan untuk mengatasi segala macam ancaman, tantangan, hambatan dan gangguan yang langsung atau tidak langsung akan membahayakan kesatuan, keberadaan, serta kelangsungan hidup bangsa dan negara. Bisa jadi ancaman-ancaman tersebut dari dalam ataupun dari luar. [1]

Ketahanan individu secara teoritis dalam ketahanan nasional belum ada, namun demikian simpul dari ketahanan daerah yang menjadi penopang dari ketahanan nasional, yakni ketahanan individu. Oleh karena itu ketahanan individu diperlukan dalam rangka menjamin eksistensi bangsa dan negara dari segala gangguan baik yang datang dari dalam maupun dari dalam negeri. Salah satu ciri yang kuat dalam ketahanan individu yaitu memiliki keuletan dan ketangguhan yang perlu dibina secara konsisten dan berkelanjutan.

2.2. *Soft skills* sebagai suatu paradigma dalam Ketahanan Individu.

Paradigma dalam filsafat ilmu merupakan pola pikir yang dapat dijadikan contoh atau panutan, sehingga memungkinkan untuk berkembangnya ilmu pengetahuan, karena memiliki obyek material dan obyek formal yang sama. Menurut Thomas Kuhn, **pengertian paradigma** adalah landasan berpikir atau pun konsep dasar yang digunakan / dianut sebagai model atau pun pola yang dimaksud para ilmuwan dalam usahanya, dengan mengandalkan studi – studi keilmuan yang dilakukannya.

Ketahanan individu merupakan turunan dari ketahanan daerah, dan apa yang ada di dalam kemampuan individu identik dengan apa yang ada di kompetensi *soft skill*. Pemahaman *soft skill* sangat umum dalam dunia pendidikan atau dunia kerja, karena apa yang ada di *soft skill* menjadi prasyarat yang mendasar pada kemampuan *hard skill*. Sebagai contoh keterampilan-keterampilan yang dimasukkan dalam kategori *soft skills*

adalah integritas, inisiatif, motivasi, etika, kerja sama dalam tim, kepemimpinan, kemauan belajar, komitmen, mendengarkan, tangguh, fleksibel, komunikasi lisan, jujur, berargumen logis, dan lainnya. Keterampilan-keterampilan tersebut umumnya berkembang dalam kehidupan bermasyarakat. *Soft skills* didefinisikan sebagai "Personal and interpersonal behaviors that develop and maximize human performance (e.g. coaching, team building, initiative, decision making etc.) [2]

Soft skill merupakan keterampilan dan kecakapan hidup, baik untuk sendiri, berkelompok, atau bermasyarakat, serta dengan sang pencipta. Selebihnya dengan mempunyai *soft skills* membuat keberadaan seseorang akan semakin terasa di masyarakat. Keterampilan akan berkomunikasi, keterampilan emosional, keterampilan bahasa, keterampilan berkelompok, memiliki etika dan moral, santun, dan keterampilan spriritual.

Soft skills merupakan jenis ketrampilan yang lebih banyak terkait dengan sensitivitas perasaan seseorang terhadap lingkungan di sekitarnya. Karena *soft skills* terkait dengan ketrampilan psikologis, maka dampak yang diakibatkan lebih abstrak namun tetap bisa dirasakan seperti misalnya perilaku sopan, disiplin, keteguhan hati, kemampuan untuk dapatbekerja sama, membantu oranglain, dan sebagainya. Konsep *soft skills* merupakan istilah sosiologis yang merepresentasikan pengembangan dari kecerdasan emosional (*emotional intelligence*) seseorang yang merupakan kumpulan karakter kepribadian, kepekaan sosial, komunikasi, bahasa, kebiasaan pribadi, keramahan, dan optimisme yang menjadi ciri hubungan dengan orang lain. *Soft skills* melengkapi hard skills, dimana *hard skills* merupakan representasi dari potensi IQ seseorang terkait dengan persyaratan teknis pekerjaan dan beberapa kegiatan lainnya.

Dalam modul uji kompetensi untuk assesor BNSP (Badan Nasional Stratdarisasi Profesi tahun 2015) ada empat *soft skills* yang terkait dengan ketahanan individu, yakni: a) integritas yang artinya kemampuan mengorbankan keinginan jangka pendek bagian/ unit kerjanya guna kebaikan jangka panjang organisasi, memiliki dan mengaplikasikan norma-norma yang sejalan dengan organisasi. b) Kemampuan menghadapi perubahan (*ability to change*) yang artinya kemampuan menyesuaikan strategi diri terhadap perubahan organisasi, menanggapi tantangan baru dan aktif menyusun strategi. c) kepemimpinan (leadership) kemampuan menyiapkan sistem dan struktur yang dibutuhkan dalam perubahan, menciptakan suasana yang mampu menggerakkan organisasi ke arah yang diinginkan dan d) kemampuan berkomunikasi (*communication skills* yakni kemampuan mengeksplorasi terhadap lawan bicara dilakukan secara tajam dan spesifik sehingga kesepakatan tidak terkesan dipaksakan, serta memiliki pengaruh yang kuat dalam organisasi.

Soft skills memiliki banyak manfaat, untuk pengembangan karir serta etika profesional. Dari sisi organisasional, *soft skills* memberikan dampak terhadap kualitas manajemen secara total, efektivitas institusional dan sinergi inovasi. Esensi *soft skills* adalah kesempatan untuk membangun alumni yang berkualitas. Terbentuknya kualitas alumni sangat sangat di tentukan dalam pemahaman soft skil, demikian juga dengan ketahanan nasional, bahwa ketahanan nasional yang berfungsi sebagai doktrin nasional perlu dipahami dengan baik guna menjamin tetap terjadinya pola pikir, pola sikap, dan pola tindak dalam menyatukan langkah bangsa.

Asas dalam [3] meliputi: a) asas kesejahteraan dan kemanan, b) asas komprehensif / menyeluruh dan terpadu, dan c) asas kekeluargaan. Sedangkan sifat dari ketahanan nasional yakni: a) manunggal, b) mawas ke dalam, c) kewibawaan, d) dinamis menitik beratkan konstitusi dan saling menghargai, dan mandiri. Dengan mendasarkan pada tiga asas dan enam fungsi ketahanan nasional maka pembentukan ketahanan individu semakin relevan, demikian pula dengan kemampuan soft skill yang dibangun.

3. METHODOLOGI PENELITIAN

Dalam penelitian ini metode yang digunakan yakni metode kuantitatif deskriptif yang bertujuan untuk memberikan gambaran tentang suatu keadaan yang menjadi fokus penelitian yang dapat berupa hubungan antara dua gejala atau lebih. Selain itu digunakan pula penelitian kepustakaan dimana peneliti melalui kepustakaan mengumpulkan data-data dan keterangan dari buku yang berhubungan dengan masalah yang diteliti.

Penentuan jumlah sampel diambil dengan rumus slovin,

$$n = \frac{N}{N(\alpha^2) + 1}$$

Keterangan:

N = Jumlah Populasi

n = jumlah sampel

α = derajat kesalahan

dengan teknik pengambilan data dilakukan dengan menggunakan kuesioner tertutup serta menggunakan lima pilihan dalam model skala Likert. Dari hasil Likert di buat tabulasi yang kemudian di cari reratanya untuk mengetahui kecenderungan jawaban. Setelah itu disajikan dalam bentuk grafik yang kemudian di analisis.

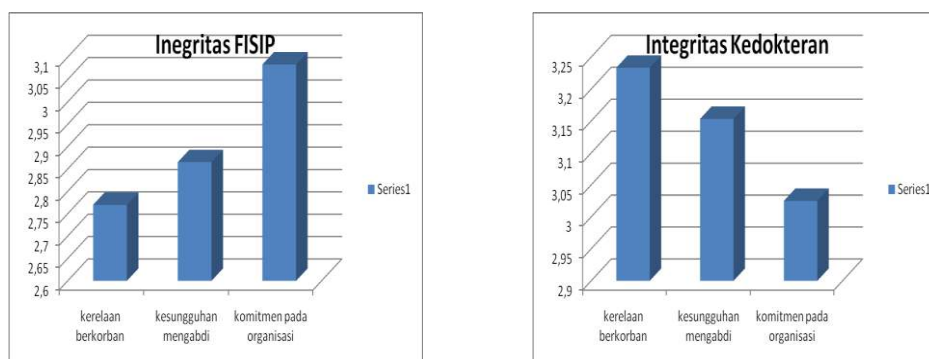
4. HASIL DAN PEMBAHASAN

Ketahanan Daerah dan Ketahanan Nasional adalah dua konsep terkait-tak terpisahkan bagi eksistensi Indonesia sebagai sebuah negara-bangsa (*nation-state*). Indonesia yang kuat adalah Indonesia yang mampu mewujudkan Tujuan Nasional-nya. Sementara Ketahanan Nasional, tak mungkin dicapai tanpa ditopang Ketahanan Daerah yang tangguh. Ketahanan Nasional yang tangguh akan terwujud manakala daerah-daerah di seluruh Indonesia dapat mengatasi sejumlah problematika bangsa dan Negara, dan hal ini tidak terlepas dari ketahanan individu.

Ketahanan individu yang diidentikkan dengan soft skill di dalam penelitian ini memberikan makna bahwa apa yang di perbuat oleh individu pada dasarnya merupakan cerminan dari kemampuan soft skill yang dimiliki. Konsep integritas yang didalamnya terdiri dari tiga indikator yakni: a) kerelaan berkorban, dalam hal ini

mahasiswa secara ikhlas mengorbankan waktu dan tenaganya untuk mengikuti proses belajar dengan maksimum, b) kesungguhan dalam pengabdian di artikan sebagai totalitas dalam belajar sebagai mahasiswa, dan c) komitmen pada organisasi di artikan sebagai kesungguhan memegang janji dan memegang teguh kehormatan lembaga dalam hal ini Fakultas dan jurusan dan Unila pada Umumnya.

Berikut disajikan secara tabulasi dengan menggunakan visual grafik untuk memlihat hasil olah data integritas di kedua fakultas yakni Fakultas kedokteran dan Fisip.

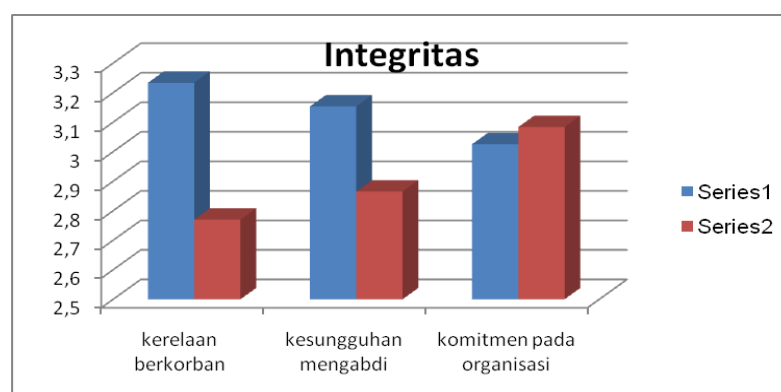


Sumber Data Diolah 2017.

Grafik 1. Gambaran Intgritis mahasiswa Fisip dan Kedokteran.

Dari kedua grafik 1 a dan b di atas, terlihat bahwa integritas yang tumbuh dan berkembang di kedua fakultas berbeda, bahkan ada kecenderungan berlawanan. Mahasiswa Fisip melihat bahwa integritas lebih dimaknai sebagai komitmen pada organisasi yang lebih dianggap penting dan kerelaan berkorban yang merupakan dasar dalam integritas kurang dimaknai dengan baik. *Soft skill* yang menjadi bagian dalam integritas khususnya kerelaan berkorban menjadi pondasi dalam ketahanan Individu.

Demikian sebaliknya dengan mahasiswa kedokteran bahwa kerelaan berkorban menjadi kekuatan dalam mengartikan integritas, sehingga hal ini pada dasarnya sejalan dengan jiwa dan semangat corps kedokteran. Berikut gambaran komparasi kedua fakultas terhadap pemahaman integrasi.



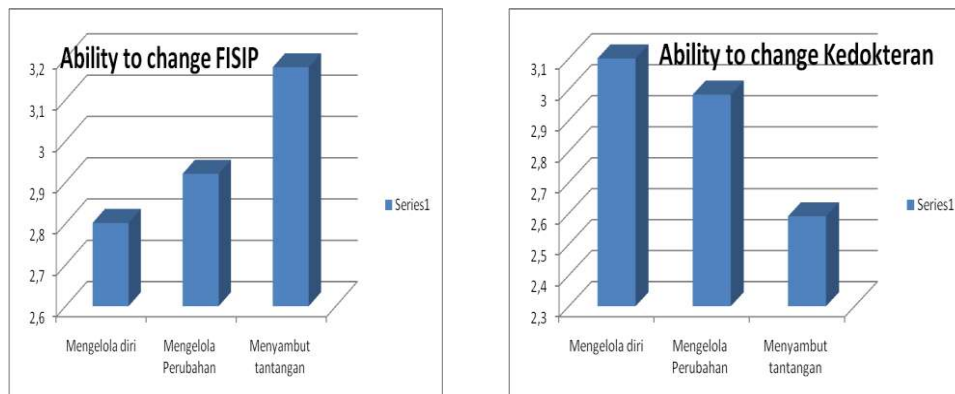
Sumber Data Diolah 2017.

Grafik 2. Gambaran komparasi Integritas

Konsep integritas merupakan parameter dasar untuk melihat karakter mahasiswa terhadap kepedulian dengan masalah di luar dirinya. Memperhatikan dari komparasi kedua fakultas yang dapat dinyatakan sebagai perbandingan pemikiran eksak dan sosial, dapat di pahami bahwa masalah integritas yang lebih pada pemahaman *soft skill* dengan baik adalah pada pemikiran eksak. Oleh karena itu di dalam kenyataan di lapangan atau masyarakat pemikiran eksak lebih mendominasi di dalam struktur organisasi.

Kemampuan menghadapi perubahan atau *ability to change* juga merupakan faktor yang mendasar dalam *soft skill*, karena tidak mungkin di dalam proses fase kehidupan tidak ada perubahan, dan perubahan akan selalu membawa dampak positif maupun negatif. Kemampuan dalam menghadapi perubahan yang terdiri dari tiga indikator yakni a) kemampuan mengelola diri. Artinya terdapat kemampuan penyesuaian terhadap lingkungan dan beradaptasi untuk stabil dan terus berkembang. Jika di gambarkan pada posisi mahasiswa adalah gambaran mahasiswa yang tidak gampang menyerah dan memiliki daya juang, b) kemampuan mengelola perubahan maknanya yakni mahasiswa mampu menjawab tantangan dari perubahan yang ada, dan c) menyambut tantangan yang ada dengan menyiapkan segala sesuatunya secara lengkap, sehingga tantangan berubah suatu harapan.

Berikut ini gambaran kemampuan menghadapi perubahan pada dua fakultas yakni Fisip dan kedokteran sebagai berikut:



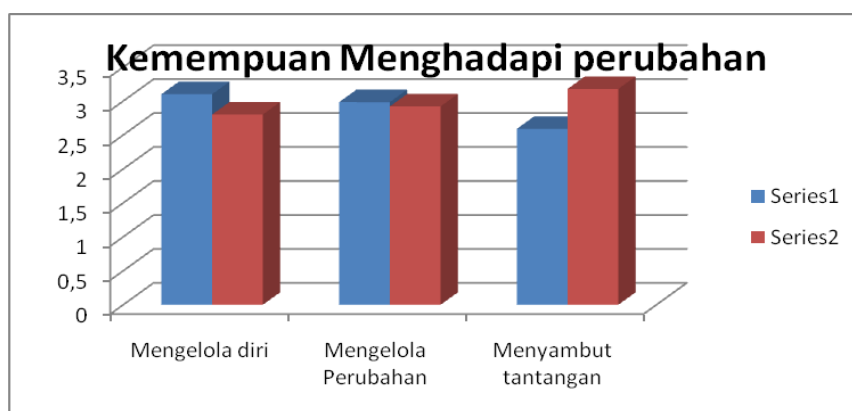
Sumber data diolah 2017

Grafik 3 Gambaran kemampuan dalam menghadapi perubahan

Dari gambaran grafik 3 di atas, dapat dinyatakan bahwa kemampuan menghadapi perubahan mahasiswa sosial dalam hal ini lebih siap dalam mengatur ataupun mencari peluang dalam fase perubahan. Bila di cermati, kondisi sosial politik yang “cenderung berubah atau abu abu” menjadikan peluang dan tantangan jauh lebih terbuka dan memberikan harapan. Tidak halnya dengan mahasiswa kedokteran yang meragukan perubahan akan bermakna positif.

Perbedaan yang cukup mendasar dalam kemampuan menghadapi perubahan yakni, jika mahasiswa sosial (Fisip) kemampuan mengelola dirinya relatif lemah, artinya *soft skill* dan ketahanan individu dalam menghadapi tantangan perubahan cenderung lemah atau rendah, justru pada aspek kesempatan lebih diutamakan. Demikian

sebaliknya dengan mahasiswa kedokteran, yang melihat bahwa central menghadapi perubahan adalah kemampuan individu yang tangguh. Hal ini sesuai, searah dengan konsep ketahanan individu.

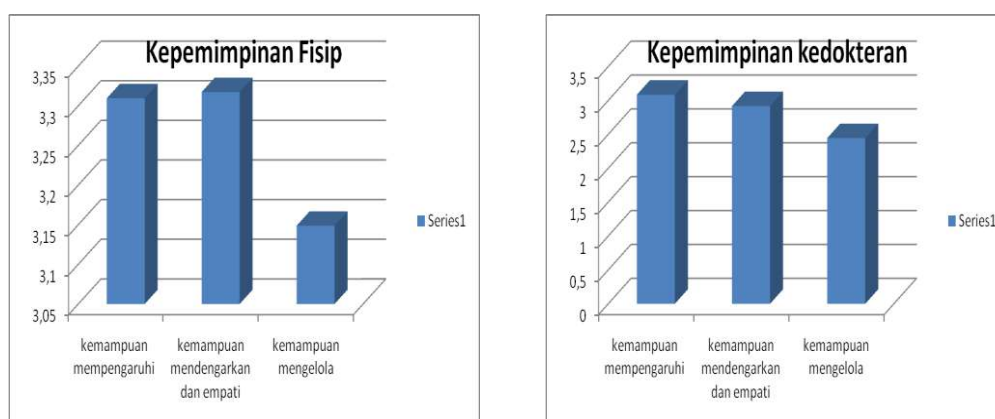


Sumber data diolah 2017

Grafik 4 Gambaran komparasi menghadapi perubahan

Memperhatikan grafik 4 berupa gambaran komparasi menghadapi perubahan, dapat dinyatakan bahwa secara ketahanan individu, mahasiswa eksak atau kedokteran lebih siap di bandingkan mahasiswa sosial, kemampuan mengelola diri dan mempersiapkan perubahan lebih mapan. Sedangkan dalam melihat peluang perubahan baik sosial maupun eksak hampir sama, artinya pola berikir positif terhadap perubahan memiliki cara pandang yang hampir sama.

Pada aspek kepemimpinan yang di artikan sebagai a) kemampuan mempengaruhi, b) kemampuan mendengarkan dan empati, dan c) kemampuan mengelola atau mengatur. Adapun gambaran kepemimpina di kalangan mahasiswa fisip dan kedokteran adalah sebagai berikut.

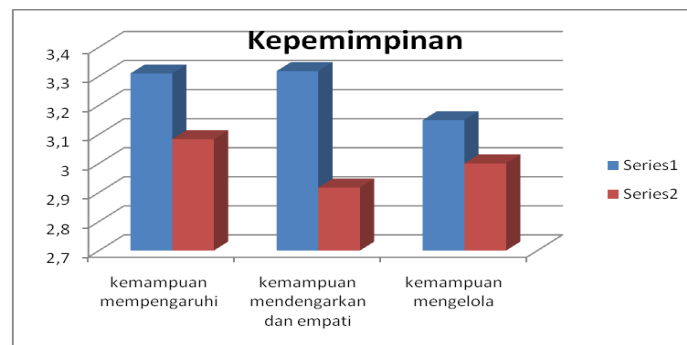


Sumber data diolah 2017

Gambar 5 Gambaran kepemimpinan di fisip dan kedokteran.

Secara statistik, faktor kepemimpinan di dalam pemahaman mahasiswa fisip lebih menguasai, bisa jadi terkait dengan faktor obyek utama dalam fisip yakni kekuasaan, maka pemahaman konsep kepemimpinan lebih termaknai dengan baik. Skor yang lebih tinggi dari pada mahasiswa kedokteran yangkni di atas 3,3 rata rata yang lebih baik daripada kedokteran yang di bawah tiga, menandakan kemampuan mempengaruhi kemampuan mendengar dan berempati mahasiswa fisip dalam hasil penelitian lebih unggul.

Sedangkan kemampuan mempengaruhi yang ada pada mahasiswa kedokteran dalam konteks organisasi tidak begitu dihiraukan. Seperti yang tersaji gambaran grafik dalam komparasi berikut ini:



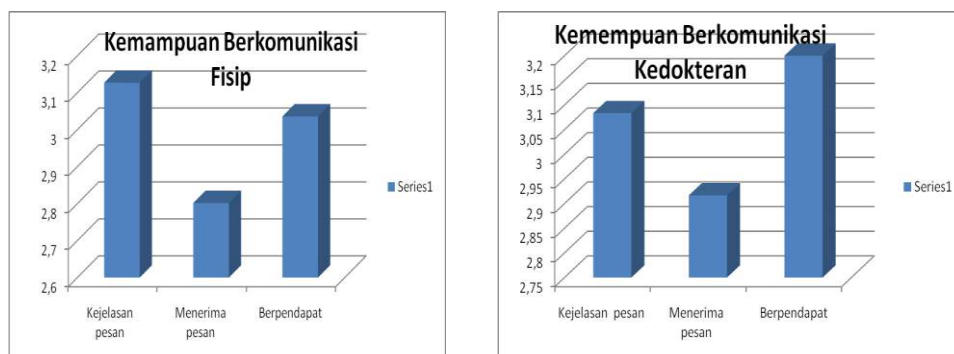
Sumber data diolah 2017

Gambar 6 Grafik komparasi kepemimpinan

Pada kepemimpinan pemahaman mahasiswa fisip lebih termaknai dengan baik.seperti dalam kemampuan mempengaruhi, mendengarkan pendapat dan kemampuan mengelola dalam arti manajerial.

Pada aspek keempat yakni kemampuan berkomunikasi yang terdiri atas: a) kejelasan pesan yang disampaikan atau dapat diartikan isi pesan dapat dipahami dengan baik antara pemberi ide dan penerima ide. b) menerima pesan yakni lebih dimaknai dengan kemampuan mendengar pendapat orang lain, sehingga kesabaran dan kejelian dalam menangkap isi pesan menjadi sangat penting. dan

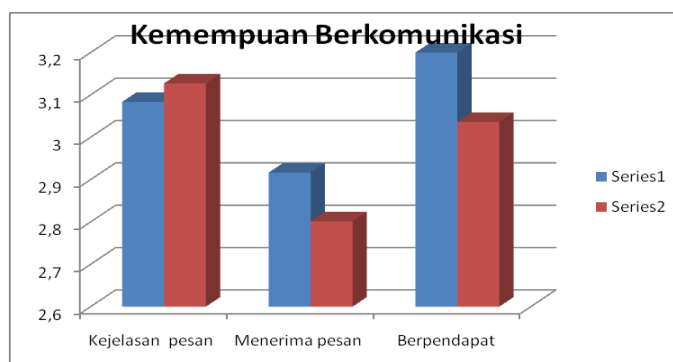
c) kemampuan berpendapat atau kemampuan berargumentasi. Perbandingan ketiga tersebut seperti tersajikan dalam grafik berikut ini.



Sumber data diolah 2017

Gambar 7. kemampuan berkomunikasi

Kemampuan berkomunikasi sebagai kemampuan soft skill yang mendasar untuk kedua fakultas fisip dan kedokteran memiliki kemiripan pola. Hanya pada faktor kemampuan mengungkapkan pendapat, kejelasan lebih unggul di mahasiswa kedokteran. Sedang komprasi keduanya dapat disajikan sebagai berikut:



Sumber data diolah 2017

Gambar 8 grafik komparasi kemampuan berkomunikasi.

Kemampuan berkomunikasi merupakan kompetensi dasar setiap individu, kemampuan dasar yang merupakan soft skill pada setiap individu dimaknai kemampuan untuk mengemukakan pendapat dan berargumentasi. Pada hasil komprasi, ternyata mahasiswa kedokteran memiliki kelebihan dalam membangun argumentasi di banding mahasiswa fisip.

Dengan demikian empat faktor yang ada dalam soft skill dalam konteks ketahanan individu, menjadi semakin njelas dan nyata, bahwa sangat erat hubungan keempat faktor *soft skill* dalam ketahanan individu guna membangun ketahanan nasional. Ketahanan individu yang bersisikan beberapa perangkat *soft skill* ternyata sangat mendukung dan mendasari dalam pembangunan ketahanan individu. Memperhatikan hasil keseluruhan khususnya pada grafik komporasi, maka dapat dinyatakan bahwa tingkat ketahanan individu mahasiswa baik dan fisip dan kedokteran memiliki level yang cukup bagus, hal ini terlihat dari rata rata yang ada yakni di kisaran 3 dengan skala 1 sampai 4. Sedangkan gambaran umum kepedulian dan tingkat ketahanan individu, mahasiswa kedokteran lebih baik.

5. SIMPULAN.

1. Ketahanan individu dalam perspektif *soft skill* masih sangat relevan untuk menjadi tolok ukur dalam mengukur kemampuan individu khususnya dalam hal integritas, kemampuan dalam menghadapi perubahan, kepemimpinan dan kemampuan berkomunikasi.
2. Mahaiswa eksak dalam hal ini kedokteran dalam kemampuan berpikir dan menterjemahkan tantangan hidup lebih baik dibanding mahasiswa sosial dalam hal ini fisip.
3. Dalam hal kepemimpinan mahasiswa fisip memiliki kemampuan menterjemahkan lebih baik di banding mahasiswa kedokteran. Demikian halnya kemampuan dalam berkomunikasi mahasiswa fisip lebih jelas ekspresinya dalam tuntutananya.

4. Pemahaman konsep integritas antara fisip dan kedokteran berbanding terbalik, artinya pada faktor kerelaan berkorban, mahasiswa kedokteran lebih tinggi di bandingkan mahasiswa fisip, dan sesuai dengan karakter pendidikan. Sedangkan mahasiswa fisip lebih memfokuskan makna integritas pada aspek kelembagaan, jadi ada perspektif yang berkembang yaitu mikro di kalangan mahasiswa kedokteran yang berarti berpusat pada individu, dan konsep makro yang berpusat pada organisasi yang berkembang di kalangan mahasiswa fisip.

KEPUSTAKAAN

- [1] Srijanti. (2009). **Ketahanan Nasional dalam suatu perspektif pendidikan**, Alfa Beta, Bandung
- [2] Saeful Zaman. (2011). **Paradigma Soft Skill dalam mencapai sukses**, Media Perubahan Housing Publisher, Jakarta.
- [3] Anonim, Lemhanas, (2000): **Konsep Ketahanan Nasional**, diktat/ fotocopy)

METODE ANALISIS HOMOTOPHI PADA SISTEM PERSAMAAN DIFERENSIAL PARSIAL LINEAR NON HOMOGEN ORDE SATU

Atika Faradilla, Suharsono S.

Jurusan Matematika FMIPA Universitas Lampung, Bandar Lampung
Jl. Prof. Dr. Sumantri Brojonegoro No. 1 Bandar Lampung. Telp 0721701609. Fax 0721702767
E-mail: atika.faradilla@yahoo.com

ABSTRAK

Pada tahun 1992, Liao mengusulkan sebuah metode analisis untuk memecahkan masalah – masalah persamaan maupun sistem persamaan, bernama metode analisis homotopi. Metode ini digunakan untuk memecahkan masalah dengan langkah yang sesuai dan memudahkan penentuan konvergensi deret. Dalam beberapa tahun belakangan ini, metode ini telah dengan sukses dapat diaplikasikan untuk menyelesaikan berbagai persamaan yang linear maupun tak linear dan yang homogen maupun non homogen. Untuk memahami lebih dalam mengenai metode ini, maka penulis mencoba menerapkan metode analisis homotopi pada sistem persamaan diferensial linear non homogen orde satu. Dengan menerapkan metode analisis homotopi ke dalam sistem persamaan diferensial parsial linear non homogen $u_t - v_x - (u - v) = -2$, $v_t + u_x - (u - v) = -2$ dengan syarat awal $u(x, 0) = 1 + e^x$ dan $v(x, 0) = -1 + e^x$, didapatkan solusi homotopi yang konvergen ke solusi $u(x, t) = 1 + e^{x+t}$ dan $v(x, t) = -1 + e^{x-t}$ apabila nilai $h = -1$.

Kata kunci: metode analisis homotopi, persamaan linear non homogen, sistem persamaan diferensial parsial

1. PENDAHULUAN

Persamaan diferensial parsial dapat dijumpai dalam kaitan dengan berbagai masalah fisik dan geometris bila fungsi yang terlibat tergantung pada dua atau lebih peubah bebas. Tidak berlebihan jika dikatakan bahwa hanya sistem fisik yang paling sederhana yang dapat dimodelkan dengan persamaan diferensial biasa, sedangkan masalah-masalah yang lebih rumit, seperti topik-topik dalam fisika lanjut, harus dimodelkan dengan persamaan diferensial parsial.

Kenyataannya, persamaan – persamaan tersebut sulit untuk dipecahkan dengan penyelesaian persamaan diferensial parsial analitik biasa. Sehingga dikembangkan berbagai metode yang dapat menyelesaikan persamaan diferensial parsial ini, seperti metode pemisahan variabel, metode transformasi Laplace, metode transformasi Fourier, metode beda hingga, metode garis, dan sebagainya.

Pada tahun 1992, Liao mengusulkan sebuah metode dalam disertasi PhD-nya, bernama metode analisis homotopi. Metode ini kemudian dimodifikasi lebih lanjut pada tahun 1997 untuk memperkenalkan parameter untuk membangun homotopi pada sistem diferensial dalam bentuk umum. Dalam beberapa tahun belakangan ini, metode ini telah dengan sukses dapat diaplikasikan untuk menyelesaikan berbagai persamaan maupun sistem persamaan linear maupun tak linear dan yang homogen maupun non homogen.

Metode analisis homotopi ini digunakan untuk memecahkan masalah dengan langkah yang sesuai dan memudahkan penentuan konvergensi deret. Selain itu, metode analisis homotopi tidak bergantung pada besar kecilnya parameter. Kenyataannya adalah metode homotopi lebih mudah digunakan dalam menyelesaikan masalah yang sulit. Berdasarkan hal-hal tersebut, maka penulis mencoba menerapkan metode analisis homotopi pada sistem persamaan diferensial linear non homogen.

Misalkan diberikan suatu sistem persamaan linear non homogen $u_t - v_x - (u - v) = -2$, $v_t - u_x - (u - v) = -2$. Dengan kondisi awal $u(x, 0) = 1 + e^x$ dan $v(x, 0) = -1 + e^x$. Untuk menyelesaikan sistem persamaan tersebut, maka akan digunakan metode analisis homotopi.

Tujuan dari penelitian ini adalah sebagai berikut:

1. Menjelaskan tentang metode analisis homotopi.
2. Menerapkan metode analisis homotopi pada sistem persamaan diferensial parsial linear non homogen $u_t - v_x - (u - v) = -2$, $v_t - u_x - (u - v) = -2$.
3. Menyelesaikan sistem persamaan diferensial parsial linear non homogen $u_t - v_x - (u - v) = -2$, $v_t - u_x - (u - v) = -2$ hingga diperoleh solusinya.

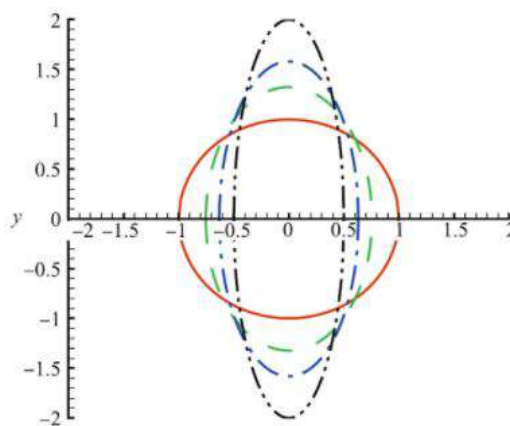
2. LANDASAN TEORI

Konsep Dasar Persamaan Diferensial Parsial

Persamaan diferensial adalah suatu persamaan yang memuat turunan terhadap satu atau lebih dari variabel-variabel bebas. Berdasarkan jumlah variabel bebasnya persamaan differensial dibagi dalam dua kelas, yaitu persamaan differensial biasa (PDB) dan persamaan differensial parsial (PDP). Jika turunan fungsi itu hanya tergantung pada satu variabel bebas, maka disebut persamaan differensial biasa. Sedangkan, jika turunan fungsi itu tergantung pada lebih dari satu variabel bebasdisebut persamaan differensial parsial[1].

2.2 Deformasi

Deformasi adalah perubahan bentuk sebuah benda dari kondisi semula ke kondisi terkini. Makna dari "kondisi" dapat diartikan sebagai serangkaian posisi dari semua partikel yang ada di dalam benda tersebut. Sebagai contoh sebuah lingkaran dapat dideformasikan secara kontinu menjadi elips, dan bentuk dari cangkir kopi dapat dideformasikan secara kontinu menjadi bentuk donat.



Gambar 1. Deformasi kontinu dari solusi $y(x; q)$ dari persamaan (1)

Keterangan

: $q = 0$	
: $q = 1/4$	
: $q = 1/2$	
: $q = 1$	

Misalkan terdapat suatu persamaan

$$\varepsilon(q) : (1 + 3q)x^2 + \frac{y^2}{(1+3q)} = 1, \quad q \in [0,1] \quad (1)$$

Dengan $q \in [0,1]$ adalah parameter homotopi.

Jika $q = 0$, maka akan didapat persamaan lingkaran

$$\varepsilon_0 : x^2 + y^2 = 1 \quad (2)$$

Yang solusinya merupakan lingkaran $y = \pm \sqrt{1 - x^2}$. Jika $q = 1$, maka didapat persamaan elips

$$\varepsilon_1 : 4x^2 + y^2/4 = 1 \quad (3)$$

Yang solusinya merupakan elips $y = \pm 2\sqrt{1 - 4x^2}$.

Sehingga, seiring dengan bertambahnya parameter homotopi q dari 0 ke 1, persamaan (1) berubah secara kontinu dari persamaan lingkaran ε_0 menjadi persamaan elips ε_1 , sementara solusinya ydeformasi secara kontinu dari lingkaran $y = \pm \sqrt{1 - x^2}$ menjadi elips $y = \pm 2\sqrt{1 - 4x^2}$, seperti dapat dilihat dalam Gambar 1.

Maka, solusi y dari (1) tidak hanya bergantung kepada x tapi juga bergantung kepada $q \in [0,1]$, sehingga persamaan (1) dapat ditulis dengan lebih tepat sebagai

$$\varepsilon(q) : (1 + 3q)x^2 + \frac{y^2(x;q)}{(1+3q)} = 1, \quad q \in [0,1] \quad (4)$$

Yang mendefinisikan dua homotopi: pertama, homotopi dari persamaan

$$\varepsilon(q) : \varepsilon_0 \sim \varepsilon_1 \quad (5)$$

Dimana ε_0 dan ε_1 secara berturut – turut dinotasikan dalam (2) dan (3). Yang kedua adalah homotopi dari fungsi

$$y(x; q) = \pm \sqrt{1 - x^2} \sim \pm 2\sqrt{1 - 4x^2} \quad (6)$$

Dengan kata lain, solusi $y(x; q)$ dari (4) juga merupakan homotopi. Deformasi kontinu secara lengkap telah dinotasikan dalam (4). Untuk lebih sederhananya, (4) dapat disebut sebagai deformasi orde – nol. Konsep yang sama dapat diterapkan kedalam berbagai jenis persamaan lainnya, seperti persamaan diferensial, persamaan integral, dan lain – lain [2].

2.3Metode Analisis Homotopi

Misalkan terdapat persamaan differensial berikut,

$$N_i[z_i(x, t)] = 0 \quad ; i = 1, 2, \dots, n \quad (7)$$

dengan N_i adalah operator nonlinear yang mewakili seluruh persamaan, x dan t adalah variabel bebas dan $z_i(x, t)$ adalah fungsi yang tidak diketahui. Dengan cara mengeneralisasi metode homotopi sederhana, Liao menyusun persamaan deformasi orde nol.

$$(1 - q) L[\phi_i(x, t; q) - z_{i,0}(x, t)] = q h_i N_i [\phi_i(x, t; q)] \quad (8)$$

Untuk $q \in [0, 1]$ adalah parameter homotopi, h_i adalah fungsi tak nol tambahan, L adalah operator linear tambahan, $z_{i,0}(x, t)$ adalah syarat awal dari $z_i(x, t)$ dan $\phi(x, t; q)$ adalah fungsi yang tidak diketahui. Penting untuk diingat bahwa, bebas untuk memilih objek tambahan seperti h_i dan L pada metode analisis homotopi. Terlihat jelas bahwa saat $q = 0$ dan $q = 1$, keduanya menghasilkan

$$\phi_i(x, t; 0) = z_{i,0}(x, t) \text{ dan } \phi_i(x, t; 1) = z_i(x, t) \quad (9)$$

Maka, seiring dengan bertambahnya nilai q dari 0 ke 1, solusi $\phi_i(x, t; q)$ berubah dari syarat awal $z_{i,0}(x, t)$ ke solusi $z_i(x, t)$. Mengekspansikan $\phi_i(x, t; q)$ ke dalam deret Taylor terhadap q , akan menghasilkan

$$\phi_i(x, t; q) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t) q^m \quad (10)$$

dengan

$$z_{i,m} = \frac{1}{m!} \left. \frac{\delta^m \phi_i(x, t; q)}{\delta q^m} \right|_{q=0} \quad (11)$$

Jika operator linear tambahan, syarat awal, parameter tambahan h_i , dan fungsi tambahan dipilih dengan benar, maka deret pada persamaan (10) konvergen ke $q = 1$ dan

$$\phi_i(x, t; 1) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t) \quad (12)$$

Yang merupakan salah satu dari solusi – solusi persamaan nonlinear, seperti yang telah dibuktikan oleh Liao. Jika $h_i = -1$, persamaan (8) menjadi

$$(1 - q) L[\phi_i(x, t; q) - z_{i,0}(x, t)] + q N_i [\phi_i(x, t; q)] = 0 \quad (13)$$

Yang merupakan persamaan yang paling sering digunakan dalam metode analisis homotopi.

Berdasarkan (11), persamaan tersebut dapat di deduksi dari persamaan deformasi orde nol (8). Didefinisikan vektor

$$\overrightarrow{z_{i,n}} = \{ z_{i,0}(x, t), z_{i,1}(x, t), \dots, z_{i,n}(x, t) \} \quad (14)$$

Mendiferensialkan (8) sebanyak m kali terhadap parameter homotopi q dan substitusikan $q = 0$ dan akhirnya membaginya dengan $m!$, maka diperoleh persamaan deformasi orde ke- m .

$$L [z_{i,m}(x, t) - \chi_m z_{i,m-1}(x, t)] = h_i R_{i,m}(\overrightarrow{z_{i,m-1}}) \quad (15)$$

untuk

$$R_{i,m}(\overrightarrow{z_{i,m-1}}) = \frac{1}{(m-1)!} \left. \frac{\delta^{m-1} N_i[\phi_i(x, t; q)]}{\delta q^{m-1}} \right|_{q=0} \quad (16)$$

dan

$$\chi_m = \begin{cases} 0, & m \leq 1 \\ 1, & m > 1 \end{cases}$$

Perhatikan bahwa $z_{i,m}(x, t) (m \geq 1)$ berdasarkan persamaan linear (15) dengan kondisi batas linear yang didapat dari masalah awalnya [3].

3. METODOLOGI PENELITIAN

Adapun langkah-langkah yang dilakukan pada penelitian ini adalah sebagai berikut:

1. Mendefinisikan operator linear L dan mengkonstruksikan persamaan deformasi orde ke-nol.
2. Menentukan rangkaian solusi dari komponen yang telah diperoleh melalui perkiraan awal jika $q=0$, dan $q=1$.
3. Mengkonstruksikan persamaan deformasi orde-m.
4. Menentukan solusi persamaan deformasi pada persamaan yang diperoleh dari langkah ke (3) untuk setiap $m = 1, 2, \dots, n$.
5. Mensubstitusikan hasil ini ke dalam deret homotopi, sehingga diperoleh solusi homotopi.

4. HASIL DAN PEMBAHASAN

Diberikan suatu sistem persamaan linear tak homogen

$$u_t - v_x - (u - v) = -2 \quad (17)$$

$$v_t + u_x - (u - v) = -2 \quad (18)$$

dengan syarat awal

$$u(x, 0) = 1 + e^x, \text{ dan } v(x, 0) = -1 + e^x \quad (19)$$

Untuk menyelesaikan sistem (1) – (3) dengan metode analisis homotopi, maka dipilih

$$u_0(x, 0) = 1 + e^x \text{ dan } v_0(x, 0) = -1 + e^x.$$

Didefinisikan operator linear sebagai berikut

$$L[\phi_i(x, t; q)] = \frac{\partial \phi_i(x, t; q)}{\partial t} \quad (20)$$

untuk

$$L[c_i] = 0 \quad (21)$$

Dan $c_i (i = 1, 2)$ merupakan konstanta integral.

Selanjutnya, berdasarkan sistem (1) dan (2) didefinisikan sistem operator nonlinear sebagai berikut.

$$N_1[\phi_i(x, t; q)] = \frac{\partial \phi_1(x, t; q)}{\partial t} - \frac{\partial \phi_2(x, t; q)}{\partial x} - \phi_1(x, t; q) + \phi_2(x, t; q) + 2$$

$$N_2[\phi_i(x, t; q)] = \frac{\partial \phi_2(x, t; q)}{\partial t} + \frac{\partial \phi_1(x, t; q)}{\partial x} - \phi_1(x, t; q) + \phi_2(x, t; q) + 2$$

Dengan menggunakan definisi diatas, dikonstruksikan persamaan deformasi orde ke-0.

$$(1 - q) L[\phi_i(x, t; q) - z_{i,0}(x, t)] = q h_i N_i[\phi_i(x, t; q)], \quad i = 1, 2 \quad (22)$$

Sehingga, saat $q = 0$ dan $q = 1$,

$$\begin{aligned} \phi_1(x, t; 0) &= z_{1,0}(x, t) = u_0(x, t), & \phi_1(x, t; 1) &= u(x, t), \\ \phi_2(x, t; 0) &= z_{2,0}(x, t) = v_0(x, t), & \phi_2(x, t; 1) &= v(x, t), \end{aligned}$$

Maka, seiring dengan bertambahnya nilai q dari 0 ke 1, solusi $\phi_i(x, t; q)$ berubah dari syarat awal $z_{i,0}(x, t)$ ke solusiz_i(x, t). Kemudian jika $\phi_i(x, t; q)$ diekspansikan ke dalam deret Taylor terhadap q , akan menghasilkan

$$\phi_i(x, t; q) = z_{i,0}(x, t) + \sum_{m=1}^{+\infty} z_{i,m}(x, t) q^m$$

untuk

$$z_{i,m} = \frac{1}{m!} \left. \frac{\delta^m \phi_i(x, t; q)}{\delta q^m} \right|_{q=0}$$

Jika, syarat awal dan parameter tambahan h_i , dipilih dengan benar, maka deret diatas akan konvergen ke $q = 1$, sehingga

$$\begin{aligned} u(x, t) &= z_{1,0}(x, t) + \sum_{m=1}^{+\infty} z_{1,m}(x, t) \\ v(x, t) &= z_{2,0}(x, t) + \sum_{m=1}^{+\infty} z_{2,m}(x, t) \end{aligned}$$

Yang mana merupakan salah satu dari solusi – solusi persamaan nonlinear, seperti yang telah dibuktikan oleh Liao. Sekarang didefinisikan vektor

$$\overrightarrow{z_{i,n}} = \{ z_{i,0}(x, t), z_{i,1}(x, t), \dots, z_{i,n}(x, t) \}$$

Sehingga, persamaan deformasi ke- m adalah

$$L[z_{i,m}(x, t) - \chi_m z_{i,m-1}(x, t)] = h_i R_{i,m}(\overrightarrow{z_{i,m-1}}) \quad (23)$$

Dengan syarat awal

$$z_{i,m}(x, 0) = 0 \quad (25)$$

dan

$$\begin{aligned} R_{1,m}(\overrightarrow{z_{1,m-1}}) &= (z_{1,m-1})_t - (z_{2,m-1})_x - z_{1,m-1} + z_{2,m-1} + 2 - 2\chi_m \\ R_{2,m}(\overrightarrow{z_{2,m-1}}) &= (z_{2,m-1})_t + (z_{1,m-1})_x - z_{1,m-1} + z_{2,m-1} + 2 - 2\chi_m \end{aligned}$$

untuk

$$\chi_m = \begin{cases} 0, & m \leq 1 \\ 1, & m > 1 \end{cases}$$

Sekarang, solusi dari persamaan deformasi orde ke- m (23) untuk $m \geq 1$ menjadi

$$z_{i,m}(x, t) = \chi_m z_{i,m-1}(x, t) + h_i \int_0^t R_{i,m}(\overrightarrow{z_{i,m-1}}) d\tau + c_i \quad (26)$$

Dimana konstanta integral c_i ($i = 1, 2$) ditentukan oleh syarat awal (25).

Sehingga untuk $h_i = h$ didapatkan

$$R_{1,1}(\overrightarrow{z_{1,0}}) = -e^x$$

$$z_{1,1}(x, t) = -hte^x$$

$$R_{2,1}(\overrightarrow{z_{2,0}}) = e^x$$

$$z_{2,1}(x, t) = hte^x$$

$$R_{1,2}(\overrightarrow{z_{2,0}}) = hte^x - he^x$$

$$z_{1,2}(x, t) = -\frac{hte^x}{2} [2 - h(t - 2)]$$

$$R_{2,2}(\overrightarrow{z_{2,1}}) = hte^x + he^x$$

$$z_{2,2}(x, t) = \frac{hte^x}{2} [2 + h(t + 2)]$$

$$R_{1,3}(\overrightarrow{z_{1,2}}) = -he^x + 2h^2te^x - h^2e^x - \frac{h^2t^2e^x}{2} + hte^x$$

$$z_{1,3}(x, t) = -\frac{hte^x}{6} [6 - 6h(t - 2) + h^2(t^2 - 6t + 6)]$$

$$R_{2,3}(\overrightarrow{z_{2,2}}) = he^x + 2h^2te^x + h^2e^x + \frac{h^2t^2e^x}{2} + hte^x$$

$$z_{2,3}(x, t) = \frac{hte^x}{6} [6 + 6h(t + 2) + h^2(t^2 + 6t + 6)]$$

$$R_{1,4}(\overrightarrow{z_{1,3}}) = -he^x + 4h^2te^x - 2h^2e^x - \frac{h^3t^2e^x}{2} + 3h^3te^x - h^3e^x - h^2t^2e^x - h^3t^2e^x + hte^x + \frac{h^3t^3e^x}{6}$$

$$z_{1,4}(x, t) = -\frac{hte^x}{24} [24 - 36h(t - 2) + 12h^2(t^2 - 6t + 6) - h^3(t^3 - 12t^2 + 36t - 24)]$$

$$R_{2,4}(\overrightarrow{z_{2,3}}) = he^x + 4h^2te^x + 2h^2e^x + \frac{h^3t^2e^x}{2} + 3h^3te^x + h^3e^x + h^2t^2e^x + h^3t^2e^x + hte^x + \frac{h^3t^3e^x}{6}$$

$$z_{2,4}(x, t) = \frac{hte^x}{24} [24 + 36h(t + 2) + 12h^2(t^2 + 6t + 6) + h^3(t^3 + 12t^2 + 36t + 24)]$$

$$R_{1,5}(\overrightarrow{z_{1,4}}) = -he^x + 6h^2te^x - 3h^2e^x - \frac{9h^3t^2e^x}{2} + 9h^3te^x - 3h^3e^x - 3h^4t^2e^x + 4h^4te^x - h^4e^x$$

$$+ \frac{2h^4t^3e^x}{3} + hte^x - \frac{3h^2t^2e^x}{2} - \frac{h^4t^4e^x}{24} + \frac{h^3t^3e^x}{2}$$

$$z_{1,5}(x, t) = -\frac{hte^x}{120} [120 - 240h(t - 2) + 120h^2(t^2 - 6t + 6) - 20h^3(t^3 - 12t^2 + 36t - 24) + h^4(t^4 - 20t^3 + 120t^2 - 240t + 120)]$$

$$\begin{aligned}
 R_{2,5}(\overrightarrow{z_{2,4}}) &= he^x + 6h^2te^x + 3h^2e^x + \frac{9h^3t^2e^x}{2} + 9h^3te^x + 3h^3e^x + 3h^4t^2e^x + 4h^4te^x + h^4e^x + \frac{2h^4t^3e^x}{3} \\
 &\quad + hte^x + \frac{3h^2t^2e^x}{2} + \frac{h^4t^4e^x}{24} + \frac{h^3t^3e^x}{2} \\
 z_{2,5}(x, t) &= \frac{hte^x}{120} [120 + 240h(t + 2) + 120h^2(t^2 + 6t + 6) + 20h^3(t^3 + 12t^2 + 36t + 24) \\
 &\quad + h^4(t^4 + 20t^3 + 120t^2 + 240t + 120)]
 \end{aligned}$$

Selanjutnya, deret solusi dengan menggunakan metode HAM dapat dituliskan dalam bentuk

$$\begin{aligned}
 u(x, t) &= z_{1,0}(x, t) + z_{1,1}(x, t) + z_{1,2}(x, t) + z_{1,3}(x, t) + z_{1,4}(x, t) + z_{1,5}(x, t) + \dots, \\
 v(x, t) &= z_{2,0}(x, t) + z_{2,1}(x, t) + z_{2,2}(x, t) + z_{2,3}(x, t) + z_{2,4}(x, t) + z_{2,5}(x, t) + \dots,
 \end{aligned}$$

Dalam kasus $h = -1$,

$$z_{1,0}(x, t) = 1 + e^x$$

$$z_{1,1}(x, t) = te^x$$

$$z_{1,2}(x, t) = \frac{t^2e^x}{2} = \frac{t^2e^x}{2!}$$

$$z_{1,3}(x, t) = \frac{t^3e^x}{6} = \frac{t^3e^x}{3!}$$

$$z_{1,4}(x, t) = \frac{t^4e^x}{24} = \frac{t^4e^x}{4!}$$

$$z_{1,5}(x, t) = \frac{t^5e^x}{120} = \frac{t^5e^x}{5!}$$

$$z_{2,0}(x, t) = -1 + e^x$$

$$z_{2,1}(x, t) = -te^x$$

$$z_{2,2}(x, t) = \frac{t^2e^x}{2} = \frac{t^2e^x}{2!}$$

$$z_{2,3}(x, t) = -\frac{t^3e^x}{6} = -\frac{t^3e^x}{3!}$$

$$z_{2,4}(x, t) = \frac{t^4e^x}{24} = \frac{t^4e^x}{4!}$$

$$z_{2,5}(x, t) = -\frac{t^5e^x}{120} = -\frac{t^5e^x}{5!}$$

Sehingga, deret solusinya dapat dituliskan sebagai

$$u(x, t) = 1 + e^x \left(t + \frac{t^2}{2!} + \frac{t^3}{3!} + \frac{t^4}{4!} + \frac{t^5}{5!} + \dots \right),$$

$$v(x, t) = -1 + e^x \left(-t + \frac{t^2}{2!} - \frac{t^3}{3!} + \frac{t^4}{4!} - \frac{t^5}{5!} + \dots \right)$$

Karena $e^t = t + \frac{t^2}{2!} + \frac{t^3}{3!} + \frac{t^4}{4!} + \frac{t^5}{5!} + \dots$ dan $e^{-t} = -t + \frac{t^2}{2!} - \frac{t^3}{3!} + \frac{t^4}{4!} - \frac{t^5}{5!} + \dots$, maka solusi tersebut konvergen ke solusi eksak-nya, yaitu

$$u(x, t) = 1 + e^{x+t} \text{ dan } v(x, t) = -1 + e^{x-t}$$

Selanjutnya, akan dibuktikan bahwa solusi tersebut memenuhi persamaan (17) dan (18).

Untuk persamaan (17),

$$u_t - v_x - (u - v) = -2$$

$$e^{x+t} - e^{x-t} - 1 - e^{x+t} - 1 + e^{x-t} = -2$$

$$-2 = -2 \text{ (Terbukti)}$$

Untuk persamaan (18),

$$v_t + u_x - (u - v) = -2$$

$$-e^{x-t} + e^{x+t} - 1 - e^{x+t} - 1 + e^{x-t} = -2$$

$$-2 = -2 \text{ (Terbukti)}$$

Sehingga, terbukti bahwa solusi $u(x, t) = 1 + e^{x+t}$ dan $v(x, t) = -1 + e^{x-t}$ telah memenuhi sistem persamaan tersebut.

5. SIMPULAN

Dengan menerapkan metode analisis homotopi ke dalam sistem persamaan diferensial parsial linear non homogen $u_t - v_x - (u - v) = -2$, $v_t + u_x - (u - v) = -2$ dengan syarat awal $u(x, 0) = 1 + e^x$ dan $v(x, 0) = -1 + e^x$, didapatkan solusi homotopi yang konvergen ke solusi $u(x, t) = 1 + e^{x+t}$ dan $v(x, t) = -1 + e^{x-t}$ apabila nilai $h = -1$.

Dalam karya tulis ini, metode analisis homotopi hanya diterapkan ke dalam sistem persamaan diferensial parsial linear non homogen orde satu. Oleh karena itu, karya tulis ini diharapkan dapat digunakan sebagai referensi dalam menerapkan metode analisis homotopi ke dalam sistem persamaan diferensial parsial lainnya yang lebih rumit.

KEPUSTAKAAN

- [1] Bronson, R. dan Costa, G. (2007). *Persamaan Differensial*. Erlangga, Jakarta.
- [2] Liao, S. (2012). *Homotopy Analysis Method in Nonlinear Differential Equation*. Higher Education Press, Beijing.
- [3] Bataineh, Noorani, dan Hashim. (2007). Approximate Analytical Solutions of Systems of PDEs by Homotopy Analysis Method. *Journal of Computers and Mathematics with Applications*, Malaysia. 55: 2913–2923.

Penerapan Neural Machine Translation untuk Eksperimen Penerjemahan secara Otomatis pada Bahasa Lampung – Indonesia

Zaenal Abidin

Program studi Sistem Informasi, Fakultas Teknik dan Ilmu Komputer, Universitas Teknokrat Indonesia
Jalan Z.A Pagaralam 9-11 Kedaton Bandar Lampung Telp (0721)702022
www.teknokrat.ac.id, zabin@teknokrat.ac.id

ABSTRAK

Bahasa Lampung adalah bahasa asli masyarakat provinsi Lampung. Siswa-siswi di sekolah mempelajari bahasa Lampung, sebagai muatan lokal, guna pelestarian bahasa tersebut. Salah satu aktifitas pembelajaran bahasa Lampung adalah penerjemahan bahasa Lampung. Penerjemahan dilakukan secara manual karena hanya tersedia kamus bahasa Lampung. Dalam menerjemahkan bahasa Lampung ke bahasa Indonesia secara manual atau kata per kata berpotensi menimbulkan persepsi salah makna karena meniadakan faktor tata bahasa dan konteks kalimat. Neural machine translation (NMT) adalah sebuah pendekatan baru dalam teknologi mesin penerjemah yang bekerja dengan memadukan encoder, sebuah komponen berupa recurrent neural network (RNN) yang mengkodekan bahasa sumber menjadi vektor-vektor yang panjangnya tetap, dan decoder, sebuah komponen berupa recurrent neural network (RNN) yang membangkitkan hasil terjemahan, secara terpadu. Penelitian diawali dengan pembuatan 3000 kalimat paralel bahasa Lampung (dialek api) – Indonesia kemudian dilanjutkan dengan penentuan parameter model NMT untuk proses training data, tahap selanjutnya adalah membangun model NMT dan menguji model NMT. Pengujian model NMT menggunakan 25 kalimat tunggal dan 25 kalimat majemuk tanpa out-of-vocabulary (OOV). Hasil eksperimen menunjukkan bahwa penerjemahan bahasa Lampung-Indonesia pada 25 kalimat tunggal tanpa OOV diperoleh nilai Bilingual Evaluation Under Study (BLEU) sebesar 41.79 dan 25 kalimat majemuk tanpa OOV diperoleh nilai Bilingual Evaluation Under Study (BLEU) sebesar 37.5.

Kata kunci: NMT, RNN, Encoder-Decoder, OOV, BLEU

1. PENDAHULUAN

1.1 LATAR BELAKANG

Badan Perserikatan Bangsa-Bangsa (PBB) menyatakan 3000 bahasa etnis yang ada di seluruh dunia terancam punah. Di Indonesia, dari 617 bahasa daerah yang teridentifikasi terdapat 15 bahasa ibu dinyatakan punah dan 139 lainnya dalam status terancam punah. Informasi ini dinyatakan oleh Badan Pembinaan dan Pengembangan Bahasa Kementerian Pendidikan dan Kebudayaan melalui hasil risetnya. Dalam konteks tersebut, pelestarian bahasa ibu menjadi teramat penting. Sebab, bahasa merupakan identitas penting suatu daerah [1].

Bahasa ibu merupakan bentuk ekspresi kultural utama suatu etnis atau daerah. Tidak terkecuali bahasa Lampung. Bahasa Lampung adalah identitas penting dari kebudayaan Lampung. Jika tidak ada upaya menjaga bahasa Lampung maka akan mengakibatkan sirnanya kebudayaan masyarakat Lampung. Upaya menjaga dan melestarikan bahasa daerah tidak hanya dilakukan oleh pemerintah pusat melainkan membutuhkan peran aktif pemerintah daerah, masyarakat, serta institusi pendidikan atau lembaga pendidikan. Setiap provinsi memiliki bahasa daerah, baik yang berbentuk tulisan maupun ucapan termasuk salah satunya provinsi Lampung yang memiliki bahasa daerah dengan dua dialek utama, yaitu dialek *Api* dan dialek *Nyo*. Upaya pemerintah provinsi Lampung untuk melestarikan bahasa Lampung, salah satunya diberikan melalui pendidikan bahasa Lampung sebagai muatan lokal wajib dari mulai sekolah tingkat dasar sampai sekolah tingkat menengah atas. Hal ini sesuai peraturan gubernur nomor 39 tahun 2014 tentang mata pelajaran bahasa dan aksara Lampung sebagai muatan lokal wajib pada jenjang satuan pendidikan dasar sampai menengah atas dan didukung oleh ketersediaan buku ajar mulai dari SD, SMP dan SMA berikut kamus bahasa Lampung.

Pemanfaatan kakas berbasis komputer dapat digunakan sebagai sebuah upaya untuk melestarikan bahasa Lampung secara digital sesuai perkembangan teknologi saat ini yaitu kamus digital dan mesin penerjemah. Kamus digital memiliki keterbatasan dalam menerjemahkan bahasa daerah ke bahasa Indonesia karena pendekatan yang digunakan adalah menerjemahkan kata per kata, contoh sebuah kakas kamus digital yang dapat diakses melalui laman www.kamusdaerah.com. Alternatif lain dari kamus adalah mesin penerjemah.

Sebuah kakas mesin penerjemah yang dikenal oleh masyarakat awam adalah mesin penerjemah *google*. Mesin ini dapat diakses dengan mudah, akan tetapi pada mesin penerjemah *google* hanya terdapat dua bahasa daerah yang tersedia yaitu bahasa Jawa dan bahasa Sunda. Bahasa Lampung belum tersedia di mesin penerjemah *google*. Sebelum November 2016, *Statistical Machine Translation* (SMT) banyak digunakan sebagai standar mesin penerjemah *google* karena memiliki beberapa kelebihan.

Kelebihan SMT dibanding dengan *rule-based machine translation* adalah pada fleksibilitas dan *robustness*. Kekurangan yang dimiliki oleh mesin penerjemah berbasis SMT, yaitu (1) pada SMT, sulit untuk memperoleh pemetaan kata yang tepat, (2) pada SMT, sulit untuk menentukan kandidat target frase terbaik dari masukan frase yang diberikan dan konteks frase yang berbeda akan memberikan makna yang berbeda, (3) pada SMT, pekerjaan yang sulit untuk memprediksi struktur derivasi terjemahan, (4) pada SMT, sulit untuk mempelajari model bahasa yang baik. Pada aspek yang paling mendasar, kekurangan SMT adalah pada *data sparseness* dan kurangnya pemodelan semantik pada kata-kata baik berupa frase atau kalimat. Alternatifnya, *deep neural networks* (DNNs) pandai dalam mempelajari representasi semantik dan pemodelan konteks yang luas. Khusus pada RNNs memiliki kapabilitas dalam merepresetasikan *history words* atau mengingat kata-kata yang telah dipelajarinya [2].

Neural machine translation (NMT) adalah sebuah pendekatan baru dalam penerjemahan mesin yang menggunakan arsitektur RNNs pada bagian *encoder* dan *decoder*-nya. [3] telah membuktikan performa NMT lebih baik dibanding SMT pada eksperimen 30 jenis penerjemahan [3]. NMT yang digunakan saat ini oleh para peneliti mesin penerjemah adalah NMT berbasis *attention*, sebagai hasil penelitian yang telah dilakukan oleh Bahdanau dan kawan-kawan di tahun 2015 [4].

Peneliti dari Indonesia tentang mesin penerjemah yang menggunakan RNN, yaitu (1) [7] telah melakukan penelitian dan melaporkannya dalam sebuah *paper* berjudul *Recurrent Neural Network Language Model for English-Indonesia Machine Translation: Experimental Study*, (2) Adiputra dan Arase (2017) telah melakukan penelitian dan melaporkannya dalam sebuah *paper* berjudul *Performance of Japanese-to-Indonesian Machine Translation on Different Models*. Bahasa daerah sebagai objek penelitian mesin penerjemah menggunakan NMT khususnya RNNs belum pernah ada.

Penelitian ini adalah sebuah upaya untuk ikut melestarikan bahasa Lampung melalui eksperimen penerjemahan bahasa Lampung – Indonesia dengan mesin penerjemah berbasis NMT dengan menggunakan masukan berupa

sebuah kalimat atau beberapa kalimat untuk kemudian diproses secara komputasi guna mendapatkan hasil berupa kalimat hasil terjemahan dari bahasa Lampung ke bahasa Indonesia.

1.2 RUMUSAN MASALAH

Berdasarkan pemaparan pada latar belakang yang telah dijelaskan diatas, beberapa rumusan masalah yang akan dibahas pada penelitian ini adalah:

1. Masalah apa saja yang ditemukan saat melakukan penerjemahan bahasa Lampung - Indonesia secara manual
2. Bagaimana mengukur hasil penerjemahan mesin penerjemah bahasa Lampung - Indonesia

1.3 TUJUAN PENELITIAN

Tujuan penelitian yang berlandaskan latar belakang adalah :

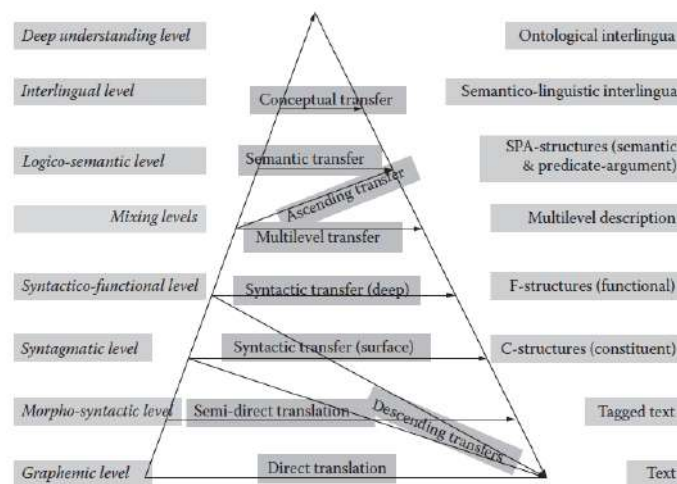
1. Mengidentifikasi masalah penerjemahan bahasa Lampung – Indonesia secara manual.
2. Membangun model NMT untuk penerjemahan bahasa Lampung – Indonesia.
3. Mengukur hasil terjemahan dari model NMT untuk penerjemahan bahasa Lampung – Indonesia

2. LANDASAN TEORI

2.1 MESIN PENERJEMAH

Paradigma penguasa mesin terjemahan saat ini sebelum datangnya *neural machine translation* (NMT) adalah *statistical machine translation* (SMT). Sebelumnya ada mesin terjemahan berbasis contoh / *example-based machine translation* (EBMT) diperkenalkan pada awal tahun 1980an. Kedua paradigma ini baik SMT dan EBMT bergantung pada ketersediaan contoh terjemahan yang disebut paralel korpus.

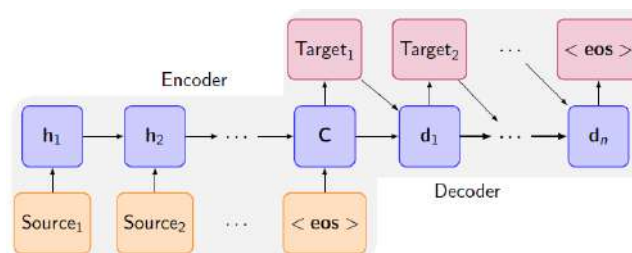
Pendekatan mesin penerjemah telah dikelompokkan pada segitiga yang terkenal yaitu segitiga Vauquois atau piramida Vauquois. Informasi yang tergambarkan dari segitiga Vauquois ini adalah dalam penerjemahan dari suatu bahasa sumber ke bahasa tujuan akan membutuhkan banyak operasi. Sisi kiri segitiga adalah sisi menaik dan sisi kanan adalah sisi turun. Sudut kiri tempat bahasa sumber berada dan sudut kanan tempat bahasa target berada. Saat kita naik ke sisi kiri segitiga tersebut, kita akan melakukan analisis berbagai jenis pada kalimat bahasa sumber. Kalimat dari bahasa sumber berpotensi akan melalui proses analisis morfologi, POS Tag, identifikasi kelompok kata kerja / kata benda, *parsing, semantics, discourse and co-reference*. Kalimat bahasa sumber yang telah melewati sisi kiri selanjutnya akan berpindah ke sisi kanan yaitu sisi membangkitkan kalimat pada bahasa tujuan [5].



Gambar 1. Segitiga Vauquois mengekspresikan pendekatan terhadap mesin penerjemah (Bhattacharyya, 2015)

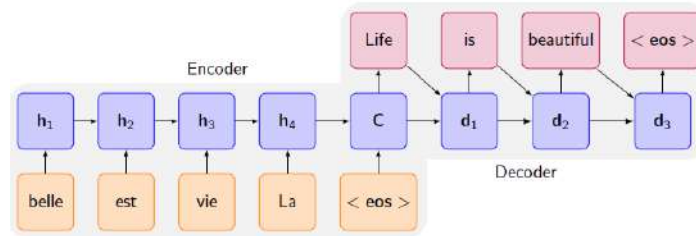
2.1 NEURAL MACHINE TRANSLATION

Untuk melakukan *machine translation* (MT) dengan jaringan saraf maka perlu untuk menghubungkan RNN *encoder* dengan RNN *decoder* model bahasa. RNN *encoder* memproses kalimat dalam kata bahasa sumber dari kata demi kata dan menghasilkan representasi kalimat masukan. RNN *decoder* atau model bahasa mengambil *output* dari *encoder* sebagai masukan dan menghasilkan terjemahan kata masukan kata demi kata. Gambar 2 mengilustrasikan bagaimana menghubungkan model *encoder* dan *decoder* secara bersamaan. Arsitektur model untuk terjemahan mesin dikenal sebagai *model encoder-decoder* [6].



Gambar 2. Penerjemahan secara *sequence to sequence* dengan menggunakan arsitektur model *encoder-decoder*

Gambar 3 mengilustrasikan bagaimana sistem MT *encoder-decoder* akan menghasilkan terjemahan bahasa Inggris dari sebuah kalimat bahasa Prancis. *Encoder* memproses kata kalimat Prancis dengan kata termasuk simbol $<eos>$, akhir dari urutan. Perhatikan bahwa dalam contoh ini kita melewati kalimat sumber di belakang, melakukan hal ini telah ditemukan memberikan hasil terjemahan yang lebih baik. Kami kemudian melewati representasi sumber kalimat yang dikodekan ke *decoder* atau model bahasa dan kami membiarkan model bahasa ini menghasilkan kata terjemahan berdasarkan kata sampai menghasilkan simbol $<eos>$, akhir dari urutan [6].



Gambar 3. Contoh penerjemahan dengan menggunakan arsitektur model encoder-decoder

Keuntungan dari arsitektur *encoder-decoder* adalah sistem memproses keseluruhan masukan sebelum mulai menerjemahkan. Ini berarti bahwa *decoder* dapat menggunakan apa yang telah dihasilkan dan keseluruhan kalimat sumber saat membuat kata berikutnya dalam terjemahan.

2.3 EVALUASI MESIN PENERJEMAH

Evaluasi hasil penerjemahan dilakukan dengan membandingkan kalimat hasil penerjemahan dengan kalimat acuan dengan menggunakan *Bilingual Evaluation Understudy* (BLEU). BLEU adalah algoritma yang berfungsi untuk mengevaluasi kualitas dari hasil terjemahan yang telah diterjemahkan oleh mesin dari suatu bahasa sumber ke bahasa tujuan. BLEU mengukur skor presisi berbasis statistik yang telah dimodifikasi antara hasil terjemahan, secara otomatis, dengan terjemahan rujukan dengan menggunakan konstanta yang dinamakan *brevity penalty* (BP) [7].

$$BP_{BLEU} = \begin{cases} 1 & \text{jika } c > r \\ e^{(1-\frac{r}{c})} & \text{jika } c \leq r \end{cases} \quad (1)$$

$$p_n = \frac{\sum_{C \in \{Candidates\}} \sum_{n\text{-gram} \in C} Count_{clip}(n\text{-gram})}{\sum_{C' \in \{Candidates\}} \sum_{n\text{-gram}' \in C'} Count_{clip}(n\text{-gram}')} \quad (2)$$

$$BLEU = BP \cdot \exp(\sum_{n=1}^N w_n \cdot \log p_n) \quad (3)$$

Persamaan (3) jika disajikan dalam domain log akan menjadi seperti persamaan (4) di bawah ini.

$$\log BLEU = \min\left(1 - \frac{r}{c}, 0\right) + \sum_{n=1}^N w_n \cdot \log p_n \quad (4)$$

BP adalah simbol *brevity penalty*, c adalah jumlah kata-kata dari hasil penerjemahan mesin, r adalah panjang terjemahan rujukan yang efektif. p_n adalah rata-rata geometris dari presisi n -gram yang dimodifikasi. Nilai N yang digunakan adalah $N = 4$ dan $w_n = \frac{1}{N}$ [8].

3. METODOLOGI PENELITIAN

3.1 ANALISIS MASALAH PENERJEMAHAN BAHASA LAMPUNG

Bahasa Lampung adalah bahasa asli masyarakat Lampung. Di provinsi Lampung secara umum terbagi menjadi dua dialek utama yaitu dialek *Api* dan dialek *Nyo*. Pada penelitian ini hanya dilakukan pengamatan secara seksama terhadap bahasa Lampung dialek *Api*. Hasil wawancara dengan ahli bahasa Lampung menunjukkan bahasa Lampung memiliki kemiripan struktur tata bahasa dengan bahasa Indonesia. Kosakata dalam bahasa Lampung jumlahnya lebih sedikit dibanding dengan bahasa Indonesia. Sebuah kamus bahasa Lampung – Indonesia yang diterbitkan oleh penerbit Gunung Pesagi Bandar Lampung berisi 4800 kata sedangkan kamus bahasa Indonesia edisi ke-5 yang diterbitkan oleh penerbit Gramedia pustaka utama memiliki 118000 kata. Dalam menerjemahkan bahasa Lampung ke bahasa Indonesia, tidak tepat menterjemahkan kata per kata. Hal tersebut akan menimbulkan persepsi salah makna / arti serta meniadakan faktor tata bahasa serta konteks kalimat. Beberapa contoh kalimat dalam bahasa Lampung dan terjemahannya dapat dilihat pada tabel 1 di bawah ini.

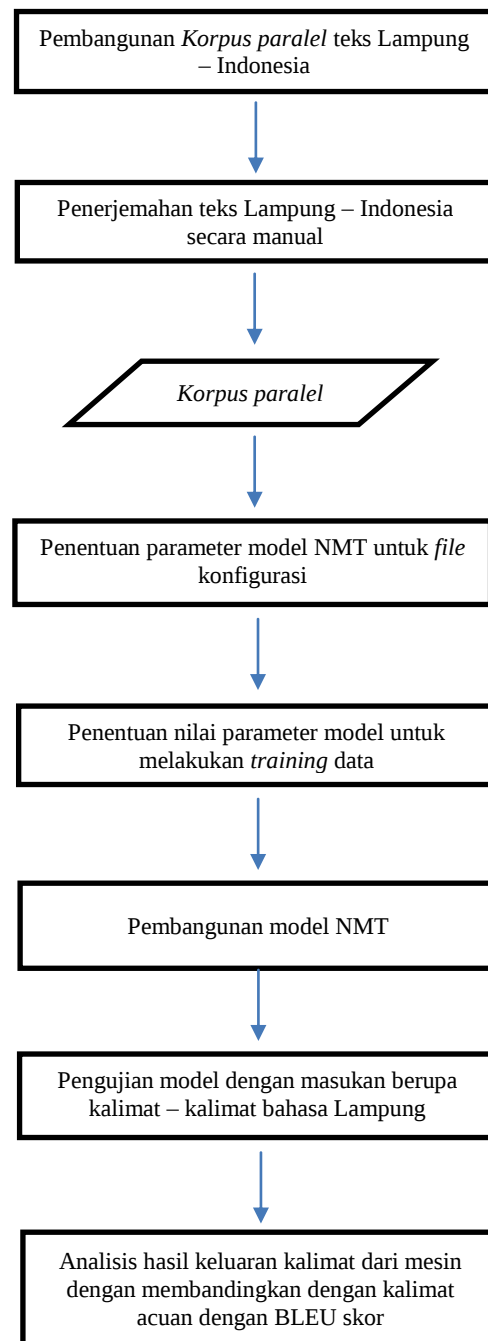
Tabel 1. Contoh Korpus Paralel Bahasa Lampung - Indonesia

Bahasa Lampung	Bahasa Indonesia
Ani tulung akukko wai sina	Ani tolong ambilkan air itu
Makku lagi mepoh kawai di wai	Ibuku sedang mencuci baju di sungai
Gham lapah mit pasagh	Kita pergi ke pasar
Lapah gham sholat di masjid	Mari kita sholat di masjid
Ketua kelas sai ngator barisan di kelasni	Ketua kelas yang mengatur barisan di kelasnya
Ya tetap di kelas sai ulah mak cakak kelas telu kali	Dia tetap di kelas satu karena tidak naik kelas tiga kali

Pada tabel 1 di atas terdapat beberapa kata yang memiliki makna berbeda tergantung konteks kalimat atau posisi pada kalimat, seperti kata *wai*, *lapah*, *sai*. Penggunaan kamus bahasa Lampung harus menerjemahkan satu per satu dan memilih makna kata yang tepat secara manual agar kalimat tersebut dapat dipahami dengan baik. Penerjemahan kalimat bahasa Lampung sangat penting memperhatikan konteks kalimat yang digunakan agar dapat menghasilkan terjemahan yang sesuai. Pada bahasa Lampung juga terdapat beberapa kata yang memiliki satu makna yaitu *lapah*, *mit*, *tandak*, *lijing* yang semuanya bermakna *pergi*. Beberapa partikel seperti *kudo*, *do*, *pai*, *tah* semuanya memiliki satu makna yaitu *lah*, sebuah partikel dalam bahasa Indonesia.

3.2 DESAIN PENELITIAN

Pada eksperimen penerjemahan bahasa Lampung – Indonesia dengan pendekatan NMT, langkah-langkah yang dilakukan disajikan pada diagram alir di bawah ini.



Gambar 4. Tahapan Pembangunan Mesin Translasi Bahasa Lampung – Indonesia

3.3 FILE KONFIGURASI

Pada eksperimen ini terdapat 3000 kalimat dalam bahasa Lampung dialek *api* dan 3000 kalimat terjemahannya dalam bahasa Indonesia. Eksperimen ini akan melibatkan dua jenis kumpulan data yang diambil dari 3000 kalimat tersebut yaitu data latih dan data validasi. Data latih dan validasi yaitu satu daftar kalimat sumber dipasangkan dengan satu atau beberapa target terjemahan yang digunakan untuk pemilihan model dan optimisasi *hyper-parameter*. Data pengujian yaitu satu daftar kalimat sumber dipasangkan dengan satu atau beberapa target

terjemahan yang digunakan untuk mengevaluasi kinerja terjemahan pada teks kalimat yang belum pernah diujicobakan.

Tabel 2 Parameter yang digunakan pada konfigurasi file untuk proses training data

Parameter NMT	Keterangan
source vocabulary size	Ukuran kosa kata pada bahasa sumber
target vocabulary size	Ukuran kosa kata pada bahasa tujuan
source word embedding dimension	Dimensi <i>word embedding</i> pada bahasa sumber
target word embedding dimension	Dimensi <i>word embedding</i> pada bahasa tujuan
encoder hidden layer dimension	Dimensi <i>hidden layer</i> pada bagian <i>encoder</i>
decoder hidden layer dimension	Dimensi <i>hidden layer</i> pada bagian <i>decoder</i>
MRT sample size	Jumlah kalimat yang dijadikan sampel pada model <i>training</i> MRT
MRT length ratio limit	Panjang rasio limit kalimat yang disampel
maximum sentence length	Batas panjang kalimat yang berada di data pelatihan
mini-batch size	Jumlah kalimat dalam sebuah <i>mini-batch</i>
mini-batch sorting size	Jumlah <i>mini-batch</i> yang diurutkan
iteration limit	Batasan iterasi pada <i>training</i>
convergence limit	Jumlah iterasi dimana program akan berhenti sebelum konvergen dimana tidak terjadi kenaikan nilai BLEU pada batasan jumlah iterasi tersebut
Optimizer	THUMT mendukung tiga metode optimisasi. Nilai 0 untuk stochastic gradient descent (SGD), 1 untuk AdaDelta (Zeiler, 2012), 2 untuk Adam (Kingma and Ba, 2014). THUMT merekomendasi menggunakan Adam.
Clip	Gradien kliping untuk mengatasi masalah ledakan gradien.
SGD learning rate	Tingkat pembelajaran / learning rate pada SGD.
AdaDelta rho	Tingkat kerusakan / decay rate pada AdaDelta.
AdaDelta epsilon	Konstanta digunakan untuk kondisi denominator yang lebih baik di AdaDelta.
Adam alpha	Step size di Adam. THUMT merekomendasikan setting hyper-parameter ini pada 0.0005 untuk MLE, 0.00001 untuk MRT, dan 0.00005 untuk SST.
Adam alpha decay	Tingkat peluruhan eksponensial untuk ukuran langkah di Adam.
Adam beta 1	Tingkat peluruhan eksponensial untuk perkiraan waktu di Adam.
Adam beta 2	Tingkat peluruhan eksponensial lainnya untuk perkiraan momen di Adam.
Adam epsilon	Sebuah konstanta kecil di Adam.
beam size	Beam size pada saat decoding.
model dumping iteration	Interval untuk membuang dan memvalidasi intermediate models.
checkpoint iteration	Interval untuk menyimpan pos pemeriksaan yang digunakan untuk melanjutkan pelatihan saat terganggu

Skema validasi data dilakukan dengan *5-fold validation* yaitu suatu skema melakukan percobaan berulang-ulang sebanyak lima kali guna mendapatkan informasi kelayakan data yang digunakan sebagai Korpus paralel. Data pengujian yaitu satu daftar kalimat sumber dipasangkan dengan satu atau beberapa target terjemahan yang digunakan untuk mengevaluasi kinerja terjemahan pada teks kalimat yang belum pernah diujicobakan.

Tabel 3 5-folds validation padaKorpus Paralel Bahasa Lampung – Indoesia

Urutan Kalimat	Fold1	Fold2	Fold3	Fold4	Fold5
Data Latih	601 – 3000	1 – 600 & 1201 – 3000	1 -1200 & 1801 – 3000	1 – 1800 & 2401 – 3000	1 – 2400
Data Validasi	601 – 3000	1 – 600 & 1201 – 3000	1 -1200 & 1801 – 3000	1 – 1800 & 2401 – 3000	1 – 2400
Data Uji	1 – 600	601 – 1200	1201 – 1800	1801 – 2400	2401 – 3000

3.4 NMT dengan Kakas THUMT

3.4.1 Proses *Training* dan *Validation Data* NMT pada Kakas THUMT

Pada bagian di bawah disajikan *syntax* / perintah yang digunakan dalam THUMT untuk melakukan proses *training, validation* [9].

Usage: trainer [--help] ...

Required arguments:

--config-file <file> configuration file
 --trn-src-file <file> training set, source file
 --trn-trg-file <file> training set, target file
 --vld-src-file <file> validation set, source file
 --vld-trg-file <file> validation set, target file
 --device {cpu, gpu0, ...} device

Optional arguments:

--training-criterion {0,1,2} training criterion
 0: MLE (default)
 1: MRT
 2: SST
 --replace-unk {0,1} replacing unknown words
 0: off
 1: on (default)
 --save-all-models {0,1} saving all intermediate models
 0: the best model (default)
 1: all intermediate models
 --mono-src-file <file> monolingual source file
 --mono-trg-file <file> monolingual target file
 --init-model-f-file <file> initialization model file
 --debug {0,1} displaying debugging info
 0: off (default)
 1: on
 --help displaying this message

3.4.2 Proses *Testing Data* NMT pada Kakas THUMT

Pada bagian di bawah disajikan *syntax* / perintah yang digunakan dalam THUMT untuk melakukan *testing data* [9].

```
Usage: test.py [--help] ...

Required arguments:
--model-file <file>                model file
--test-src-file <file>             test set, source file
--test-trg-file <file>            translation of the test set
--device {cpu,gpu0,...}           the device for running this script

Optional arguments:
--test-ref-file <file>            test set, reference file(s)
--replace-unk                     replacing unknown words
                                0: off
                                1: on (default)
--help                            displaying helping message
```

3.4.3 Visualisasi Hasil Uji NMT pada Kakas THUMT

Pada bagian di bawah disajikan *syntax* / perintah yang digunakan dalam THUMT untuk melakukan visualisasi hasil penerjemahan [9].

```
Usage: visualize.py [--help] ...

Required arguments:
--model-file <file>                model file
--input-file <file>               input source file
--output-dir <dir>                output directory
--device {cpu,gpu0,...}           the device for running this script

Optional arguments:
--help                            displaying helping message
```

4. HASIL DAN PEMBAHASAN

4.1 KORPUS PARALEL DAN *FILE* KONFIGURASI

Langkah awal yang telah dilakukan dalam proses pembangunan model penerjemahan bahasa Lampung – Indonesia adalah menyiapkan korpus paralel yaitu teks padanan kalimat bahasa Lampung – Indonesia. Kalimat-kalimat tersebut telah dilakukan prapemrosesan yaitu dengan melakukan pemisahan menggunakan spasi termasuk tanda baca seperti koma, tanda tanya dan titik. Korpus paralel yang telah melewati tahap prapemrosesan maka akan digunakan pada tahap *training data*.

Training data dilakukan dengan dengan tujuan untuk mendapatkan model translasi dan kosa kata bahasa Lampung dan bahasa Indonesia. Informasi jumlah kemunculan setiap kata yang terdapat pada Korpus paralel dapat diperoleh melalui tahapan *training data*. Model translasi yang didapatkan digunakan untuk melakukan validasi data, melalui percobaan menerjemahkan serangkaian kalimat dari bahasa Lampung ke bahasa Indonesia dan membandingkan hasilnya terhadap terjemahan acuan yang telah disediakan. Nilai skor BLEU digunakan

untuk menilai hasil terjemahan dari sistem dengan hasil terjemahan acuan. Pada eksperimen penerjemahan bahasa Lampung – Indonesia dengan menggunakan model NMT membutuhkan penetapan nilai-nilai pada banyak parameter yang digunakan dan disatukan dalam sebuah *file* konfigurasi.

Pengaturan *file* konfigurasi, yang digunakan pada model NMT, dimulai dari penetapan ukuran maksimal kosakata pada bahasa Lampung dan bahasa Indonesia. Selama masa tahapan *training data*, pendugaan parameter yang digunakan adalah pendugaan *maximum likelihood*. Kosakata bahasa Lampung dan bahasa Indonesia yang terdapat di paralel korpus tidak melebihi 5000 kata. Nilai-nilai dimensi pada arsitektur jaringan yang digunakan pada model NMT *unreplacing unknown words* dan model NMT *replacing unknown words* meliputi dimensi *word embedding* pada bahasa Lampung dan bahasa Indonesia yaitu 310. Kerangka kerja yang digunakan adalah *encoder* dan *decoder RNNsearch*. Pada *encoder RNNsearch* terdiri dari *forward RNN* dan *backward RNN* dengan masing-masing bernilai dimensi 500 *hidden unit*. Pada *decoder* juga bernilai dimensi 500 *hidden unit*.

Untuk melakukan *training data* pada setiap model digunakan algoritma *minibatch stochastic gradient descent* (SGD) bersama Adadelta. Setiap arahan *update* SGD dihitung menggunakan *minibatch* 80 kalimat. Panjang kalimat yang bisa dimasukan pada sistem adalah 50 kata. Lama *training data* dengan masing-masing 4600 iterasi, menghabiskan alokasi waktu antara 4 sampai 5 jam pada setiap model. *Optimizer* yang digunakan adalah *Adam optimizer*. Selama sebuah model dilakukan *training*, sebuah algoritma *beam search* digunakan untuk memaksimalkan aproksimasi nilai peluang bersyarat, dengan ukuran *beam* yang digunakan adalah 20.

Pada pelaksanaan *training data* sebuah model, di kasus THUMT terdapat mekanisme pengecekan nilai validasi pada titik iterasi tertentu guna mengetahui nilai BLEU yang diperoleh pada titik iterasi tersebut. Di penelitian ini ditetapkan pada titik iterasi dengan nilai 1500.

4.2 HASIL EKSPERIMEN PADA NMT

Eksperimen NMT menghabiskan waktu lebih dari 20 jam untuk proses *training data* secara berulang-ulang sebanyak lima kali. Mekanisme pengulangan lima kali tersebut guna mencoba melakukan percobaan penerjemahan dengan 600 kalimat yang selalu berbeda-beda pada tiap *fold*. Jumlah iterasi yang digunakan rata-rata menggunakan 4600 iterasi kecuali pada *fold3* yaitu 4500. Pada *fold3* terjadi kondisi yang tidak optimal dengan munculnya sebuah peringatan pada layar monitor bertuliskan *NaN* atau *Not a Number* pada saat iterasi antara 4500 sampai 4600. Eksperimen diulang kembali pada *fold3* dengan mencoba berbagai nilai iterasi yaitu 3000, 3500 dan 4500 iterasi. Pada *fold3* titik optimal iterasi terdapat pada angka 4500 iterasi.

Hasil eksperimen NMT pada tahap validasi menghabiskan waktu lebih dari 20 jam untuk proses *training data* dirangkum pada tabel IV.2 di bawah ini. Pada saat dilakukan uji validasi, nilai BLEU *validation*, pada masing-masing *fold* dengan menggunakan kalimat yang sama seperti pada data *training*, NMT konvensional rata-rata memberikan nilai BLEU diatas 95 dikarenakan kalimat-kalimat yg digunakan telah dipelajari oleh sistem NMT dengan baik. Pada setiap *fold* dilakukan pengujian menggunakan *data testing* sebanyak 600 kalimat. Enam ratus

kalimat yang digunakan sebagai *data testing* pada setiap *fold*-nya berisi kalimat-kalimat yang berbeda-beda pada setiap *fold*-nya dan memberikan hasil nilai BLEU *testing* yang berbeda-beda. Nilai BLEU pada fold1, fold2 dan fold5 sangat kecil dibandingkan fold3 dan fold4 dikarenakan terdapat OOV lebih dari 45 % dari kosa kata yang digunakan pada kalimat *data testing*.

Tabel 4. Hasil 5-fold validation sebagai validasi pada NMT

Jumlah Kalimat	Fold1	Fold2	Fold3	Fold4	Fold5	Rata-rata Fold
Data Training Kalimat	2400	2400	2400	2400	2400	2400
Data Validation Kalimat	2400	2400	2400	2400	2400	2400
Data Testing Kalimat	600	600	600	600	600	600
BLEU Validation (Close experiment)	94,57	98,6	98,74	98,73	98,2	97,525
BLEU Testing (Open experiment)	5,51	7,28	15,10	12,15	7,98	8,23

Lima puluh kalimat yang digunakan pada pengujian pada NMT dengan rincian terdiri dari 25 kalimat tunggal tanpa OOV dan 25 kalimat majemuk tanpa OOV. Hasil pengujian disajikan pada tabel 5. Nilai-nilai BLEU yang diperoleh pada tabel 5 juga sangat dipengaruhi oleh tingginya angka persentase kosa kata yang sangat jarang digunakan pada paralel korpus.

Tabel 5. Hasil Eksperimen pada NMT

Hasil Eksperimen	BLEU
Kalimat Tunggal Tanpa OOV	41,79
Kalimat Majemuk Tanpa OOV	37,5

4.3 VISUALISASI DAN ANALISIS HASIL EKSPERIMEN NMT

Sampel yang digunakan pada bagian visualisasi dan analisis hasil eksperimen NMT yaitu 1 kalimat dari 25 kalimat tunggal tanpa OOV. Penyajian sampel tersebut disajikan dalam bentuk tabel untuk memudahkan dalam proses analisisnya. Sampel pertama bahasa Lampung pada kalimat tunggal tanpa OOV adalah kalimat '*sanak sai disuntik sina mak sihat*'. Tabel 6 memberikan informasi dari hasil penerjemahan oleh mesin bahwa sistem NMT gagal dalam menerjemahkan kata ke-3, dan menghilangkan keberadaan dari kata ke-4, ke-5 dan ke-6.

Tabel 6. Sampel pertama kalimat tunggal tanpa OOV pada uji NMT

Bahasa Lampung	sanak sai disuntik sina mak sihat
Bahasa Indonesia	anak yang disuntik itu tidak sehat
Hasil NMT Konvensional	anak yang termahal

Merujuk pada tabel 6 salah satu penyebab ketidakmampuan sistem NMT konvensional ini adalah ada 2 kata yang terindikasi sebagai OOV, yaitu kata *disuntik* dan *sihat*, walaupun dua kata tersebut ada dalam korpus tetapi frekuensi kemunculannya yang sangat rendah sehingga sistem membacanya sebagai *unknown word*.

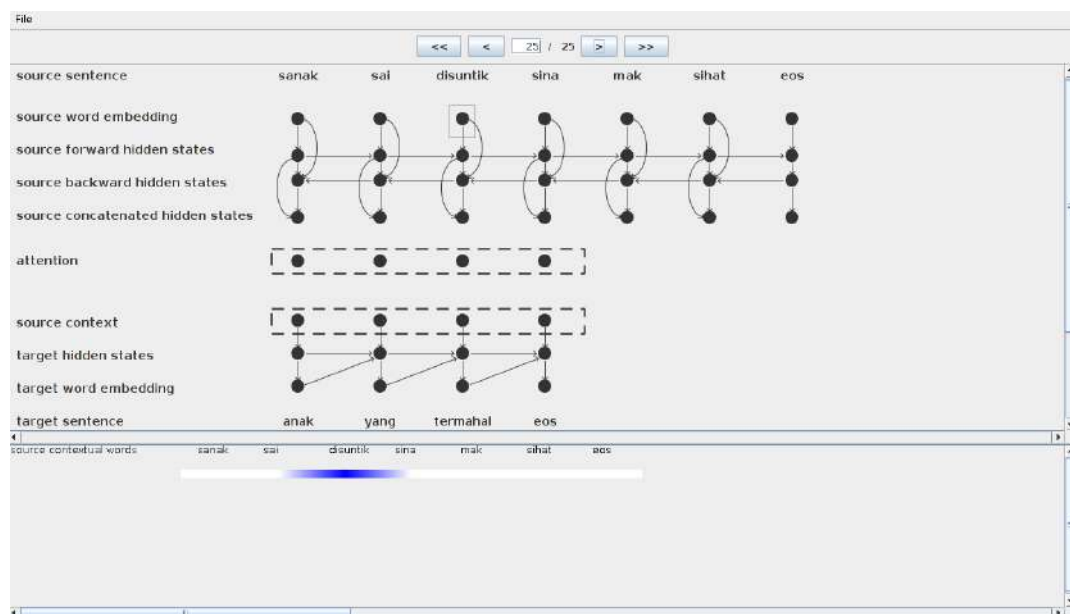
Kalimat dalam bahasa Lampung :*sanak sai disuntik sina mak sihat*

Sistem NMT membacanya sebagai berikut :*sanak sai unk₁ sina mak unk₂*

Tabel 7 Frekuensi kemunculan kata pada sampel pertama kalimat tunggal tanpa OOV dalam korpus paralel

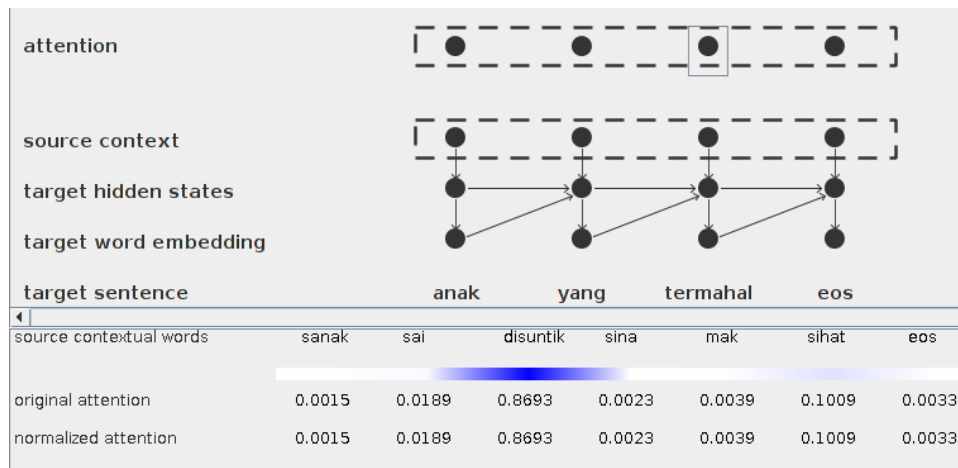
Kata	Frekuensi kemunculan dalam korpus
Disuntik	2
sina	152
mak	202
Sihat	10

Sampel pertama bahasa Lampung pada kalimat tunggal tanpa OOV adalah kalimat '*sanak sai disuntik sina mak sihat*' akan dianalisis melalui visualisasi guna menginterpretasikan apa yang terjadi pada bagian internal dari NMT pada sampel kalimat pertama ini. Kalimat '*sanak sai disuntik sina mak sihat*' artinya '*anak yang disuntik itu tidak sehat*', hasil NMT memberikan hasil '*anak yang termahal*'.



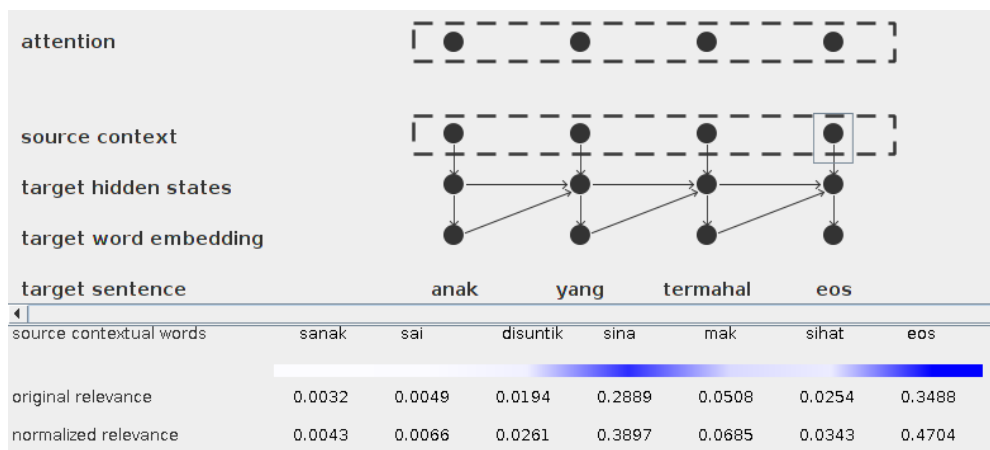
Gambar 5. Visualisasi hasil dari NMT konvensional pada sampel kalimat pertama

Kesalahan penerjemahan terjadi pada kata ketiga dalam bahasa Indonesia '*termahal*' adalah terjadinyaword omission. Informasi dari gambar 5 simbol '*eos*' muncul terlalu cepat dari yang diharapkan. Visualisasi lebih detail diamati melalui *target hidden state*, nilai informasi *attention* paling dominan pada kata '*disuntik*', seperti tergambar pada gambar 6 di bawah ini. Hal ini menunjukkan dengan benar bahwa kata '*disuntik*' penting untuk membangkitkan kata setelahnya.



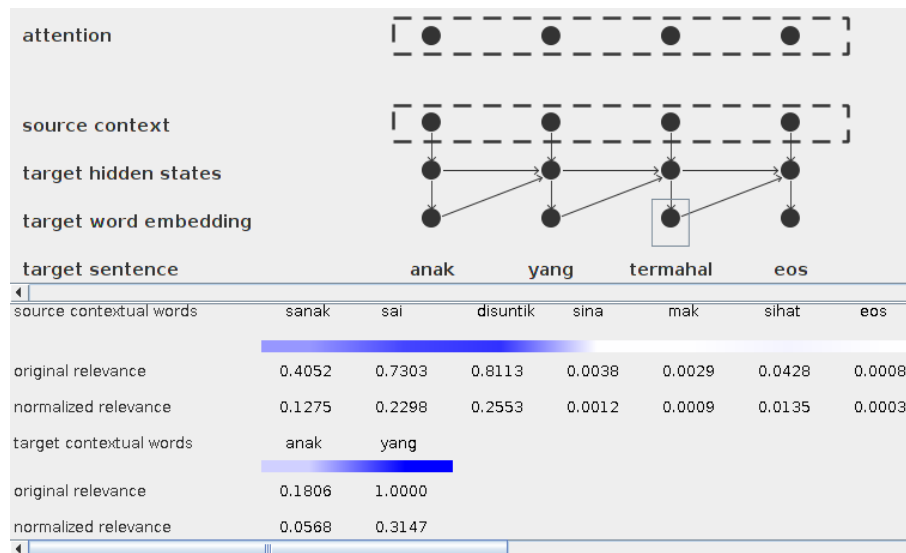
Gambar 6. Visualisasi pada bagian attention pada sampel kalimat pertama kalimat tunggal tanpa OOV

Nilai relevansi pada *source context* menunjukkan nilai terbesar pada bagian akhir kalimat 'eos'.



Gambar 7. Visualisasi pada bagian source context di bagian 'eos' pada sampel kalimat pertama

Akhirnya nilai relevansi dari *target word* pada kata 'termahal' mengungkap fakta bahwa kata 'termahal' paling dominan untuk membangkitkan *target word* pada *softmax layer*, seperti terlihat pada gambar 8 di bawah ini.



Gambar 8. Visualisasi nilai relevansi pada kata 'termahal' pada sampel kalimat pertama

5.SIMPULAN

Penerjemahan bahasa Lampung ke bahasa Indonesia secara otomatis dengan menggunakan *Neural Machine Translation* memberikan hasil penerjemahan bahasa Lampung-Indonesia pada 25 kalimat tunggal tanpa OOV diperoleh nilai Bilingual Evaluation Under Study (BLEU) sebesar 41.79 dan 25 kalimat majemuk tanpa OOV diperoleh nilai Bilingual Evaluation Under Study (BLEU) sebesar 37.5. NMT mampu mengatasi makna kontekstual yang terdapat pada bahasa Lampung seperti pada beberapa kata yang memiliki makna berbeda tergantung konteks kalimat atau posisi pada kalimat, seperti kata *wai*, *lapah*, *sai*. Tingkat kemunculan kata-kata pada korpus paralel akan memberikan pengaruh yang signifikan sehingga kata-kata yang jarang muncul akan terdeteksi sebagai OOV dan sistem NMT gagal dalam menterjemahkannya. Penelitian ini dapat dilanjutkan pada kalimat tunggal dan majemuk yang mengandung OOV.

6. UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada ketua yayasan pendidikan Teknokrat dan Rektor Universitas Teknokrat Indonesia yang telah memberikan kesempatan kepada penulis untuk melanjutkan studi jenjang magister Informatika di STEI ITB sehingga penulis mendapatkan kesempatan menulis sebuah karya tentang *Neural Machine Translation*, kemudian kepada Dekan FTIK dan ketua LPPM Universitas Teknokrat Indonesia yang selalu memberi dukungan kepada penulis untuk selalu berkarya lewat karya ilmiah.

KEPUSTAKAAN

- [1] Anonymous. (2017). *Nestapa Guru Bahasa Lampung* [diakses 1Agustus 2017]. Tersedia dari <http://www.lampost.co/berita-nestapa-guru-bahasa-lampung>
- [2] Zhang, J., & Zong, C. (2015). Deep Neural Networks in Machine Translation: An Overview. *IEEE Intelligent Systems*. Vol.30, Issue: 5, pp. 16 – 25.

- [3] Junczys-Dowmunt, M., Dwojak, T., & Hoang, H. (2016). *Is Neural Machine Translation Ready for Deployment ? A Case Study on 30 Translation Directions* [diakses 1 Agustus 2017]. Tersedia dari https://workshop2016.iwslt.org/downloads/IWSLT_2016_paper_4.pdf
- [4] Bahdanau, D., Cho, K., & Bengio, Y. (2015). Neural Machine Translation by Jointly Learning to Align and Translate. *Proceeding of the 3rd International Conference on Learning Representation*. San Diego, CA. 7 - 9 Mei 2015.
- [5] Bhattacharyya, P. (2015). *Machine Translation*. Mumbai, India.
- [6] Kelleher, J. D. (2016), Fundamental of Machine Learning for Neural Machine Translation, *Translationg Europe Forum 2016: Focusing on Translation Technologies*. European Commision Directorate-General for Translation.
- [7] Hermanto, A., Adji, T. B., & Setiawan, N. A. (2015) . Recurrent Neural Network Language Model for English-Indonesian Machine Translation: Experimental Study. *Proceeding of International Conference on Science in Information Technology*. Yogyakarta, Indonesia. 27 – 28 October 2015.
- [8] Papineni, K., Roukos, S., Ward, T., & Zhu, W.J. (2002). BLEU: a method for automatic evaluation of machine translation. *Proceeding of the 40th annual meeting on association or computational linguistics*. Philadelphia, Pennsylvania. 7 - 12 July 2002.
- [9] Zhang, J., Ding, Y., Shen, S., Cheng, Y., Sun, M., Luan, H., dan Liu, Y. (2017). *THUMT: An Open Source Toolkit for Neural Machine Translation* [diakses 1 Agustus 2017]. Tersedia dari <https://arxiv.org/abs/1706.06415>

Ukuran Risiko Cre-VaR

Insani Putri¹ dan Khreshna I. A. Syuhada²

Program Studi Magister Aktuaria - Institut Teknologi Bandung
Jalan Ganesha 10 Bandung 40132 Indonesia

¹insaniputriedi@gmail.com, ²khreshna@math.itb.ac.id

ABSTRAK

Asuransi adalah bentuk perlindungan dengan meminimalisir risiko. Agar kegiatan berasuransi berjalan lancar, perusahaan asuransi harus mampu mengelola risiko berupa klaim yang diajukan pemegang polis. Ukuran risiko adalah alat untuk memprediksi risiko dimasa yang akan datang dan membantu perusahaan mempersiapkan cadangan. Value-at-risk (VaR) adalah ukuran risiko yang sering digunakan. VaR adalah nilai risiko maksimum yang dapat terjadi pada suatu tingkat kepercayaan tertentu. Jika tersedia suatu sampel acak risiko maka dapat dihitung nilai VaR dari risiko tersebut. Mean sampel dapat diperoleh dari sampel acak risiko. Karena mean sampel adalah fungsi dari peubah acak maka dapat dihitung VaR mean sampel. VaR mean sampel adalah VaR dengan melibatkan informasi sebelumnya. VaR mean sampel lebih terpercaya daripada VaR karena mean sampel yang digunakan mewakili sekelompok sampel. Namun mean sampel yang digunakan dalam perhitungan VaR mean sampel, dianggap belum cukup baik dalam memprediksi risiko karena sifat mean sampel yang kurang stabil jika data sangat bervariasi. Dalam asuransi, teori kredibilitas memprediksi risiko (credibility premium) dengan mengkombinasikan mean sampel dengan informasi lain. Dapat dihitung VaR dari credibility premium yang disebut sebagai Credible VaR (Cre-VaR). Ukuran risiko yang optimal dapat dilihat dari keakuratannya. Keakuratan VaR, VaR mean sampel, dan Cre-VaR dapat diketahui dengan menghitung peluang cakupannya. Semakin dekat peluang cakupan dengan tingkat kepercayaan semakin akurat nilai VaR, VaR mean sampel, dan Cre-VaR yang dihitung.

Kata kunci: kredibilitas, mean sampel, prediksi risiko, volatilitas.

1. PENDAHULUAN

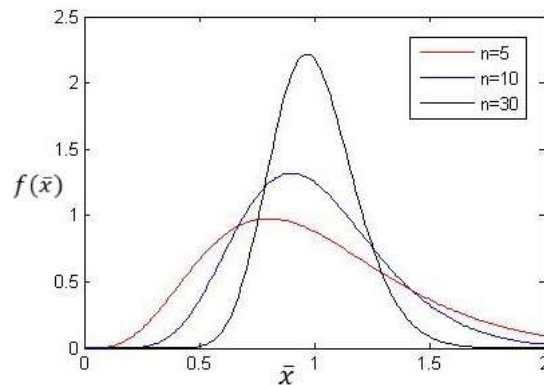
Setiap kegiatan dalam kehidupan selalu memiliki risiko. Risiko yang bersifat merugikan tentu diharapkan dapat diminimalisir keberadaannya. Salah satu cara yang digunakan masyarakat dalam meminimalisir risiko adalah dengan berasuransi. Berasuransi adalah kegiatan memindahkan risiko dari pemegang polis ke perusahaan asuransi. Risiko yang dapat dipindahkan ke perusahaan asuransi adalah risiko dapat dikuantifikasi, sehingga besar risiko yang mungkin akan ditanggung dapat diprediksi. Contohnya adalah besarnya biaya pengobatan saat sakit, besarnya biaya perbaikan untuk rumah yang mengalami kebakaran atau kendaraan yang mengalami kerusakan, dan besarnya biaya untuk mengembalikan kegiatan usaha yang mengalami kegagalan karena peristiwa yang tak terduga.

Agar kegiatan berasuransi lancar, perusahaan asuransi harus mampu mengelola risiko berupa klaim yang diajukan pemegang polis. Perusahaan asuransi harus memprediksi risiko dimasa yang akan datang dengan akurat. Ukuran risiko adalah alat untuk memprediksi risiko dimasa yang akan datang dan membantu perusahaan mempersiapkan cadangan. Ukuran Risiko yang sering digunakan adalah Value-at-Risk (VaR). VaR adalah nilai risiko maksimum yang dapat terjadi pada suatu tingkat kepercayaan tertentu.

Dalam menghitung ukuran risiko VaR, perusahaan harus memperhatikan apakah data yang tersedia kredibel atau terpercaya. Pada asuransi terdapat teori kredibilitas yang memperhatikan tingkat kredibilitas data. Dalam teori kredibilitas, model risiko yang baik adalah model yang mengkombinasikan data dengan informasi lain. Akan dihitung VaR menggunakan model risiko pada teori kredibilitas yang selanjutnya disebut Credible-VaR (Cre-VaR). Ukuran risiko yang optimal dapat dilihat dari keakuratannya. Keakuratan VaR dan Cre-VaR dapat diketahui dengan menghitung peluang cakupan. Semakin dekat peluang cakupan dengan tingkat kepercayaan, semakin akurat nilai VaR dan Cre-VaR yang dihitung.

2. TEORI KREDIBILITAS

Perusahaan asuransi harus memprediksi risiko (klaim) dimasa yang akan datang secara akurat. Misalkan sampel acak risiko $\mathbf{X}$ berukuran n dan diketahui data $X_1, \dots, X_n$. Diharapkan dapat diprediksi risiko X_{n+1} . Misalkan $E(X_i) = \mu$ dan $Var(X_i) = \sigma^2$ untuk $i = 1, 2, \dots, n$. Mean sampel $\bar{X} = \sum_{i=1}^n X_i / n$ dapat mewakili sampel acak $\mathbf{X}$. Mean sampel $\bar{X}$ merupakan fungsi dari peubah acak sehingga $\bar{X}$ adalah peubah acak dengan $E(\bar{X}) = \mu$ dan $Var(\bar{X}) = \sigma^2/n$. Mean sampel memiliki variansi yang lebih kecil. Semakin besar ukuran sampel semakin kecil variansi dari mean sampel. Gambar 1 adalah ilustrasi yang menunjukkan pengaruh ukuran sampel terhadap variansi mean sampel.



Gambar 1: Ilustrasi fungsi peluang mean sampel

Mean sampel $\bar{X}$ dapat digunakan untuk memprediksi X_{n+1} . Tetapi $\bar{X}$ belum cukup baik dalam memprediksi risiko karena $\bar{X}$ kurang stabil jika data $X_1, \dots, X_n$ sangat bervariasi. Dalam asuransi, terdapat teori kredibilitas yang memodelkan risiko sebagai kombinasi dari mean sampel dan informasi lain.

$$C = Z\bar{X} + (1 - Z)M, \quad (1)$$

dengan

C adalah prediksi risiko (credibility premium)
 $\bar{X}$ adalah mean sampel
 M adalah informasi lain (manual rate)
 Z adalah faktor kredibilitas
 $Z \in (0, 1)$

Pada persamaan (1), Z adalah proporsi dari $\bar{X}$ dan $(1 - Z)$ adalah proporsi dari M . Jika nilai Z mendekati 1 artinya memberikan proporsi yang lebih besar untuk $\bar{X}$ atau dengan kata lain 'lebih percaya' kepada mean sampel. Begitu sebaliknya, jika nilai Z mendekati 0 artinya proporsi yang lebih kecil untuk $\bar{X}$ atau dengan kata lain 'kurang percaya' pada mean sampel.

2.1. Manual Rate

Manual rate adalah informasi dari risiko yang hampir sama tapi tidak identik. Manual rate bisa berasal dari

- prediksi risiko awal yang ditetapkan perusahaan.
- prediksi risiko dari ketetapan pemerintah (manual rate yang dibuat dengan melibatkan data dalam jumlah besar untuk risiko yang hampir sama).

Berikut adalah ilustrasi manual rate berupa tarif premi yang ditetapkan oleh Otoritas Jasa Keuangan (OJK)^[1].

Tabel 1: Tarif Premi Asuransi Kendaraan Bermotor dengan Jenis Kendaraan Non Bus dan Non Truk untuk Pertanggungan Comprehensive pada wilayah DKI Jakarta, Jawa Barat, dan Banten.

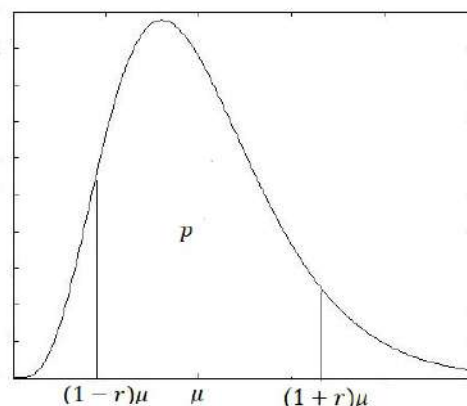
Kategori	Uang Pertanggungan	Rate Batas	
		Bawah	Atas
1	0 s.d. 125.000.000	3,26%	3,59%
2	>125.000.000 s.d. 200.000.000	2,47%	2,72%
3	>200.000.000 s.d. 400.000.000	2,08%	2,29%
4	>400.000.000 s.d. 800.000.000	1,20%	1,32%
5	>800.000.000	1,05%	1,16%

Misalkan untuk suatu polis asuransi kendaraan bermotor di kota Bandung memiliki uang pertanggungan sebesar 120.000.000 dan memiliki mean sampel 3.000.000 dari data klaim yang tersedia. Misalkan pula manual rate untuk polis ini adalah 3,26% (dipilih) $\times$ 120.000.000 = 3.912.000. Perusahaan harus menentukan apakah prediksi risiko untuk polis ini adalah 3.000.000 (full credibility, $Z = 1$), 3.912.000 (no credibility, $Z = 0$), atau kombinasi dari keduanya (partial credibility, $Z \in (0, 1)$). Selanjutnya akan dibahas bagaimana cara menentukan faktor kredibilitas Z .

2.2. Faktor Kredibilitas

Salah satu cara untuk menentukan faktor kredibilitas Z adalah dengan teori kredibilitas klasik^[2]. Pada dasarnya, teori kredibilitas klasik adalah menentukan standar untuk mencapai full credibility ($Z = 1$). Standar untuk full credibility dapat dinyatakan sebagai banyaknya data yang dibutuhkan agar menemukan kestabilan dari mean sampel $\bar{X}$. Mean sampel $\bar{X}$ dikatakan stabil jika $\bar{X}$ dekat dengan mean populasi μ dengan peluang yang tinggi. Secara matematis, $\bar{X}$ dikatakan stabil jika $\bar{X}$ berada antara $\pm r$ dari mean populasi μ dengan peluang lebih besar dari p (r mendekati 0 dan p mendekati 1).

$$P(\mu - r\mu \leq \bar{X} \leq \mu + r\mu) \geq p \quad (2)$$



Gambar 2: Ilustrasi kestabilan $\bar{X}$

Dari persamaan (2) dengan memanfaatkan pendekatan normal didapatkan

$$n \geq \left(\frac{\Phi^{-1}\left(\frac{p+1}{2}\right)}{r} \right)^2 \left(\frac{\sigma}{\mu} \right)^2, \quad (3)$$

yang disebut sebagai standar dari full credibility.

Jika standar dari full credibility tidak terpenuhi, digunakanlah partial credibility dan nilai $Z < 1$ harus ditentukan. Asumsi untuk menurunkan Z , peluang $Z\bar{X}$ berada dalam interval $[Z\mu - r\mu, Z\mu + r\mu]$ adalah p untuk nilai r yang diberikan.

$$P(Z\mu - r\mu \leq Z\bar{X} \leq Z\mu + r\mu) = p, \quad (4)$$

sehingga didapatkan

$$Z = \frac{r}{\Phi^{-1}\left(\frac{p+1}{2}\right)} \frac{\mu}{\frac{\sigma}{\sqrt{n}}} \quad (5)$$

3. VALUE-AT-RISK : PENDETEKSI RISIKO POTENSIAL

3.1. Ukuran Risiko Value-at-Risk

Value-at-risk (VaR) adalah ukuran risiko yang sering digunakan oleh berbagai institusi keuangan termasuk perusahaan asuransi. VaR digunakan untuk mendeteksi risiko yang potensial dan untuk mempersiapkan cadangan perusahaan.

VaR didefinisikan sebagai nilai risiko maksimum yang dapat terjadi pada suatu tingkat kepercayaan α dengan $\alpha \in (0, 1)$. Misalkan $X_1, \dots, X_n$ sampel acak risiko. Nilai VaR yang dihitung merupakan nilai prediksi untuk risiko X_{n+1} . Nilai VaR yang dihitung pada tingkat kepercayaan misalnya $\alpha = 0,95$, bermakna terdapat peluang 0,95 risiko yang akan didapat, kurang dari nilai VaR.

VaR pada tingkat kepercayaan α sama dengan kuantil ke- α dari distribusi X_{n+1} yang dapat dinyatakan dalam persamaan

$$P(X_{n+1} \leq x_\alpha; \theta) = \alpha \quad (6)$$

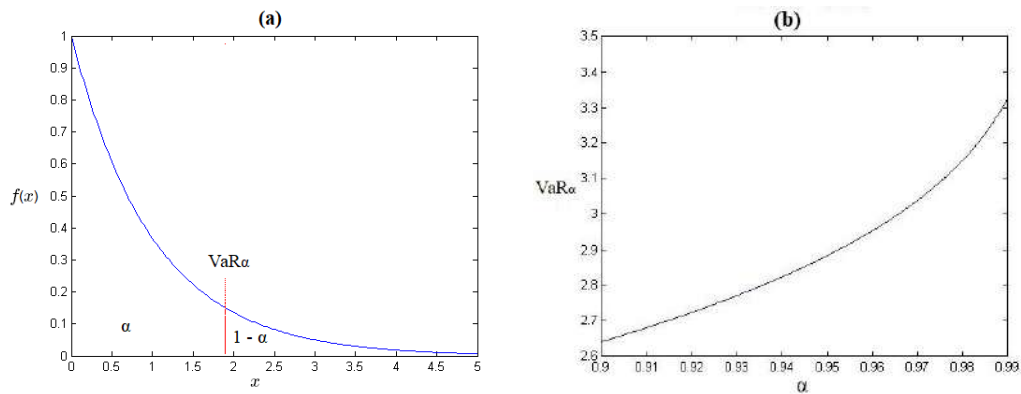
$$F_{X_{n+1}}(x_\alpha; \theta) = \alpha \quad (7)$$

sehingga

$$\text{VaR}_\alpha(X_{n+1}) = x_\alpha = F_{X_{n+1}}^{-1}(\alpha; \theta), \quad (8)$$

dengan VaR_α merupakan VaR pada tingkat kepercayaan α dan $F_{X_{n+1}}^{-1}$ merupakan invers dari fungsi distribusi $F_{X_{n+1}}$.

Gambar 3 adalah ilustrasi dari VaR pada distribusi eksponensial. Distribusi eksponensial yang peubah acaknyanya memiliki realisasi non negatif dapat digunakan sebagai distribusi besar klaim^[2].



Gambar 3: (a) VaR_α pada distribusi Eksponensial, (b) VaR dengan berbagai α pada distribusi Eksponensial

Prediksi VaR dapat dinyatakan sebagai

$$\widehat{VaR}_\alpha(X_{n+1}) = F_{X_{n+1}}^{-1}(\alpha; \hat{\theta}) \quad (9)$$

Dalam menaksir parameter θ akan terdapat variabilitas parameter sehingga mempengaruhi nilai prediksi VaR. Oleh karena itu, setelah memperoleh nilai prediksi VaR, perlu dievaluasi keakuratannya. Keakuratan prediksi VaR dapat dilihat dari nilai peluang cakupannya.

Peluang cakupan adalah peluang bahwa nilai risiko yang akan didapatkan kurang dari nilai prediksi VaR^[3]. Peluang cakupan dapat dituliskan sebagai

$$P(X_{n+1} \leq \widehat{VaR}_\alpha(X_{n+1}); \hat{\theta}) = \alpha + O(n^{-1}), \quad (10)$$

dengan $O(n^{-1})$ adalah efek dari variabilitas parameter terhadap peluang cakupan. Semakin dekat nilai peluang cakupan terhadap tingkat kepercayaan semakin akurat nilai prediksi VaR yang dihitung.

3.2. VaR Agregat Kerugian Acak

Dalam investasi, terdapat investasi portofolio dua aset. Maksudnya, investasi dengan menggabungkan dua aset dengan proporsi tertentu untuk mendapatkan imbal hasil (return) yang lebih besar. Investasi portofolio dua aset juga memiliki risiko yaitu risiko mendapatkan return negatif. Investor dapat mengetahui seberapa tinggi risiko dari portofolio dua aset dengan menghitung nilai VaR.

Misalkan X adalah return untuk aset pertama dan Y adalah return untuk aset kedua. Return R dari portofolio yang disusun oleh dua aset tersebut dapat dimodelkan sebagai

$$R = aX + (1 - a)Y, \quad (11)$$

dengan a adalah proporsi dari X , $(1 - a)$ adalah proporsi dari Y , dan $a \in (0, 1)$.

Untuk menghitung VaR dari portofolio dua aset harus ditentukan terlebih dahulu fungsi distribusi dari R . Fungsi distribusi R adalah

$$\begin{aligned} F_R(r) &= \int_{-\infty}^{\infty} P(R \leq r | Y = y) f_Y(y) dy \\ &= \int_{-\infty}^{\infty} P(aX + (1-a)Y \leq r | R = r) f_Y(y) dy \\ &= \int_{-\infty}^{\infty} P\left(X \leq \frac{r}{a} + \frac{(1-a)y}{a} \middle| Y = y\right) f_Y(y) dy \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\frac{r}{a} - \frac{(1-a)y}{a}} f_{X,Y}(x, y) dx dy, \end{aligned} \quad (12)$$

dengan $f_{X,Y}(x, y)$ adalah fungsi peluang bersama X dan Y .

VaR_α dari R adalah invers dari fungsi distribusi R

$$\text{VaR}_\alpha(R) = F_R^{-1}(\alpha; \theta) \quad (13)$$

Jika diperhatikan, model return portofolio dua aset mirip dengan model credibility premium pada teori kredibilitas. Sehingga, VaR dari portofolio dua aset inilah yang mendasari pengembangan Credible-VaR.

4. CREDIBLE VALUE-AT-RISK

Misalkan sampel acak risiko $X_1, \dots, X_n$. Jika terdapat m buah sample acak $\mathbf{X}$ berukuran n , maka akan terdapat sampel acak $\bar{X}_1, \bar{X}_2, \dots, \bar{X}_m$. Setiap $\bar{X}_j$ untuk $j = 1, \dots, m$ mewakili suatu sampel acak. Sehingga, diharapkan penggunaan $\bar{X}_1, \bar{X}_2, \dots, \bar{X}_m$ dalam perhitungan VaR akan menghasilkan peluang cakupan yang lebih baik dari pada $X_1, \dots, X_n$. VaR dari mean sampel pada tingkat kepercayaan α adalah

$$\text{VaR}_\alpha(\bar{X}_{m+1}) = F_{\bar{X}_{m+1}}^{-1}(\alpha; \theta), \quad (14)$$

dengan $F_{\bar{X}_{m+1}}^{-1}$ merupakan invers dari fungsi distribusi $F_{\bar{X}_{m+1}}$. Prediksi VaR dari mean sampel adalah

$$\widehat{\text{VaR}}_\alpha(\bar{X}_{m+1}) = F_{\bar{X}_{m+1}}^{-1}(\alpha; \hat{\theta}), \quad (15)$$

dengan peluang cakupan

$$P(\bar{X}_{m+1} \leq \widehat{\text{VaR}}_\alpha(\bar{X}_{m+1}); \hat{\theta}) = \alpha + O(m^{-1}) \quad (16)$$

Penggunaan mean sampel dalam perhitungan VaR mungkin belum cukup baik karena mean sampel yang kurang stabil jika data sangat bervariasi. Oleh karena itu dipakailah model risiko pada teori kredibilitas dalam perhitungan VaR. Credible Value-at-Risk (Cre-VaR) adalah VaR dari model risiko kredibilitas. Untuk menentukan Cre-VaR, harus ditentukan terlebih dahulu fungsi distribusi dari $C = Z\bar{X} + (1-Z)M$. Fungsi distribusi dari C adalah

$$\begin{aligned} F_C(c) &= P(C \leq c) \\ &= P(Z\bar{X} + (1-Z)M \leq c) \\ &= P\left(\bar{X} \leq \frac{c - (1-Z)M}{Z}\right) \\ &= F_{\bar{X}}\left(\frac{c - (1-Z)M}{Z}\right) \\ &= \int_{-\infty}^{\frac{c - (1-Z)M}{Z}} f_{\bar{X}}(\bar{x}) d\bar{x}, \end{aligned} \quad (17)$$

dengan $f_{\bar{X}}(\bar{x})$ adalah fungsi peluang dari mean sampel.

Cre-VaR pada tingkat kepercayaan α didefinisikan sebagai

$$\text{Cre-VaR}_{\alpha} = \text{VaR}_{\alpha}(C) = F_C^{-1}(\alpha; \theta), \quad (18)$$

dengan F_C^{-1} merupakan invers dari fungsi distribusi F_C .

Prediksi Cre-VaR pada tingkat kepercayaan α adalah

$$\widehat{\text{Cre-VaR}}_{\alpha} = F_C^{-1}(\alpha; \hat{\theta}) \quad (19)$$

5. PENERAPAN

Misalkan suatu polis pada asuransi kendaraan bermotor (mobil pribadi) di kota Bandung memiliki uang pertanggungan sebesar 120.000.000. Manual rate untuk polis ini adalah 3,26% (dipilih) $\times$ 120.000.000 = 3.912.000. Data diasumsikan berdistribusi Eksponensial dengan parameter $\theta = 3 \times 10^6$. Dengan teknik fungsi distribusi didapatkan mean sampel berdistribusi Gamma (500,500/ θ). Perhitungan VaR, VaR mean sampel, dan Cre-VaR dilakukan dengan simulasi Monte Carlo^[4].

Tabel 2: Hasil simulasi VaR, VaR mean sampel, dan Cre-VaR

α	0,90	0,95	0,99
$\text{VaR}_{\alpha}(X)$	$6,6021 \times 10^6$	$8,6422 \times 10^6$	$13,233 \times 10^6$
$\text{VaR}_{\alpha}(\bar{X})$	$3,0493 \times 10^6$	$3,0994 \times 10^6$	$3,1974 \times 10^6$
Cre-VaR_{α}	$3,3502 \times 10^6$	$3,3835 \times 10^6$	$3,4468 \times 10^6$

Tabel 3: Peluang Cakupan dari VaR, VaR mean sampel, dan Cre-VaR

α	0,90	0,95	0,99
$\text{VaR}_{\alpha}(X)$	0.8992	0.9504	0.9899
$\text{VaR}_{\alpha}(\bar{X})$	0.8996	0.9492	0.9899
Cre-VaR_{α}	0.9002	0.9501	0.9901

Dari hasil simulasi, nilai Cre-VaR lebih besar daripada VaR mean sampel tapi lebih kecil daripada nilai VaR. Akan tetapi, peluang cakupan Cre-VaR lebih akurat dibandingkan dengan VaR dan VaR mean sampel.

KEPUSTAKAAN

- [1] Surat Edaran Otoritas Jasa Keuangan Nomor 6/SEOJK.05/2017. Available from www.ojk.go.id/id/kanal/iknb/regulasi/asuransi/surat-edaran-ojk/Pages/Surat-Edaran-Otoritas-Jasa-Keuangan-Nomor-6-SEOJK.05-2017.aspx
- [2] Tse, Yiu-Kuen. (2009). *Nonlife Actuarial Models Theory, Methods, and Evaluation*. Cambrige University Press.
- [3] Cox, D.R. (1975). Prediction Intervals and Empirical Bayes Confidence Intervals. In: Gani, J. (Ed), *Prespective in Probability and Statistics*. Academic Press. London, pp. 47-55.
- [4] McNeil, A.J., Frey, R., dan Embrechts, P. (2005). *Quantitative Risk Management*. Princeton University Press.

PENENTUAN RISIKO INVESTASI $V@R$ DAN $CV@R$ DENGAN MOMEN ORDE TINGGI

Marianik¹ dan Khreshna Syuhada²

Program Studi Magister Matematika - Institut Teknologi Bandung
Jalan Ganesha 10 Bandung 40132 Indonesia

¹marianik@s.itb.ac.id, ²khreshna@math.itb.ac.id

Abstrak

Investasi merupakan kegiatan penanaman modal yang diharapkan dapat memberikan keuntungan bagi pihak yang terlibat. Investasi dalam bidang produksi dan jasa secara tidak langsung dapat meningkatkan perekonomian Indonesia. Investasi erat kaitannya dengan risiko. Risiko yang tidak dikelola dengan baik dapat memberikan dampak negatif (kerugian). Untuk menghitung risiko diperlukan suatu ukuran. Dua ukuran risiko kuantitatif yang digunakan adalah Value-at-Risk ($V@R$) dan Conditional- $V@R$ ($CV@R$). $V@R$ dikenal sebagai batas risiko maksimum yang dapat ditanggung oleh pelaku risiko, sedangkan $CV@R$ menyatakan mean dari nilai risiko yang lebih dari $V@R$. Penentuan nilai $V@R$ - $CV@R$ dapat dihitung dengan menggunakan basis kuantil. Ekspansi Cornish-Fisher (CF) dan Transformasi Johnson SU (JSU) yang melibatkan momen orde tinggi digunakan untuk menentukan kuantil distribusi risiko tersebut yang berujung pada penentuan nilai $V@R$ - $CV@R$. Ukuran risiko yang optimal ditentukan melalui nilai peluang cakupannya, semakin dekat nilai peluang cakupan dengan tingkat kepercayaan maka semakin akurat ukuran risiko tersebut. Nilai peluang cakupan dipengaruhi oleh metode penaksiran parameternya. Metode moment matching digunakan untuk menaksir parameter CF- $V@R$ - $CV@R$, sedangkan penaksiran parameter JSU- $V@R$ - $CV@R$ menggunakan metode LM-KK.

Kata kunci: ekspansi Cornish-Fisher, momen orde tinggi, peluang cakupan, transformasi Johnson SU, $V@R$ - $CV@R$.

1. PENDAHULUAN

Investasi adalah suatu tindakan penanaman modal (dapat berupa uang) pada suatu usaha atau kegiatan yang diharapkan dapat memberikan keuntungan bagi pelakunya (KBBI, 2017). Investasi menyumbang peranan yang cukup penting dalam kemajuan ekonomi Indonesia. Misalkan saja kegiatan penanaman modal pada pembangunan jalan dan pengadaan transportasi umum. Dalam berinvestasi ada tiga kemungkinan hasil yang diperoleh yaitu untung, rugi dan diantara keduanya. Untuk mengetahui hasil tersebut dibutuhkan suatu ukuran. Elton dan Gruber (1995) mendefinisikan imbal hasil atau return sebagai tingkat pengembalian dari suatu investasi pada suatu periode tertentu. Return positif menggambarkan bahwa nilai aset investasi naik (untung), nol menunjukkan tidak ada perubahan pada nilai aset investasi, sedangkan negatif menunjukkan bahwa nilai aset investasi turun (rugi) pada suatu periode. Golongan *risk taker* percaya bahwa semakin besar return (positif) yang diinginkan maka semakin besar pula risiko yang mungkin terjadi.

Dalam berinvestasi, kerugian dikenal sebagai risiko. Risiko memiliki sifat yang tidak pasti. Ketidakpastian ini menggiring kepada analisa lebih lanjut yang melibatkan analisis stokastik. Risiko dalam suatu periode tertentu dapat membentuk sebuah barisan. Barisan risiko ini yang nantinya akan dimodelkan sehingga dapat digunakan untuk memprediksi risiko yang akan datang. Selain itu, keberadaan risiko juga perlu diminimalisir, sehingga risiko perlu dikuantifikasi. McNeil dkk. (2005) dalam buku manajemen risiko kuantitatif-nya, mendefinisikan suatu ukuran risiko kuantitatif yang merepresentasikan risiko maksimum yang masih dapat ditolerir oleh pelaku risiko pada tingkat kepercayaan tertentu yang disebut value-at-risk ($V@R$). Sebagai alternatif, untuk mewakili kemungkinan prediksi nilai risiko yang jatuh pada nilai melebihi $V@R$ akan digunakan conditional- $V@R$ ($CV@R$). Nilai $CV@R$ dapat dipandang sebagai nilai rata-rata fungsi realitas risiko (luas rata-rata daerah dibawah fungsi) yang diwakili oleh perkalian nilai realisasi risiko dengan fungsi peluangnya. Dalam perkembangannya, nilai $V@R$ - $CV@R$ dapat ditentukan dengan basis prediksi berdasarkan kuantil (Christoffersen dan Gonçalves, 2004).

Kuantil ke- k distribusi didefinisikan sebagai nilai ke- $100 \times k$ dari statistik terurut peubah acaknya. Nilai kuantil ke- k juga dapat dipandang sebagai nilai invers fungsi distribusi saat k , sehingga penentuan fungsi distribusi yang cocok sangatlah penting. Teknik pendekatan yang digunakan untuk menentukan fungsi distribusi tersebut adalah ekspansi Cornish-Fisher (melalui ekspansi Edgeworth) dan transformasi Johnson SU yang melibatkan momen orde tinggi. Dua teknik ini memanfaatkan distribusi normal dalam penentuan fungsi distribusinya, sehingga memberikan kemudahan dalam mencari kuantil. Teknik tersebut pernah dilakukan Alexander, dkk. (2013) untuk menentukan $V@R$ berbasis kuantil untuk 5,10,20 periode kedepan pada model risiko GJR-GARCH(1,1). Dengan memanfaatkan hubungan antara $V@R$ dan $CV@R$ maka $CV@R$ berbasis kuantil dengan momen orde tinggi juga dapat ditentukan. Selain penentuan nilai $V@R$ - $CV@R$, pengukuran keakuratan dari kedua ukuran tersebut juga diperlukan. Salah satu ukuran yang dapat digunakan adalah nilai peluang cakupan.

Kabaila dan Syuhada (2010) memanfaatkan ekspansi Taylor untuk menentukan nilai peluang cakupan. Ekspansi Taylor dikenal dapat memberikan kemudahan untuk mendekati nilai fungsi pada suatu nilai dengan fungsi polinom. Tentunya pendekatan nilai tersebut melibatkan nilai 'sisa'. Dengan memanfaatkan konsep yang sama, ekspansi Taylor stokastik orde dua akan digunakan untuk mendekati fungsi distribusi dari suatu peubah acak dengan parameter fungsi distribusi tersebut sebagai variabel bebas. Orde pertama dan kedua ekspansi Taylor stokastik ini berturut-turut menyatakan Bias dan MSE parameter. Suku sisa pada ekspansi Taylor tersebut dianggap tidak memberikan pengaruh yang signifikan dengan orde ketelitian sebesar $O(n^{-1})$.

2. RISIKO STOKASTIK DAN UKURAN RISIKO

Risiko yang terjadi secara terus menerus dapat membentuk suatu barisan. Sifat ketidakpastian nilai risiko menggiring kepada analisa risiko secara stokastik. Oleh karena itu, risiko dalam periode waktu tertentu dapat membentuk suatu proses stokastik, yang berujung pada fakta bahwa risiko dapat diprediksi. Dalam penerapannya, risiko dapat dinyatakan sebagai fungsi dari volatilitas. Nilai volatilitas menyatakan besarnya perubahan risiko relatif terhadap 'informasi' sebelumnya pada suatu periode tertentu. Semakin besar nilai volatilitas maka semakin besar pula risiko yang mungkin ditanggung. Misalkan risiko setiap waktu dipetakan ke suatu peubah acak, maka kita akan mendapatkan suatu barisan peubah acak (proses stokastik). Misalkan $\{R_t\}_{t \geq 1}$ menyatakan barisan risiko acak tersebut. Model risiko tersebut memiliki formula $R_t = f(\sigma_t, \varepsilon_t) = \sigma_t \varepsilon_t$ dengan ε_t merupakan proses stokastik saling bebas dan identik dengan mean 0 dan variansi 1.

Volatilitas memiliki sifat yang tidak terobservasi secara langsung, sehingga tidak ada kesepakatan yang pasti untuk menentukan nilainya. Dalam melakukan prediksi risiko diperlukan prediksi nilai volatilitas. Salah satu formula yang dapat digunakan untuk menentukan nilai prediksi volatilitas satu langkah adalah volatilitas heteroskedastik yang memiliki formula sebagai berikut.

$$\sigma_{t+1}^2 = \alpha_0 + \alpha_1 R_t^2,$$

dengan parameter α_0 α_1 dapat ditentukan menggunakan metode gaussian-LM. Formula risiko dan bentuk volatilitas heteroskedastik di atas akan digunakan untuk menentukan bentuk umum ukuran risiko. Misalkan terdapat sebanyak n observasi risiko. Bentuk umum prediksi satu langkah $(n+1)$ ukuran risiko momen orde tinggi $V@R$ - $CV@R$ berbasis kuantil dapat dinyatakan dalam formula berikut ini.

$$V@R_{n+1;\hat{\theta}}^\alpha = \sigma_{n+1} \cdot H_{\varepsilon; \hat{m}_\varepsilon}^{-1(i)}(1 - \alpha) \quad (1)$$

$$CV@R_{n+1;\hat{\theta}}^\alpha = \sigma_{n+1} \cdot g_\varepsilon^\alpha(\hat{m}_\varepsilon^{(i)}) \quad (2)$$

dengan $H_{\varepsilon; \hat{m}_\varepsilon^{(i)}}^{-1}(1 - \alpha)$ merupakan nilai kuantil ke- $(1 - \alpha)$ dari distribusi inovasi (ε_t) ; $g_\varepsilon^\alpha(\hat{m}_\varepsilon^{(i)})$ menyatakan rata-rata nilai inovasi yang melebihi $H_{\varepsilon; \hat{m}_\varepsilon^{(i)}}^{-1}(1 - \alpha)$ relatif terhadap fungsi distribusinya.

Dalam penerapannya, bentuk $H_{\varepsilon; \hat{m}_\varepsilon^{(i)}}^{-1}(1 - \alpha)$ dan $g_\varepsilon^\alpha(\hat{m}_\varepsilon^{(i)})$ dapat ditentukan melalui dua pendekatan yaitu melalui ekspansi Cornish-Fisher dan transformasi Johnson SU. Berikut ini merupakan beberapa keunggulan dari kedua pendekatan ini: Pertama, penerapan momen orde tinggi dapat memberikan pendekatan kecocokan distribusi populasi lebih baik daripada distribusi normal. Kedua, keistimewaan bentuk kuantil dari distribusi normal baku memberikan kemudahan untuk menentukan kuantil dari ekspansi Cornish-Fisher dan transformasi Johnson SU. Dari bentuk umum $V@R-CV@R$ (1) dan (2) didapatkan prediksi momen orde tinggi $V@R-CV@R$ satu langkah sebagai berikut.

- Prediksi satu langkah Cornish Fisher $V@R-CV@R$

$$\begin{aligned} V@R_{n+1, \alpha; \hat{\theta}}^{CF} &= \sigma_{n+1} \cdot \left[q_\alpha + \frac{\hat{\beta}_3}{6} ((q_\alpha)^2 - 1) + \frac{\hat{\beta}_4}{24} q_\alpha ((q_\alpha)^2 - 3) - \frac{\hat{\beta}_3^2}{36} q_\alpha (2(q_\alpha)^2 - 5) \right] \\ CV@R_{n+1, \alpha; \hat{\theta}}^{CF} &= \sigma_{n+1} \cdot \frac{\phi(q_\alpha^{CF})}{\alpha} \left[1 + \frac{\hat{\beta}_3}{6} (q_\alpha^{CF})^3 + \frac{\hat{\beta}_4}{24} [(q_\alpha^{CF})^4 - 2(q_\alpha^{CF})^2 - 1] \right] \\ &\quad - \sigma_{n+1} \cdot \frac{\phi(q_\alpha^{CF})}{\alpha} \left[\frac{\hat{\beta}_3^2}{72} [(q_\alpha^{CF})^6 - 9(q_\alpha^{CF})^4 + 9(q_\alpha^{CF})^2 - 3] \right] \end{aligned}$$

dengan q_α dan q_α^{CF} berturut-turut menyatakan kuantil ke- $(1 - \alpha)$ dari distribusi normal baku dan ekspansi Cornish Fisher yang bersesuaian dengan nilai q_α .

- Prediksi satu langkah Johnson SU $V@R-CV@R$

$$\begin{aligned} V@R_{n+1, \alpha; \hat{\theta}}^{JSU} &= \sigma_{n+1} \cdot \left[\hat{\xi} + \hat{\lambda} \sinh \left(\frac{q_\alpha - \hat{\xi}}{\hat{\delta}} \right) \right] \\ CV@R_{n+1, \alpha; \hat{\theta}}^{JSU} &= \sigma_{n+1} \left[q_\alpha^{JSU} \phi \left(\hat{\xi} + \hat{\delta} \sinh^{-1} \left(\frac{q_\alpha^{JSU} - \hat{\xi}}{\hat{\lambda}} \right) \right) + 1 - \Phi \left(\hat{\xi} + \hat{\delta} \sinh^{-1} \left(\frac{q_\alpha^{JSU} - \hat{\xi}}{\hat{\lambda}} \right) \right) \right]. \end{aligned}$$

dengan q_α menyatakan kuantil ke- $(1 - \alpha)$ distribusi normal baku;

q_α^{JSU} menyatakan kuantil ke- $(1 - \alpha)$ transformasi Johnson SU yang bersesuaian dengan q_α .

3. EKSPANSI TAYLOR STOKASTIK: PENENTUAN PELUANG CAKUPAN $V@R-CV@R$

Ekspansi Taylor biasanya digunakan untuk mendekati nilai fungsi pada suatu nilai dengan fungsi polinom. Fungsi polinom ini memberikan kemudahan dalam perhitungan nilai tersebut. Dengan mengadopsi ide tersebut, ekspansi Taylor akan digunakan untuk mencari nilai peluang cakupan dari momen orde tinggi $V@R-CV@R$. Peluang cakupan merupakan suatu nilai yang merepresentasikan besarnya peluang bahwa nilai prediksi risiko akan jatuh pada daerah yang lebih kecil dari nilai $V@R-CV@R$. Berdasarkan bentuk umum $V@R-CV@R$ pada persamaan (1), peluang cakupan dari $V@R$ untuk proses stokastik risiko $\{R_t\}$ dapat ditentukan melalui peluang cakupan dari barisan peubah acak inovasinya $(\{\varepsilon_t\})$. Berikut ini merupakan bentuk umum peluang cakupan dari momen orde tinggi $V@R$.

$$P_\theta \left(\varepsilon_{t+1} \leq q_{t+1; \hat{\theta}}^\alpha \right) = E_\theta \left[H_\theta \left(q_{t+1; \hat{\theta}}^\alpha \right) \right]. \quad (3)$$

Sedangkan, bentuk umum peluang cakupan $CV@R$ menggunakan teknik yang pernah dilakukan oleh Rohmawati dan Syuhada (2015) yang memandang $CV@R$ sebagai nilai kuantil ke- U dari distribusi

inovasinya dengan U merupakan peubah acak berdistribusi uniform (0,1) yang lebih besar dari $(1 - \alpha)$. Oleh karena itu, bentuk umum peluang cakupan CV@R sama dengan V@R.

Ekspansi Taylor digunakan untuk mendekati nilai harapan $H_\theta \left(q_{t+1;\hat{\theta}}^\alpha \right)$ dengan $H_\theta \left(q_{t+1;\theta}^\alpha \right)$. Berdasarkan Bartle (2011), turunan ke- n dari H_θ dijamin ada di $\hat{\theta}$ dikarenakan turunan ke- $(n-1)$ dari H_θ selalu ada sepanjang interval I . Hal ini dapat terjadi karena H_θ merupakan fungsi distribusi dari peubah acak kontinu yang nilai turunan pertamanya (fungsi peluang) terjamin selalu ada di sepanjang interval I .

Misalkan $\hat{\theta}$ menyatakan peubah acak dari estimasi parameter. Misalkan $G_{\hat{\theta}}(\hat{\theta}; \alpha) = H_\theta(q_{n+1;\hat{\theta}}^\alpha)$. Ekspansi Taylor digunakan untuk menentukan fungsi $G_{\hat{\theta}}(\hat{\theta}; \alpha)$ sebagai berikut

$$G_{\hat{\theta}}(\hat{\theta}; \alpha) = G_{\hat{\theta}}(\theta; \alpha) + \frac{1}{1!} \frac{\partial G_{\hat{\theta}}(\hat{\theta}; \alpha)}{\partial \hat{\theta}_i} \Big|_{\hat{\theta}=\theta} (\hat{\theta}_i - \theta_i) + \frac{1}{2!} \frac{\partial^2 G_{\hat{\theta}}(\hat{\theta}; \alpha)}{\partial \hat{\theta}_i \partial \hat{\theta}_j} \Big|_{\hat{\theta}=\theta} (\hat{\theta}_i - \theta_i)(\hat{\theta}_j - \theta_j) + \dots \quad (4)$$

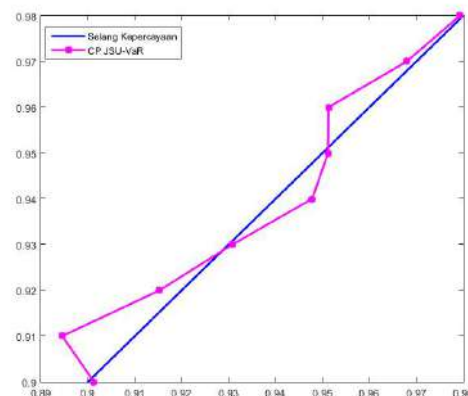
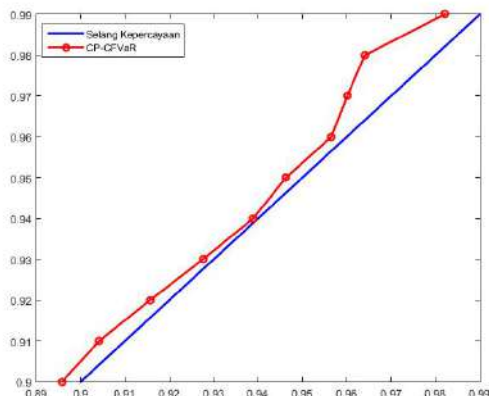
dengan $\hat{\theta}_i$ dan θ_i komponen ke- i dari vektor parameter $\hat{\theta}$ dan θ , sehingga dari persamaan (3) dan menerapkannya pada ekspansi Taylor pada persamaan (4) didapatkan

$$\begin{aligned} P_\theta \left(\varepsilon_{n+1} \leq q_{t+1;\hat{\theta}}^\alpha \right) &= E_\theta [G_{\hat{\theta}}(\hat{\theta}; \alpha)] \\ &= 1 - \alpha + c_\alpha(\theta)n^{-1}. \end{aligned} \quad (5)$$

Dari persamaan (5) di atas terlihat bahwa semakin dekat nilai peluang cakupan dengan selang kepercayaan $(1 - \alpha)$, maka semakin akurat ukuran risiko V@R-CV@R.

Bentuk peluang cakupan diatas memberikan implikasi bahwa penentuan peluang cakupan sangat bergantung terhadap metode penaksiran parameter yang digunakan. Pada metode CF-V@R digunakan metode penaksiran *moment matching* yaitu dengan mendekati parameter β_3 dengan momen ke-3 dari peubah acak inovasi dan parameter β_4 dengan momen ke-4 dari peubah acak inovasi. Sedangkan metode penaksiran yang digunakan untuk menaksir keempat parameter transformasi Johnson SU adalah gabungan dari metode likelihood maksimum dan kuadrat terkecil (LM-KK). Metode LM (likelihood maksimum) digunakan untuk menaksir parameter ζ dan δ , sedangkan metode KK (kuadrat terkecil) untuk menaksir parameter ξ dan λ . Dengan menggunakan kedua metode penaksiran tersebut didapatkan bentuk eksplisit taksiran parameter.

Penentuan seberapa jauh nilai peluang cakupan dengan selang kepercayaan dapat dilihat dari visualisasinya. Berikut ini merupakan grafik peluang cakupan V@R dengan ekspansi Cornish Fisher dan transformasi Johnson SU dengan data bangkitan sebanyak 500 data.



Gambar 1: Peluang Cakupan V@R (dengan CF) Gambar 2: Peluang Cakupan V@R (dengan JSU)

4. ANALISIS V@R-CV@R: DATA ASET SAHAM

Pada bab ini dilakukan penentuan prediksi nilai satu langkah V@R-CV@R dengan momen orde tinggi pada suatu aset saham. Data sampel yang digunakan adalah data aset saham PT Jasa Marga periode 1 Oktober 2013 sampai dengan 10 November 2017. Data yang digunakan adalah return majemuk. Return majemuk memiliki kekhasan yaitu jumlahan return selama periode tertentu akan sama dengan return pada akhir periode relatif terhadap return awal periode. Hal tersebut memberikan kemudahan dalam menganalisa return yang akan didapat pada akhir investasi.

Bentuk umum ukuran risiko V@R-CV@R pada persamaan (1) menunjukkan ukuran risiko di sebelah kanan, sehingga dalam pengolahannya return yang dipandang sebagai risiko akan ditransformasi terlebih dahulu. Transformasi yang dilakukan adalah dengan mencerminkan data return ke sumbu-x (negatif return), sehingga didapatkan data risiko sebelah kanan (dari sumbu-y). Berikut ini adalah hasil numerik nilai V@R-CV@R dengan ekspansi Cornish-Fisher dan transformasi Johnson SU dengan berbagai selang kepercayaan.

Table 1: Prediksi V@R dan CV@R Return Majemuk Saham PT Jasa Marga

$1 - \alpha$	CF-V@R	JSU-V@R	CF-CV@R	JSU-CV@R
0.90	0.0246	0.02634	0.03905	0.0406
PC 90%	89.62	90.25	94.52	94.86
0.95	0.0352	0.03708	0.0576	0.0534
PC 95%	95.12	94.91	95.81	96.34
0.99	0.0644	0.0655	0.1285	0.0808
PC 99%	98.25	99.51	98.99	99.99

Dari Tabel 1, didapatkan bahwa nilai V@R-CV@R kedua metode memberikan nilai yang hampir mirip dengan peluang cakupan V@R lebih dekat dibandingkan dengan nilai cakupan CV@R.

Berikut ini ilustrasi dari nilai prediksi V@R-CV@R pada suatu investasi. Misalkan investor memiliki nilai investasi pada tanggal 1 Oktober 2013 adalah satu juta rupiah. Risiko maksimum yang dapat ditanggung selama menginvestasikan sampai 10 November 2017 (V@R) dengan tingkat kepercayaan 95% dengan ekspansi Cornish-Fisher adalah sebesar 0.0352.

5. SIMPULAN

Pada artikel ini telah diformulasikan bentuk umum ukuran risiko V@R-CV@R berbasis kuantil yang di-representasikan dengan pendekatan momen orde tinggi yaitu CF V@R-CV@R dan JSU V@R-CV@R. Dengan fakta bahwa risiko merupakan fungsi dari volatilitas dan inovasinya, maka penentuan kedua ukuran risiko tersebut cukup mengacu pada pendekatan distribusi inovasinya. Penentuan ukuran risiko yang optimal dilihat dari kedekatan peluang cakupan dan selang kepercayaan yang digunakan. Dengan memodifikasi teknik yang dilakukan oleh Kabaila dan Syuhada (2010), peluang cakupan dari kedua ukuran risiko tersebut cukup dilihat dari ekspansi distribusi inovasinya. Dari Gambar 1 dan Gambar 2 terlihat bahwa V@R dengan ekspansi Cornish Fisher memberikan nilai peluang cakupan yang cukup dekat pada selang kepercayaan 90%-99%.

KEPUSTAKAAN

- [1] Elton, E.J. dan Gruber, M.J. (1995): *Modern Portfolio : Theory and Investment Analysis 5th ed.* John Wiley & Sons.
- [2] McNeil, A.J., Frey, R. dan Embrechts, P. (2005): *Quantitative Risk Management : Concepts, Techniques and Tools.* Princeton University Press.
- [3] Alexander, C., Lazar, E., dan Stanescu, S. (2013) : Forecasting VaR Using Analytic Higher Moments for GARCH Processes, *International Review of Financial Analysis*, 30, 36-45.
- [4] Kabaila, P. dan Syuhada, K. (2010). The Asymptotic Efficiency of Improved Prediction Intervals, *Statistics and Probability Letters*, 80, 1348-133
- [5] Rohmawati, A.A. dan Syuhada, K. (2015). Value-at-Risk and Expected Shortfall Relationship, *International Journal of Applied Mathematics and Statistics*, Vol. 53, No.5
- [6] Bartle, R.G. dan Sherbert, D.R. (2011): *Introduction to Real Analysis 4th ed.* John Wiley & Sons, Inc.

SIMULASI KOMPUTASI ALIRAN PANAS PADA MODEL PENGERING KABINET DENGAN METODE BEDA HINGGA

¹⁾Vivi Nur Utami, ²⁾Tiryono Ruby, ²⁾Subian Saidi, ²⁾Amanto

¹Jurusan Matematika, FMIPA, Universitas Lampung

Jl. Soemantri Brojonegoro no.1 Gedung meneng, Bandar Lampung

email: vivinurutami@gmail.com

ABSTRAK

Perpindahan panas (*heat transfer*) merupakan proses perpindahan energi yang terjadi karena adanya perbedaan suhu diantara benda atau material. Hal ini mengindikasikan bahwa perpindahan panas tidak hanya menjelaskan bagaimana perpindahan energi panas dari suatu benda ke benda lainnya, tetapi juga dapat meramalkan laju perpindahan panas yang terjadi pada kondisi-kondisi tertentu. Karena energi ini tidak dapat diukur ataupun diamati secara langsung tetapi dapat diamati arah perpindahan dan pengaruhnya. Penelitian ini bertujuan untuk mengembangkan ilmu matematika yaitu metode beda hingga untuk mendapatkan simulasi numerik perpindahan panas pada pengering kabinet menggunakan metode beda hingga agar terlihat simulasi perpindahan panas yang ada didalam oven. Dari penelitian ini di peroleh kesimpulan bahwa persebaran aliran panas ini memutar dari sumber yang berada paling bawah ke arah kanan dan kiri tidak langsung ke tengah oven dikarenakan ditengah merupakan ruang hampa tidak ada penghantar panas yang langsung dari sumber ke tengah oven.

Kata Kunci: Beda Hingga, solusi numerik, dan perpindahan aliran panas.

1. PENDAHULUAN

Perpindahan panas (*heat transfer*) merupakan proses perpindahan energi yang terjadi karena adanya perbedaan suhu diantara benda atau material. Hal ini mengindikasikan bahwa perpindahan panas tidak hanya menjelaskan bagaimana perpindahan energi panas dari suatu benda ke benda lainnya, tetapi juga dapat meramalkan laju perpindahan panas yang terjadi pada kondisi-kondisi tertentu. Karena energi ini tidak dapat diukur ataupun diamati secara langsung tetapi dapat diamati arah perpindahan dan pengaruhnya.

Jika kita melihat ke studi kasus, permasalahan yang berkaitan dengan persamaan model aliran panas yaitu perihai pengeringan pada suatu produk makanan ataupun bahan industri. Dikarenakan pengeringan sendiri erat kaitannya dengan perpindahan panas. Jika kita melihat ke studi kasus, permasalahan yang berkaitan dengan persamaan model aliran panas yaitu perihai pengeringan pada suatu produk makanan ataupun bahan industri. Dikarenakan pengeringan sendiri erat kaitannya dengan perpindahan panas. Untuk itu akan dikaji perihai perpindahan panas yaitu simulasi komputasi aliran panas pada model pengering dengan metode beda hingga sehingga didapatkan simulasi penyebaran panasnya menggunakan aplikasi matematika, matlab.

2. LANDASAN TEORI

2.1 Persamaan Diferensial Parsial

Persamaan diferensial parsial (PDP) adalah persamaan diferensial yang menyangkut turunan parsial dari satu atau lebih variabel tak bebas terhadap satu atau lebih variabel bebas. [1]

Persamaan diferensial parsial merupakan persamaan dengan dua variabel bebas / penentu atau lebih. Rumus- Rumus Forward difference dan *backward difference* serta *central difference*:

Beda Maju :

$$\frac{dy}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

Beda Mundur :

$$\frac{dy}{dx} = \lim_{h \rightarrow 0} \frac{f(x) - f(x-h)}{h}$$

Beda Tengah :

$$\frac{dy}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x-h)}{2h}$$

$$\frac{dy}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) + f(x-h) - 2f(x)}{h^2}$$

Berdasarkan definisi tersebut, maka dapat diketahui definisi dari turunan parsial sebagai berikut:

$$\text{Beda Maju} : \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x+h,y) - f(x,y)}{h} \quad \text{dan} \quad \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x,y+h) - f(x,y)}{h}$$

$$\text{Beda Mundur} : \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x,y) - f(x-h,y)}{h} \quad \text{dan} \quad \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x,y) - f(x,y-h)}{h}$$

$$\text{Beda Tengah} : \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x+h,y) - f(x-h,y)}{2h} \quad \text{dan} \quad \frac{\partial f}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x,y+h) - f(x,y-h)}{2h}$$

Dan definisi turunan Parsial Tingkat Dua, yaitu :

$$\frac{\partial^2 f}{\partial x^2} = \lim_{h \rightarrow 0} \frac{f(x+h,y) + f(x-h,y) - 2f(x,y)}{h^2} \quad \text{dan} \quad \frac{\partial^2 f}{\partial y^2} = \lim_{h \rightarrow 0} \frac{f(x,y+h) + f(x,y-h) - 2f(x,y)}{h^2} \quad [2]$$

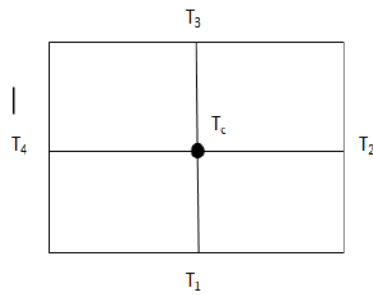
2.2 Metode Beda Hingga

Metode beda hingga atau yang lebih dikenal dengan *finite difference method*. Adalah metode numerik yang umum digunakan untuk menyelesaikan persoalan teknis dan problem matematis dari suatu gejala fisis. Secara umum metode beda hingga adalah metode yang mudah digunakan dalam penyelesaian problem fisis yang mempunyai bentuk geometri yang teratur, seperti interval dalam satu dimensi, domain kotak dalam dua dimensi, dan kubik dalam ruang tiga dimensi. [3]

Aplikasi penting dari metode beda hingga adalah dalam analisis numeric, khususnya pada persamaan diferensial biasa dan persamaan diferensial parsial. Prinsipnya adalah mengganti turunan yang ada pada persamaan diferensial dengan diskritisasi beda :

- Nilai fungsi di titik C

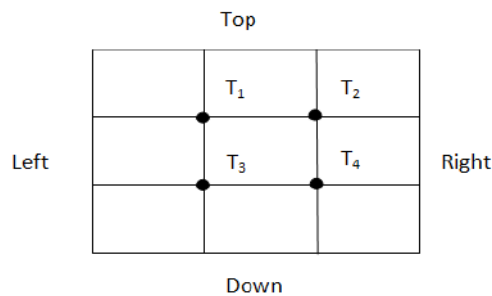
- Satu titik T_c (Temperatur di c)



Gambar 1. Ilustrasi untuk 1 Titik Tengah

$$T_1 + T_2 + T_3 + T_4 - 4T_c = 0$$

- Untuk empat titik T_c



Gambar 2. Ilustrasi untuk 4 Titik Tengah

$$l + T_2 + T_3 + T - 4T_1 = 0$$

$$T_1 + R + T + T_4 - 4T_2 = 0$$

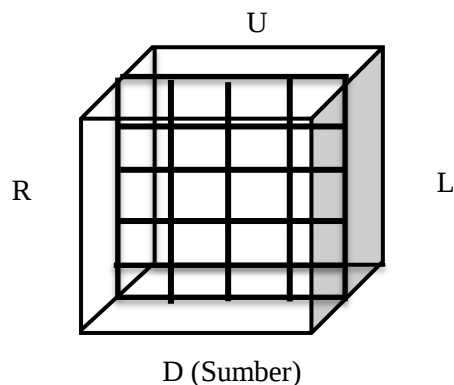
$$L + T_4 + T_1 + D - 4T_3 = 0$$

$$T_3 + R + T_2 + D - 4T_4 = 0$$

3. METODOLOGI PENELITIAN

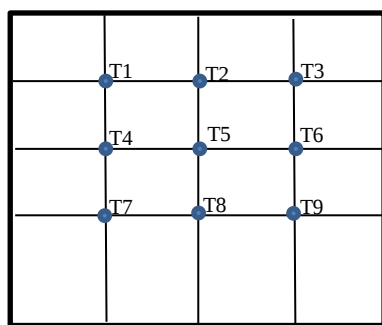
3.1 Alat dan Bahan

Objek yang diteliti yaitu oven rumahan dengan menggunakan alat perhitungan yaitu termometer oven yang digunakan untuk mengukur suhu dipinggir-pinggir oven yang mencapai 300°C , sebuah stopwatch untuk mengukur waktu sampai suhu stabil dan sebuah penggaris untuk mengukur panjang oven. Dengan menggunakan oven berukuran 30cm x 20cm x 30 cm.



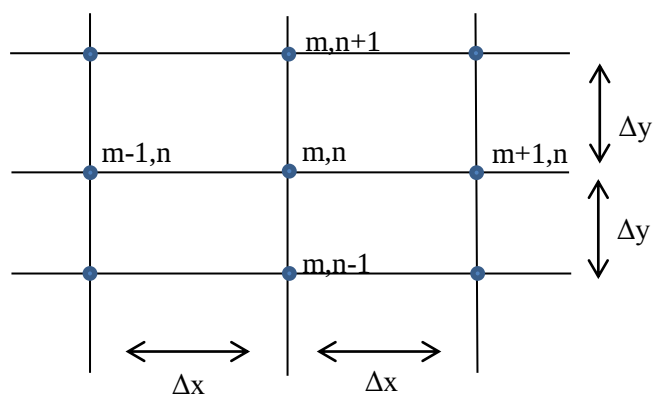
Gambar 3. Ilustrasi Oven

Untuk model aliran panas dari pengering kabinet ini sendiri dikarenakan kita ingin melihat aliran panas didalam oven dari sumber panas ke loyang di tingkat yang ketiga maka diambilah ilustrasi untuk penampang dalam pada oven bagian yang vertikal yaitu berukuran 30cm x 30cm.



Gambar 4. Ilustrasi Penampang Dalam Oven

Misalkan untuk plat berbentuk bujursangkar tersebut dibuat dibagi atas sejumlah kisi-kisi dengan pendekatan 9 node seperti gambar berikut :



Gambar 5. Desain dalam menerapkan metode beda hingga

Dengan menerapkan persamaan dari metode beda hingga, diperoleh :

$$T_{m+1,n} + T_{m-1,n} + T_{m,n+1} + T_{m,n-1} - 4_{m,n} = 0$$

4. HASIL DAN PEMBAHASAN

Dari hasil penelitian dari objek yang diteliti yaitu oven rumahan dengan menggunakan alat perhitungan yaitu termometer oven yang digunakan untuk mengukur suhu dipinggir-pinggir oven yang mencapai 300°C, sebuah stopwatch untuk mengukur waktu sampai suhu stabil dan sebuah penggaris untuk mengukur panjang oven. Dengan menggunakan oven berukuran 30cm x 20cm x 30 cm. Dengan ukuran loyang 30cm x 20cm. Maka, suhu-suhu yang didapat mengikuti waktu yang ditentukan yaitu :

Tabel 1. Waktu dan suhu yang didapat

Ket.	Suhu					
	t ₀ 0 Menit	t ₁ 5 Menit	t ₂ 10 Menit	t ₃ 15 Menit	t ₄ 20 Menit	t _{end}
Bawah (D)	0°	130°	185°	200°	210°	210°
Atas (U)	0°	70°	120°	150°	170°	210°
Kiri (L)	0°	95°	140°	165°	190°	210°
Kanan (R)	0°	90°	137°	160°	185°	210°

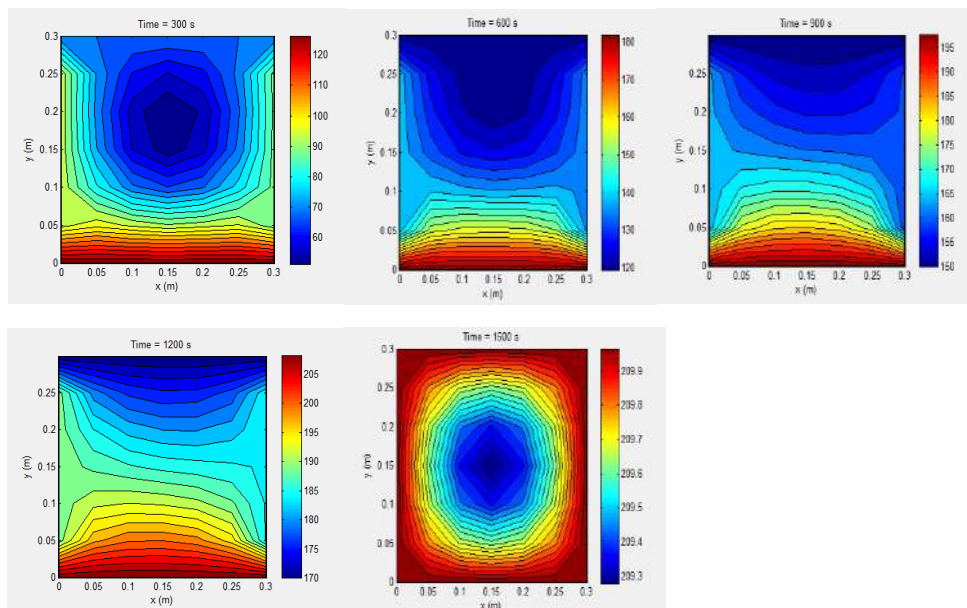
Diambil dari data yang sudah didapat saat penelitian bisa langsung kita terapkan pada ilustrasi penampang dalam dan ke persamaan metode beda hingga. Dengan menggunakan algoritma penyelesaian metode beda hingga secara iteratif dan kita terapkan batas awal untuk $T_1=T_2=T_3=T_4=T_5=T_6=T_7=T_8=T_9=0$. Hasil perhitungan dengan menerapkan algoritma metode beda hinggaseperti terlihat pada tabel berikut :

Tabel 2. Hasil perhitungan temperatur ilustrasi penampang dalam dengan metode beda hingga per 5 menit sampai waktu stabil

Ket.	Suhu (dalam °C)					
	t ₀ 0 Menit	t ₁ 5 Menit	t ₂ 10 Menit	t ₃ 15 Menit	t ₄ 20 Menit	t _{end}
T ₁	0	84.64	133.36	158.93	180.71	210
T ₂	0	84.33	133.32	159.60	180.49	210
T ₃	0	86.43	134.43	160.71	182.50	210
T ₄	0	94.24	143.11	166.12	187.37	210

T_5	0	96.25	145.50	168.75	188.75	210
T_6	0	96.38	144.39	168.26	189.51	210
T_7	0	106.07	156.57	176.79	195.00	210
T_8	0	110.04	161.18	181.03	197.63	210
T_9	0	107.86	157.64	178.57	196.79	210

4.1 Analisis Arah Aliran Panas Ilustrasi Penampang Dalam



Gambar 6. Aliran panas ilustrasi penampang dalam per 5 menit sampai suhu stabil

Dapat dilihat dari gambar yang didapat, terlihat jelas bahwa perpindahan panas terjadi dari bawah (sumber) memutar ke samping kanan dan kiri tidak langsung ke tengah lempengan dikarenakan di sebelah kanan dan kiri memiliki penghantar panas yaitu plat-plat sedangkan bagian tengah oven tidak ada penghantar panas yang berhubungan langsung dengan sumber panas dari kompor yang ada dibawah oven tersebut.

5. KESIMPULAN

Berdasarkan Hasil Penelitian dan Pembahasan dengan menggunakan Aplikasi Matlab, maka dapat ditarik kesimpulan bahwa dengan menggunakan teknik komputasi dengan metode beda hingga, didapatkan simulasi aliran perpindahan panas pada penampang tengah dari suatu oven rumahan yang menunjukkan persebaran aliran panas ini memutar dari sumber yang berada paling bawah ke arah kanan dan kiri tidak langsung ke tengah oven dikarenakan ditengah merupakan ruang hampa tidak ada penghantar panas yang langsung dari sumber ke tengah oven.

KEPUSTAKAAN

- [1] Ross, Shepley. (1984). *Differential Equation*. John Wiley and Sons, New York.
- [2] Chapra, S.C & Canale R.P. (1990). *Numerical Methods for Engineers* 2nd Ed, pp 707-749. McGraw-Hill Book Co., New York.
- [3] Li, Zhilin. (2010). *Finite Difference Methods Basics*. Southampton, Berlin

Segmentasi Provinsi Berdasarkan Derajat Kesehatan Dengan Menggunakan *Finite Mixture Partial Least Square* (FIMIX-PLS)

Agustina Riyanti

Statistisi di BPS Kabupaten Pesawaran

E-mail: justafity@gmail.com

ABSTRAK

Kesehatan merupakan salah satu investasi dalam pembangunan suatu bangsa. Indeks yang dapat digunakan untuk melihat kemajuan pembangunan kesehatan daerah adalah Indeks Pembangunan Kesehatan Masyarakat (IPKM) yang diukur dengan derajat kesehatan. Derajat kesehatan dalam penelitian ini diduga dipengaruhi oleh variabel laten lingkungan, variabel laten perilaku kesehatan, variabel laten pelayanan kesehatan, dan variabel laten genetik. Pendekatan yang digunakan dalam identifikasi wilayah dengan segmentasi wilayah berdasarkan derajat kesehatan yang terdiri dari beberapa variabel laten adalah dengan *Finite Mixture Partial Least Square* (FIMIX-PLS). Metode clustering biasa tidak dapat diterapkan dalam pengelompokan wilayah berdasarkan derajat kesehatan karena adanya hubungan antar variabel laten yang digunakan dalam struktur derajat kesehatan. Metode FIMIX-PLS (*Finite Mixture Partial Least Square*) dapat digunakan untuk pengelompokan wilayah berdasarkan derajat kesehatan dengan memperhitungkan hubungan antar variabel laten dan kelompok-kelompok yang ada dalam wilayah tersebut. Dengan metode PLS, didapatkan 10 indikator yang valid dan reliabel terhadap model derajat kesehatan. Pengelompokan wilayah berdasarkan derajat kesehatan menggunakan metode FIMIX-PLS berdasarkan nilai AIC, BIC, dan CIAC terendah serta nilai EN tertinggi menghasilkan 2 kelompok wilayah.

Kata kunci: *Finite Mixture, Heterogeneity, derajat kesehatan, Partial Least Squares, Structural Equation Modeling*

1. PENDAHULUAN

Pembangunan kesehatan harus dipandang sebagai suatu investasi untuk peningkatan kualitas sumber daya manusia. Menurut UU Kesehatan, bahwa setiap hal yang menyebabkan terjadinya gangguan kesehatan pada masyarakat Indonesia akan menimbulkan kerugian ekonomi yang besar bagi negara. Derajat kesehatan merupakan investasi pembangunan sumber daya manusia yang produktif secara sosial dan ekonomi. Dalam Undang-Undang Nomor 17 Tahun 2007 tentang RPJPN Tahun 2005-2025 dinyatakan bahwa dalam rangka mewujudkan Sumber Daya Manusia (SDM) yang berkualitas dan berdaya saing, maka kesehatan bersama-sama dengan pendidikan dan peningkatan daya beli masyarakat adalah tiga pilar utama untuk meningkatkan kualitas sumber daya manusia. Komposit dari tiga pilar utama ini selanjutnya dikenal dengan nama Indeks Pembangunan Manusia (IPM). IPKM adalah kumpulan indikator kesehatan yang dapat dengan mudah dan langsung diukur untuk menggambarkan masalah kesehatan. Serangkaian indikator kesehatan ini secara langsung maupun tidak langsung dapat berperan meningkatkan umur harapan hidup yang panjang dan sehat [1].

Penelitian terkait derajat kesehatan telah dilakukan oleh banyak pihak. Ummi, Sholiha, dan Salamah [2] menggunakan empat variabel laten eksogen (lingkungan, perilaku kesehatan, pelayanan kesehatan, dan genetik (keturunan)) dan satu variabel laten endogen (derajat kesehatan). Penelitian tersebut menyimpulkan bahwa seluruh indikator pada variabel lingkungan signifikan, tiga dari lima indikator pada variabel perilaku kesehatan signifikan, empat dari lima indikator pada variabel pelayanan kesehatan signifikan, dan dua dari tiga indikator pada variabel genetik signifikan. Anuraga, Sulistiyawan, dan Munadhiroh [3] meneliti tentang pengukuran derajat kesehatan di Jawa Timur. Ningsih, Jayanegara, dan Kencana [4] menganalisis kesehatan masyarakat Provinsi Bali dengan metode *Generalized Structured Component Analysis* (GSCA).

Permasalahan terkait derajat kesehatan masyarakat merupakan masalah yang kompleks dan multidimensional. Berbagai upaya dilakukan pemerintah untuk mengidentifikasi masalah derajat kesehatan masyarakat di Indonesia. Identifikasi wilayah berdasarkan derajat kesehatan sangat diperlukan oleh pemerintah ataupun pengambil kebijakan untuk penyusunan program dan kebijakan serta intervensi pada daerah prioritas.

Indonesia merupakan daerah yang terdiri dari beberapa provinsi yang memiliki suatu karakteristik tertentu. Penyusunan program pembangunan kesehatan akan lebih efektif dan efisien jika wilayah-wilayah tersebut dikelompokkan berdasarkan kemiripan karakteristik berdasarkan dimensi dalam derajat kesehatan masyarakat. Salah satu metode untuk pengelompokan wilayah adalah dengan clustering. Namun, metode clustering biasa tidak dapat diterapkan karena analisis berdasarkan derajat kesehatan masyarakat pada penelitian ini menggunakan beberapa variabel laten (dimensi) yang terdapat dalam derajat kesehatan masyarakat. Pada penelitian ini, metode yang digunakan untuk membuat pengelompokan wilayah berdasarkan derajat kesehatan masyarakat dengan memperhitungkan hubungan antar variabel laten dan adanya kelompok-kelompok wilayah di Indonesia adalah dengan metode *Finite Mixture Partial Least Squares* (FIMIX-PLS).

FIMIX-PLS menghasilkan pengelompokan wilayah berdasarkan hubungan antar variabel laten yang lebih homogen. *Finite Mixture Partial Least Squares* diperkenalkan pertama kali oleh Hahn, Carter, Johnson, dan Michael [5]. Ringle, Sarstedt, dan Mooi [6] mengembangkan metode sebelumnya dan mengaplikasikannya untuk data indeks kepuasan konsumen Amerika. Riyanti [7] menggunakan analisis FIMIX-PLS untuk melakukan segmentasi wilayah rawan pangan di Pulau Papua.

Tujuan penelitian ini adalah untuk mengetahui variabel dengan indikator yang valid dan reliabel apa sajakah yang berpengaruh signifikan terhadap derajat kesehatan masyarakat dan bagaimana pengelompokan wilayah berdasarkan variabel derajat kesehatan masyarakat dengan pendekatan metode FIMIX-PLS.

2. LANDASAN TEORI

A. *Structural Equation Modeling-Partial Least Square*

Partial Least Squares sering dikenal dengan sebutan *soft modeling* karena merupakan analisis yang *powerfull* dan meniadakan asumsi-asumsi OLS regresi. Asumsi OLS regresi yang ditiadakan dalam PLS antara lain asumsi data harus berdistribusi multivariate normal dan tidak adanya masalah multikolinearitas antar variabel eksogen [8]. Wold mengembangkan PLS untuk jumlah sampel yang kecil atau adanya masalah normalitas data. Selain digunakan untuk menjelaskan hubungan antar variabel laten, PLS dapat digunakan untuk mengkonfirmasi teori.

Model dalam PLS terdiri dari dua, yaitu *outer model* (model pengukuran) dan *inner model* (model struktural). Model pengukuran adalah model yang menggambarkan hubungan antara variabel laten dengan variabel pengukurannya (indikator), model struktural adalah model yang menghubungkan hubungan antar variabel latennya. Persamaan untuk model pengukuran dan model struktural adalah sebagai berikut:

A.1. Persamaan untuk *outer model* (model pengukuran)

1. Persamaan untuk model pengukuran variabel endogen:

$$\mathbf{Y} = \Lambda_y \boldsymbol{\eta} + \boldsymbol{\varepsilon} \quad (1)$$

atau

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_p \end{bmatrix} = \begin{bmatrix} \lambda_{y_{11}} & \lambda_{y_{12}} & \dots & \lambda_{y_{1m}} \\ \lambda_{y_{21}} & \lambda_{y_{22}} & \dots & \lambda_{y_{2m}} \\ \vdots & \vdots & \ddots & \vdots \\ \lambda_{y_{p1}} & \lambda_{y_{p2}} & \dots & \lambda_{y_{pm}} \end{bmatrix} \begin{bmatrix} \eta_1 \\ \eta_2 \\ \vdots \\ \eta_m \end{bmatrix} + \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_p \end{bmatrix}$$

dengan:

$\mathbf{Y}$: vektor dari variabel manifes endogen

Λ_y : matrik koefisien pengukuran (*loading factor*)

$\boldsymbol{\varepsilon}$: vektor dari kesalahan pengukuran

p : banyaknya indikator variabel endogen

m : banyaknya variabel endogen

2. Persamaan untuk model pengukuran variabel eksogen

$$\mathbf{X} = \Lambda_x \boldsymbol{\xi} + \boldsymbol{\delta} \quad (2)$$

atau

$$\begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_q \end{bmatrix} = \begin{bmatrix} \lambda_{x_{11}} & \lambda_{x_{12}} & \dots & \lambda_{x_{1n}} \\ \lambda_{x_{21}} & \lambda_{x_{22}} & \dots & \lambda_{x_{2n}} \\ \vdots & \vdots & \ddots & \vdots \\ \lambda_{x_{q1}} & \lambda_{x_{q2}} & \dots & \lambda_{x_{qn}} \end{bmatrix} \begin{bmatrix} \xi_1 \\ \xi_2 \\ \vdots \\ \xi_n \end{bmatrix} + \begin{bmatrix} \delta_1 \\ \delta_2 \\ \vdots \\ \delta_q \end{bmatrix}$$

dengan:

- X : vektor variabel manifest eksogen
 Λ_x : matrik koefisien pengukuran (*loading factor*)
 δ : vektor dari eror pengukuran
 q : banyaknya indikator variabel eksogen
 n : banyaknya variabel eksogen

A.2. Persamaan inner model (model struktural)

$$\eta = B \eta + \Gamma \xi + \zeta \quad (3)$$

atau

$$\begin{bmatrix} \eta_1 \\ \vdots \\ \eta_m \end{bmatrix} = \begin{bmatrix} 0 & 0 & \dots & 0 \\ \beta_{21} & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{m1} & \dots & \beta_{m(m-1)} & 0 \end{bmatrix} \begin{bmatrix} \eta_1 \\ \vdots \\ \eta_m \end{bmatrix} + \begin{bmatrix} \gamma_{11} & \gamma_{12} & \dots & \gamma_{1n} \\ \gamma_{21} & \gamma_{22} & \dots & \gamma_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \gamma_{m1} & \gamma_{m2} & \dots & \gamma_{mn} \end{bmatrix} \begin{bmatrix} \xi_1 \\ \vdots \\ \xi_n \end{bmatrix} + \begin{bmatrix} \zeta_1 \\ \vdots \\ \zeta_m \end{bmatrix}$$

dengan:

- η : vektor dari variabel endogen
 ξ : vektor dari variabel eksogen
 B dan Γ : matrik dari koefisien struktural
 ζ : vektor dari kesalahan (*error*) struktural
 m : banyaknya variabel laten endogen
 n : banyaknya variabel indikator

B. Finite-Mixture Partial Least Square (FIMIX-PLS)

Estimasi FIMIX-PLS berdasarkan asumsi bahwa heterogenitas terjadi pada model struktural dengan asumsi η_i berdistribusi *finite mixture* dengan fungsi kepadatan multivariate normal $f_{i|k}(\cdot)$

$$\eta_i \sim \sum_{k=1}^K \rho_k f_{i|k}(\eta_i | \xi_i, B_k, \Gamma_k, \Psi_k) \quad (4)$$

dengan:

- η_i : vektor variabel endogen pada inner model ($i = 1, 2, \dots, I$)
 ρ_k : Proporsi Mixing kelas laten k , dimana $\rho_k > 0$ dan $\sum_{k=1}^K \rho_k = 1$
 $f_{i|k}(\cdot)$: Peluang untuk kasus i given kelas k dan parameter $(\cdot)$
 B_k : matrik koefisien jalur pada *inner model* untuk kelas laten k yang menunjukkan hubungan antar variabel laten endogen berukuran $M \times M$
 Γ_k : matrik koefisien jalur pada *inner model* untuk kelas laten k yang menunjukkan hubungan antara variabel laten eksogen dengan endogen berukuran $M \times J$
 Ψ_k : matrik $M \times M$ untuk kelas laten k yang mengandung varians regresi
 I : jumlah total kasus/observasi
 i : kasus /observasi i ($i = 1, 2, 3, \dots, I$)
 K : jumlah keseluruhan kelas
 k : kelas atau segmen k dengan $k = 1, 2, 3, \dots, K$

Estimasi model pada FIMIX-PLS mengikuti prinsip *likelihood*. Fungsi *Likelihood* pada FIMIX-PLS dimaksimumkan dengan *Expectation-Maximization* (EM) algorithm. *EM-algorithm* merupakan kombinasi dari *Expectation* (E) step dan *Maximization* (M) step.

Jumlah segmen optimum pada FIMIX-PLS dapat ditentukan dengan kriteria seperti *Akaike's information criterion* (AIC), *Bayesian Information Criterion* (BIC), *Consistent AIC* (CAIC) dan *Normed Entropy Criterion* (EN), dengan EN sebagai berikut:

$$EN_k = 1 - \frac{[\sum_i \sum_k -P_{ik} \ln(P_{ik})]}{I \ln(K)} \quad (5)$$

EN : *Normal Entrophy*, ukuran relatif antara 0-1
 P_{ik} : Peluang observasi ke-i pada kelas ke-k
 k : kelas atau segmen dengan $k=1,2,\dots,K$
 i : observasi ke-i dengan $i=1,2,\dots,I$

C. Derajat Kesehatan Masyarakat

Indikator-indikator yang dapat diuraikan dalam derajat kesehatan diantaranya adalah mortalitas (angka kematian), status gizi, dan morbiditas (angka kesakitan). Batlibangkes pada tahun 2013 dalam Indeks Kesehatan Masyarakat mengungkapkan beberapa dimensi dalam derajat kesehatan masyarakat. Dimensi tersebut antara lain kesehatan masyarakat, kesehatan reproduksi, pelayanan kesehatan, perilaku kesehatan, penyakit tidak menular, penyakit menular, dan kesehatan lingkungan.

3. METODOLOGI PENELITIAN

A. Sumber Data dan Variabel Penelitian

Data yang digunakan dalam penelitian ini adalah data sekunder yaitu, data dari publikasi Riskesdas (Riset Kesehatan Dasar) 2013, Unit observasi adalah 33 provinsi di Indonesia. Variabel yang digunakan dalam penelitian ini terdiri dari satu variabel laten endogen (derajat kesehatan) dan empat variabel laten eksogen (kesehatan lingkungan, perilaku kesehatan, pelayanan kesehatan, dan genetik)

Tabel 1. Variabel Penelitian

Variabel	Indikator	
Kesehatan Lingkungan	X1	Proporsi rumah tangga yang memiliki akses terhadap fasilitas sanitasi
	X2	Proporsi rumah tangga berdasarkan akses ke sumber air minum
Perilaku Kesehatan	X3	Proporsi penduduk berperilaku benar dalam mencuci tangan
	X4	Proporsi penduduk berperilaku benar dalam buang air besar
Pelayanan Kesehatan	X5	Persentase persalinan yang ditolong oleh tenaga kesehatan di fasilitas kesehatan
	X6	Proporsi kecamatan yang mempunyai cakupan dokter per penduduk
	X7	Proporsi desa yang mempunyai kecukupan posyandu
Genetik	X8	Prevalensi Obesitas Sentral
	X9	Prevalensi Hipertensi
	X10	Prevalensi Diabetes Melitus
Derajat Kesehatan	Y1	Angka Kematian Bayi
	Y2	Prevalensi Balita sangat pendek/ pendek

B. Tahapan Penelitian

Analisis data dilakukan melalui tahapan sebagai berikut:

- Melakukan estimasi parameter SEM-PLS terhadap data yang meliputi:
 - Konseptualisasi model meliputi merancang *outer* dan *inner model*.
 - Mengkonstruksi diagram jalur.
 - Mengkonversi diagram jalur ke dalam persamaan
 - Mengestimasi parameter model yang meliputi *path coefficient*, *loading*, dan *weight*.
 - Mengevaluasi *outer* dan *inner model*. Jika *outer model* valid dan reliable maka dilanjutkan dengan evaluasi *inner model*, jika tidak memenuhi, kembali mengkonstruksi diagram jalur
 - Mendapatkan nilai skor faktor dari model yang signifikan
- Melakukan pendekatan FIMIX-PLS yaitu:
 - Nilai skor faktor pada *inner model* digunakan untuk prosedur FIMIX-PLS yaitu untuk menentukan jumlah kelompok.
 - Evaluasi hasil dan identifikasi jumlah segmen berdasarkan *fit indicator*

c) Evaluasi dan interpretasi hasil spesifik pengelompokan dengan FIMIX-PLS

4. HASIL DAN PEMBAHASAN

1. Evaluasi Model Pengukuran

Pengujian model pengukuran meliputi pengujian validitas (uji diskriminan dan konvergensi) dan realibilitas (*cronbach's alpha*). Hipotesis yang diuji adalah :

H₀: $\lambda_i = 0$ (loading faktor tidak signifikan mengukur variabel laten)

H₁: $\lambda_i \neq 0$ (loading faktor signifikan mengukur variabel laten)

dengan $i = 1, 2, 3, \dots, p$ merupakan jumlah indikator.

a. Validitas

Evaluasi validitas dapat dilihat dengan validitas konvergen dan validitas diskriminannya. Validitas ditunjukkan oleh *loading factor*, AVE (*Average Variance Extracted*), dan communality. Ukuran reflektif individual dikatakan valid jika memiliki korelasi loading dengan konstruk yang diukur memiliki nilai $> 0,5$. Validitas diskriminan dapat dilihat nilai loading factor tertinggi pada konstruk yang dituju dibandingkan dengan loading factor untuk konstruk lainnya

Tabel 2. Nilai Loading Factor Indikator Reflektif pada Variabel Laten Derajat Kesehatan Masyarakat

Indikator	Loading Factor	Keterangan
X1	0,96	Valid
X2	0,83	Valid
X3	0,57	Valid
X4	0,98	Valid
X5	0,96	Valid
X6	0,56	Tidak Valid
X7	0,82	Valid
X8	0,74	Valid
X9	0,40	Tidak Valid
X10	0,97	Valid
Y1	0,87	Valid
Y2	0,90	Valid

Pada Tabel 2. terlihat bahwa ada dua indikator yang digunakan memiliki nilai *loading factor* $< 0,5$, yaitu Proporsi kecamatan yang mempunyai cakupan dokter per penduduk dan Prevalensi Hipertensi. Pada tahap selanjutnya, kembali mengkonstruksi diagram jalur dengan mengeluarkan indikator yang tidak valid.

Hasil setelah indikator yang tidak valid dikeluarkan ditampilkan seperti pada Tabel 3. Pada Tabel 3. terlihat bahwa nilai *loading factor* $> 0,5$ untuk semua indikator. Hal ini mengindikasikan bahwa indikator yang digunakan telah valid konvergen. Selanjutnya, dilakukan evaluasi untuk validitas diskriminan, pada tabel juga terlihat bahwa nilai *loading factor* yang menuju variabel laten yang dituju lebih besar dari pada nilai variabel laten yang menuju ke variabel laten lainnya. Hal ini berarti bahwa indikator yang digunakan telah memenuhi kriteria validitas diskriminan.

Tabel 3. Nilai Loading Factor dan Cross Loading Indikator Reflektif pada Variabel Laten Derajat Kesehatan Masyarakat Setelah Indikator Tidak Valid Dikeluarkan

Indikator	Variabel Laten				
	Derajat_Kesehatan	Genetik	Kesehatan_Lingkungan	Pelayanan_Kesehatan	Perilaku_Kesehatan
x1	-0,81	0,51	0,96	0,46	0,84
x10	-0,40	0,96	0,43	0,45	0,54
x2	-0,40	0,17	0,83	0,07	0,41
x3	-0,15	0,24	0,35	0,07	0,57

Indikator	Variabel Laten				
	Derajat_Kesehatan	Genetik	Kesehatan_Lingkungan	Pelayanan_Kesehatan	Perilaku_Kesehatan
x4	-0,68	0,52	0,76	0,55	0,98
x5	-0,61	0,48	0,39	0,94	0,50
x7	-0,51	0,21	0,27	0,92	0,43
x8	-0,18	0,79	0,30	0,12	0,33
y1	0,88	-0,17	-0,64	-0,48	-0,56
y2	0,90	-0,46	-0,67	-0,60	-0,57

b. Realibilitas

Uji reliabilitas dapat menggunakan dua metode, yaitu *Cronbach's alpha* dan *composite reliability*, Hair, dkk (2009) menyatakan bahwa *rule of thumb* nilai dan *Composite Reliability* (CR) adalah $> 0,7$, meskipun nilai 0,6 masih dapat diterima, sedangkan nilai untuk *Cronbachs Alpha* $> 0,5$. Nilai CR dan *Cronbachs Alpha* ditampilkan pada Tabel 4. di bawah ini:

Tabel 4. Nilai *Composite Reliability* dan *Cronbachs Alpha*
Variabel Laten Derajat Kesehatan Masyarakat

Variabel Laten	Composite Reliability	Cronbachs Alpha	Keterangan
Derajat_Kesehatan	0,88	0,73	Reliabel
Genetik	0,87	0,74	Reliabel
Kesehatan_Lingkungan	0,89	0,79	Reliabel
Pelayanan_Kesehatan	0,93	0,85	Reliabel
Perilaku_Kesehatan	0,77	0,58	Reliabel

Pada Tabel 4. terlihat bahwa semua indikator yang digunakan reliabel, maka 10 indikator yang dipakai merupakan indikator yang valid dan reliabel untuk model derajat kesehatan masyarakat.

2. Evaluasi dan Pengujian Inner Model (Model struktural)

Evaluasi *inner model* dilakukan dengan *bootstrapping*, untuk model derajat kesehatan masyarakat dengan evaluasi terhadap koefisien determinasi R^2 , nilai statistik t, dan koefisien parameter.

Tabel 5 Nilai R^2 Variabel Laten Model
Derajat Kesehatan Masyarakat

Variabel Laten	R^2
Derajat_Kesehatan	0,68
Genetik	
Kesehatan_Lingkungan	
Pelayanan_Kesehatan	
Perilaku_Kesehatan	

Nilai R^2 untuk derajat kesehatan masyarakat adalah 0,68 yang berarti bahwa variasi variabel derajat kesehatan masyarakat dapat dijelaskan sebesar 68 persen oleh variabel Genetik, kesehatan lingkungan, pelayanan kesehatan, dan perilaku kesehatan, sedangkan 32 persen lainnya dipengaruhi oleh faktor lain yang tidak terdapat dalam model penelitian ini.

Pengujian terhadap parameter dilakukan dengan estimasi *resampling bootstrap*. Pengaruh langsung diantara variabel laten dapat dilihat pada nilai *path coefficient*. Apabila nilai statistik-t $> 1,96$ ($\alpha=5\%$) maka variabel laten tersebut mempengaruhi variabel laten lainnya. Tabel 6 menunjukkan bahwa terdapat 2 pengaruh langsung antar variabel, yaitu variabel kesehatan lingkungan dan pelayanan kesehatan berpengaruh signifikan terhadap derajat kesehatan masyarakat. Sedangkan variabel perilaku kesehatan dan genetik tidak berpengaruh signifikan terhadap variabel derajat kesehatan masyarakat dalam penelitian ini.

Tabel 6. Nilai Path Coefficient dan T Statistics Variabel Laten Model Derajat Kesehatan Masyarakat

Jalur	Original Sample (β_{awal})	Sample Mean ($\beta_{Bootstrap}$)	Standard Error $\beta_{Bootstrap}$	T Statistics
Genetik -> Derajat_Kesehatan	0,06	0,08	0,14	0,41
Kesehatan_Lingkungan -> Derajat_Kesehatan	-0,61	-0,64	0,16	3,90
Pelayanan_Kesehatan -> Derajat_Kesehatan	-0,41	-0,40	0,14	3,02
Perilaku_Kesehatan -> Derajat_Kesehatan	0,00	0,00	0,22	0,00

Nilai koefisien parameter dijelaskan pada *total effect* yang menggambarkan besarnya pengaruh total yang diterima suatu variabel laten dari variabel laten lainnya. Nilai *total effect* ditampilkan pada Tabel7. di bawah ini:

Tabel 7. Nilai Total Effect Path Coefficient dan T-Statistics Variabel Laten Model Derajat Kesehatan Masyarakat

Jalur	Original Sample (β_{awal})	Sample Mean ($\beta_{Bootstrap}$)	Standard Error $\beta_{Bootstrap}$	T Statistics
Genetik -> Derajat_Kesehatan	0,06	0,08	0,14	0,41
Kesehatan_Lingkungan -> Derajat_Kesehatan	-0,61	-0,64	0,16	3,90
Pelayanan_Kesehatan -> Derajat_Kesehatan	-0,41	-0,40	0,14	3,02
Perilaku_Kesehatan -> Derajat_Kesehatan	0,00	0,00	0,22	0,00

Hasil pengujian berdasarkan nilai efek total menunjukkan bahwa terdapat 2 jalur yang signifikan yaitu pengaruh kesehatan lingkungan terhadap derajat kesehatan masyarakat dan pengaruh pelayanan kesehatan terhadap derajat kesehatan. Hasil yang diperoleh berdasarkan tabel diatas adalah:

- Variabel genetik tidak berpengaruh signifikan terhadap derajat kesehatan masyarakat
- Variabel kesehatan lingkungan berpengaruh negatif terhadap derajat kesehatan masyarakat, yang berarti perubahan kesehatan lingkungan akan berpengaruh terhadap derajat kesehatan masyarakat yang buruk, jika kondisi kesehatan lingkungan meningkat maka derajat kesehatan masyarakat yang buruk akan menurun
- Variabel pelayanan kesehatan berpengaruh negatif terhadap derajat kesehatan masyarakat, yang berarti perubahan pelayanan kesehatan akan berpengaruh terhadap derajat kesehatan masyarakat yang buruk, jika kondisi pelayanan kesehatan meningkat maka derajat kesehatan masyarakat yang buruk akan menurun.
- Variabel genetik tidak berpengaruh signifikan terhadap derajat kesehatan masyarakat

3. Penentuan Kelompok dengan FIMIX-PLS

FIMIX-PLS menghasilkan sejumlah kelompok berdasarkan kriteria statistik yang telah ditentukan, yaitu nilai AIC,BIC, CAIC, dan EN. Nilai AIC, BIC, CAIC yang diperoleh dengan melakukan iterasi sebanyak 10.000 disajikan pada Tabel 8. di bawah ini:

Tabel 8. Kriteria AIC, BIC, CAIC, dan EN untuk k=2,3,4, 5, dan 6

Kriteria	k=2	k=3	k=4	k=5	k=6
AIC	92,35	132,49	151,46	161,20	171,41
BIC	102,82	148,95	173,91	189,63	205,83
CAIC	103,03	149,28	174,36	190,20	206,52
EN	0,44	0,20	0,25	0,34	0,39

Pada Tabel 8 ditampilkan perbandingan untuk k=2,3, 4,5, dan 6. Pada saat k=2, Nilai AIC, BIC, CAIC memiliki nilai yang terkecil dengan nilai EN sebesar 0,44. Hal ini mengindikasikan bahwa jumlah kelompok yang terbentuk adalah 2. Pada jumlah kelompok 2, ukuran jumlah untuk kelompok 1 adalah sebesar 71 persen dari jumlah provinsi di Indonesia. Kelompok 2 memiliki ukuran jumlah provinsi sebanyak 29 persen dari jumlah provinsi di Indonesia.

Kelompok wilayah provinsi yang terbentuk dari FIMIX-PLS dapat dilihat pada tabel 9. berikut ini:

Tabel 9. Segmentasi Provinsi dengan FIMIX-PLS

Provinsi Kelompok 1					
1 Aceh	8 Kepulauan Riau	15 Nusa Tenggara Timur	22 Sulawesi Barat		
2 Sumatera Utara	9 DKI Jakarta	16 Kalimantan Barat	23 Maluku		
3 Sumatera Barat	10 Jawa Barat	17 Kalimantan Tengah	24 Maluku Utara		
4 Riau	11 Jawa Tengah	18 Kalimantan Selatan	25 Papua Barat		
5 Jambi	12 Jawa Timur	19 Sulawesi Tengah	26 Papua		
6 Lampung	13 Banten	20 Sulawesi Selatan			
7 Kep. Bangka Belitung	14 Bali	21 Gorontalo			
Provinsi Kelompok 2					
1 Sumatera Selatan					
2 Bengkulu					
3 DI Yogyakarta					
4 Nusa Tenggara Barat					
5 Kalimantan Timur					
6 Sulawesi Utara					
7 Sulawesi Tenggara					

Setiap kelompok yang terbentuk memiliki kecenderungan variabel laten yang berpengaruh. Masing-masing pengaruh variabel laten dapat dilihat pada Tabel 10. di bawah ini:

Tabel 10. Path Coefficient FIMIX-PLS pada Variabel Laten Derajat Kesehatan

Jalur	Kelompok 1	Kelompok 2	Global
Kesehatan_Lingkungan-> Derajat Kesehatan	-0,44	-1,00	0,29
Pelayanan_Kesehatan -> Derajat Kesehatan	-0,60	0,02	0,07

Pada Tabel 10. terlihat pada kelompok 1 yang terdiri dari 26 provinsi memiliki karakteristik pengaruh pelayanan kesehatan terhadap derajat kesehatan yang lebih besar daripada pengaruh kesehatan lingkungan. Hal ini mengindikasikan bahwa peningkatan pelayanan kesehatan akan memberi dampak yang besar terhadap penurunan derajat kesehatan yang buruk. Kelompok 2 terdiri dari 7 provinsi memiliki karakteristik pengaruh kesehatan lingkungan terhadap derajat kesehatan masyarakat yang lebih besar dibandingkan dengan kelompok 1. Hal ini mengindikasikan bahwa peningkatan kesehatan lingkungan akan lebih berpengaruh besar pada penurunan derajat kesehatan yang kurang baik pada provinsi-provinsi yang berada di kelompok kedua.

5. SIMPULAN

Dengan metode PLS, dalam penelitian ini didapatkan 10 indikator yang valid dan reliabel terhadap model derajat kesehatan, yaitu indikator Proporsi rumah tangga yang memiliki akses terhadap fasilitas sanitasi, Proporsi rumah tangga berdasarkan akses ke sumber air minum, Proporsi penduduk berperilaku benar dalam mencuci tangan, Proporsi penduduk berperilaku benar dalam buang air besar, Persentase persalinan yang ditolong oleh tenaga kesehatan di fasilitas kesehatan, Proporsi desa yang mempunyai kecukupan posyandu, Prevalensi Obesitas Sentral, Prevalensi Diabetes Melitus, Angka Kematian Bayi dan Prevalensi Balita sangat pendek/ pendek. Variabel laten yang berpengaruh terhadap derajat kesehatan masyarakat dalam penelitian ini adalah kesehatan lingkungan dan pelayanan kesehatan. Pengelompokan wilayah berdasarkan derajat kesehatan menggunakan metode FIMIX-PLS berdasarkan nilai AIC, BIC, dan CIAC terendah serta nilai EN tertinggi menghasilkan 2 kelompok wilayah.

KEPUSTAKAAN

- [1] Badan Penelitian dan Pengembangan Kesehatan Kementrian Kesehatan RI. (2014). Indeks Pembangunan Kesehatan Masyarakat. Batlibangkes, Jakarta
- [2] Umami, E., Sholihah, N., Salamah, M. (2015). *Structural Equation Modeling-Partial Least Square* untuk Pemodelan Derajat Kesehatan Kabupaten/Kota di Jawa Timur (Studi Kasus Data Indeks Pembangunan Kesehatan Masyarakat Jawa Timur (2013). Jurnal Sains dan Seni ITS, Vol.4, No. 2
- [3] Anuraga, G., Sulistiyawan, E., Munadhiroh, S. (2017). Structural Equation Modeling – Partial Least Square Untuk Pemodelan Indeks Pembangunan Kesehatan Masyarakat (Ipkm) Di Jawa Timur, *Seminar Nasional Matematika dan Aplikasinya*. Universitas Airlangga. 21 Oktober 2017
- [4] Ningsih, P.N.P., Jayanegara, K., Kencana, I.P.E.N. (2013). Analisis Derajat Kesehatan Masyarakat Provinsi Bali dengan Menggunakan Generalized Structured Component Analysis . e-jurnal matematika, Vol.2, No. 2, hal 54-58
- [5] Hahn, Carsten, Johnson, Michael, D. (2011), “Capturing Customer Heterogeneity Using a Finite Mixture PLS Approach”, *Schmalenbach Business Review*, Vol. 54, No. 3, hal 243-269.
- [6] Ringle, C.M, Sarstedt, M., Mooi, E.A. (2010), “Respon-Based Segmentation Using Finite Mixture Partial Least Square. Theoretical Foundations and an Application to American Customer Satisfaction Index Data”, dalam *Data Mining*, eds. Stahlbock, R., Crone, S.F., Lessmann, S., New York, hal. 19-49)
- [7] Riyanti, A. (2016), *Pengelompokan Wilayah Rawan Pangan di Pulau Papua dengan Pendekatan Finite Mixture Partial Least Square*, Tesis, Institut Teknologi Sepuluh Nopember, Surabaya.
- [8] Loureiro, S. dan Miranda F.J. (2011), Brand Equity and Brand Loyalty in the Internet Banking Context: FIMIX-PLS Market Segmentation , *Journal of Science and Management*, Vol. 4, No. 4, hal. 476-485.

Representasi Operator Linear dari Ruang barisan l_3 ke Ruang Barisan $l_{3/2}$

Risky Aulia Ulfa¹⁾, Muslim Ansori²⁾, Suharsono³⁾, Agus Sutrisno⁴⁾

Jurusan Matematika, FMIPA, Universitas Lampung
Jln. Soemantri Brodjonegoro No. 1 Bandar Lampung 35145
E-Mail : riskyaul9@gmail.com

¹⁾Mahasiswa Jurusan Matematika FMIPA Universitas Lampung
^{2,3,4)}Dosen Jurusan Matematika FMIPA Universitas Lampung

ABSTRAK

Suatu pemetaan pada ruang vector khususnya ruang bernorma disebut operator. Salah satu kajian tentang operator, dalam hal ini operator linear, merupakan suatu operator yang bekerja pada ruang barisan. Banyak kasus pada operator linear dari ruang barisan ke ruang barisan dapat diwakili oleh suatu matriks tak hingga. Matriks tak hingga yaitu suatu matriks berukuran tak hingga kali tak hingga.

Sebagai contoh, suatu matriks $A : l_3 \rightarrow l_{3/2}$, dengan $A = \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \ddots \end{bmatrix}$, $l_3 = \{x = (x_i) \mid (\sum_{i=1}^{\infty} |x_i|^3)^{\frac{1}{3}} < \infty\}$ dan $l_{3/2} = \left\{x = (x_i) \mid \left(\sum_{i=1}^{\infty} |x_i|^{\frac{3}{2}}\right)^{\frac{2}{3}} < \infty\right\}$ merupakan barisan bilangan real. Selanjutnya dikonstruksikan operator A dari ruang barisan l_3 ke ruang barisan $l_{3/2}$ dengan basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$. dan ditunjukkan bahwa koleksi semua operator membentuk ruang banach.

Kata Kunci : Operator, Ruang Barisan Terbatas

1. PENDAHULUAN

Salah satu kajian tentang operator, dalam hal ini operator linear, merupakan suatu operator yang bekerja pada ruang barisan. Banyak kasus pada operator linear dari ruang barisan ke ruang barisan dapat diwakili oleh suatu matriks tak hingga. Matriks tak hingga yaitu suatu matriks berukuran tak hingga kali tak hingga.

Untuk setiap bilangan real p dengan $1 \leq p < \infty$ didefinisikan

$$l^p = \left\{x \in \{x_j\} \in \omega : \sum_{j=1}^{\infty} |x_j|^p < \infty\right\}$$

Sebagai contoh, suatu matriks $A : l_3 \rightarrow l_{3/2}$ dengan $A = \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \ddots \end{bmatrix}$,

$l_3 = \{x = (x_i) \mid (\sum_{i=1}^{\infty} |x_i|^3)^{\frac{1}{3}} < \infty\}$ dan

$l_{3/2} = \left\{x = (x_i) \mid \left(\sum_{i=1}^{\infty} |x_i|^{\frac{3}{2}}\right)^{\frac{2}{3}} < \infty\right\}$ merupakan barisan bilangan real.

Jika $x = (x_i) \in l_3$ maka

$$\begin{aligned}
 A(x) = Ax &= \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \vdots \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix} \\
 &= \begin{bmatrix} a_{11}x_1 + a_{12}x_2 + \dots \\ a_{21}x_1 + a_{22}x_2 + \dots \\ \vdots \end{bmatrix} \\
 &= \begin{bmatrix} \sum_{j=1}^{\infty} a_{1j}x_j \\ \sum_{j=1}^{\infty} a_{2j}x_j \\ \vdots \end{bmatrix}
 \end{aligned}$$

Sehingga timbul suatu permasalahan, syarat apa yang harus dipenuhi supaya $A(x) \in l_{3/2}$.

2. LANDASAN TEORI

2.1 Operator

Definisi 2.1.1

Suatu pemetaan pada ruang vector khususnya ruang bernorma disebut operator [2].

Definisi 2.1.2

Diberikan ruang Bernorm X dan Y atas *field* yang sama.

- Pemetaan dari X dan Y disebut operator.
- Operator $A : X \rightarrow Y$ dikatakan linear jika untuk setiap $x, y \in X$ dan setiap skalar α berlaku $A(\alpha x) = \alpha Ax$ dan $A(x + y) = Ax + Ay$ [2].

2.2 Barisan

Definisi 2.2.1

Diberikan ω yaitu koleksi semua barisan bilangan *real*, jadi :

$$\omega = \{\bar{x} = \{x_k\} : x_k \in \mathbb{R}\}$$

- Untuk setiap bilangan *real* p dengan $1 \leq p < \infty$ didefinisikan

$$l^p = \{x \in \{x_j\} \in \omega : \sum_{j=1}^{\infty} |x_j|^p < \infty\} \quad (1)$$

dan norm pada l^p yaitu

$$\|x\|_p = \left(\sum_{j=1}^{\infty} |x_j|^p \right)^{\frac{1}{p}}$$

b. Untuk $p = \infty$ didefinisikan

$$l_{\infty} = \{\bar{x} = \{x_k\} \in \omega : \sup_{k \geq 1} |x_k| < \infty\} \quad (2)$$

dan norm pada l_{∞} yaitu

$$\|x\|_{\infty} = \sup_{k \geq 1} |x_k|$$

[1].

Definisi 2.2.2

Misal $p, q \in (1, \infty)$ dengan $\frac{1}{p} + \frac{1}{q} = 1$ (q konjugat p), untuk $x \in l^p$ dan $y \in l^p$

$$(x_j y_j)_{j \in N} \in l^{\infty} \text{ dan } \sum_{j=1}^{\infty} |x_j y_j| \leq \|x\|^p \|y\|^q \quad (3)$$

[1].

2.3 Ruang Vektor

Definisi 2.3.1

Ruang vector adalah suatu himpunan tak kosong X yang dilengkapi dengan fungsi penjumlahan $(+): X \times X \rightarrow X$ dan fungsi perkalian skalar $(\cdot): F \times X \rightarrow X$ sehingga untuk setiap skalar λ, μ dengan elemen $x, y, z \in X$ berlaku :

- i. $x + y = y + x$
- ii. $(x + y) + z = x + (y + z)$
- iii. ada $\theta \in X$ sehingga $x + \theta = x$
- iv. ada $-x \in x$ sehingga $x + (-x) = \theta$
- v. $1 \cdot x = x$
- vi. $\lambda(x + y) = \lambda x + \lambda y$
- vii. $(\lambda + \mu)x = \lambda x + \mu x$
- viii. $\lambda(\mu x) = (\lambda \mu)x$ (Maddox, 1970).

2.4 Basis

Definisi 2.4.1

Ruang vector V dikatakan terbangkitkan secara hingga (*finitely generated*) jika ada vektor-vektor $x_1, x_2, \dots, x_n \in V$ sehingga $V = [x_1, x_2, \dots, x_n]$. Dalam keadaan seperti itu $\{x_1, x_2, \dots, x_n\}$ disebut pembangkit (*generator*) ruang vector V .

Menurut definisi di atas, ruang vektor V terbangkitkan secara hingga jika dan hanya jika ada vektor-vektor $x_1, x_2, \dots, x_n \in V$ sehingga untuk setiap vektor $x \in V$ ada skalar-skalar $\alpha_1, \alpha_2, \dots, \alpha_n$ sehingga

$$x = \alpha_1 x_1 + \alpha_2 x_2 + \cdots + \alpha_n x_n \quad (4)$$

Secara umum, jika $B \subset V$ dan V terbangkitkan oleh B , jadi $|B| = V$ atau B pembangkit V , maka untuk setiap $x \in V$ terdapat vektor-vektor $x_1, x_2, \dots, x_n \in B$ dan skalar $\alpha_1, \alpha_2, \dots, \alpha_n$ sehingga

$$x = \sum_{k=1}^n \alpha_k x_k$$

[1].

Definisi 2.4.2

Diberikan ruang vektor V . Himpunan $B \subset V$ dikatakan bebas linear jika setiap himpunan bagian hingga di dalam B bebas linear [1].

Definisi 2.4.3

Diberikan ruang vektor V atas lapangan $\mathcal{F}$. Himpunan $B \subset V$ disebut basis (*base*) V jika B bebas linear dan $|V| = B$.

Contoh :

Himpunan $\{\tilde{e}_1, \tilde{e}_2, \dots, \tilde{e}_n\}$, dengan $\tilde{e}_k$ vektor di dalam R^n yang komponen ke- k sama dengan 1 dan semua komponen lainnya sama dengan 0, merupakan basis ruang vektor R^n [1].

2.5 Ruang Metrik

Ruang metric merupakan ruang abstrak, yaitu ruang yang dibangun oleh aksioma-aksioma tertentu. Ruang metric merupakan hal yang fundamental dalam analisis fungsional, sebab memegang peranan yang sama dengan jarak pada *real line* R .

Definisi 2.5.1

Misal X adalah himpunan tak kosong, suatu metriks di X adalah suatu fungsi $d: X \times X \rightarrow [0, \infty)$, sehingga untuk setiap pasangan $(x, y) \in X \times X$ berlaku :

- i. $d(x, y) \geq 0$ untuk setiap $x, y \in X$
- ii. $d(x, y) = 0$ jika dan hanya jika $x = y$
- iii. $d(x, y) = d(y, x)$ untuk setiap $x, y \in X$ (sifat simetri)
- iv. $d(x, y) \leq d(x, z) + d(z, y)$ untuk setiap $x, y, z \in X$ (ketidaksamaan segitiga)

Selanjutnya pasangan (X, d) dengan d adalah metric pada X disebut ruang metrik. Setiap anggota X disebut titik dan nilai $d(x, y)$ disebut jarak (*distance*) dari titik x ke titik y atau jarak antara titik x dan titik y [2].

2.6 Ruang Bernorma

Definisi 2.6.1

Diberikan ruang linear X . Fungsi $x \in X \rightarrow \|x\| \in \mathbb{R}$ yang mempunyai sifat-sifat :

- i. $\|x\| \geq 0$ untuk setiap $x \in X$
- ii. $\|x\| = 0$, jika dan hanya jika $x = 0$, (0 vektor nol)
- iii. $\|\alpha x\| = \|\alpha\| \cdot \|x\|$ untuk setiap skalar α dan $x \in X$
- iv. $\|x + y\| \leq \|x\| + \|y\|$ untuk setiap $x, y \in X$

Disebut norma (*norm*) pada X dan bilangan non negative $\|x\|$ disebut norma vektor x . Ruang linear X yang dilengkapi dengan suatu norma $\|\cdot\|$ disebut ruang bernorma (*norm space*) dan dituliskan singkat dengan $X, \|\cdot\|$ atau X saja asalkan normanya telah diketahui [1].

2.7 Ruang Banach

Definisi 2.7.1

Ruang Banach (*Banach space*) adalah ruang bernorma yang lengkap (sebagai ruang metrik yang lengkap) jika dalam suatu ruang bernorma X berlaku kondisi bahwa setiap barisan Cauchy di X adalah konvergen [1].

3. METODOLOGI PENELITIAN

3.1 Waktu dan Tempat Penelitian

Penelitian ini dilakukan pada semester ganjil tahun ajaran 2017/2018 di jurusan Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lampung.

3.2 Metode Penelitian

Operator A dikonstruksikan dari ruang barisan l_3 ke ruang barisan $l_{3/2}$ dengan

basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$. Selanjutnya, mengkonstruksikan norma operator A . Jika pendefinisian operator dapat dilakukan maka akan diselidiki apakah koleksi semua operator tersebut membentuk ruang Banach. Sebagai aplikasi, operator A direpresentasikan sebagai matriks takhingga yang dikerjakan pada barisan barisan l_3 ke ruang barisan $l_{3/2}$ dengan basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$.

4. HASIL DAN PEMBAHASAN

Berdasarkan pada pembahasan sebelumnya, yang dimaksud dengan ruang barisanterbatas l_3 dapat didefinisikan sebagai :

$$l_3 = \left\{ x = (x_i) \left| \left(\sum_{i=1}^{\infty} |x_i|^3 \right)^{\frac{1}{3}} < \infty \right. \right\}$$

dengan $(x_i) = (x_1, x_2, \dots)$ merupakan barisan bilangan real $\mathbb{R}$. Ruang barisan l_3 merupakan ruang Banach dengan ruang dual $(l_3)^* = \{x^*: l_3 \rightarrow \mathbb{R}\}$ yaitu koleksi semua fungsional linear dan kontinu pada l_3 .

Untuk sebarang $x^* \in (l_3)^*$ dan $x \in l_3$, penulisan $\langle x, x^* \rangle$ dimaksudkan sebagai fungsional x^* pada x atau $x^*(x)$. Barisan vektor $\{e_n\} \subset l_3$ dinamakan basis pada l_3 jika untuk setiap vektor $x \in l_3$ terdapat barisan skalar yang tunggal $\{a_n\}$ sehingga

$$x = \sum_{n=1}^{\infty} a_n e_n$$

Barisan $\{e_n^*\} \in (l_3)^*$ dengan $\|e_n^*\| = 1$ untuk setiap n dikatakan biortonormal terhadap basis $\{e_n\} \subset l_3$ jika

$$\langle e_m, e_n^* \rangle = \delta_{mn}$$

dengan $\delta_{mn} = 1$ untuk $m = n$ dan $\delta_{mn} = 0$ untuk $m \neq n$. Selanjutnya, pasangan $\{\{e_n\}, \{e_n^*\}\}$ disebut system biortonormal pada l_3 maka

$$x = \sum_{n=1}^{\infty} a_n e_n$$

dengan $\langle x, e_n^* \rangle = a_n$.

Jika $A \in \mathcal{L}_c(l_3, l_{3/2})$ maka operator $A^* \in \mathcal{L}_c((l_3)^*, (l_{3/2})^*)$ disebut operator pendamping (adjoint operator)

A jika dan hanya jika untuk setiap $x \in l_3$ dan $y^* \in (l_{3/2})^*$, berlaku

$$\langle A(x), y^* \rangle = \langle x, A^*(y^*) \rangle$$

Jadi, jika $\{e_n\} \subset l_3$ dan $\{d_m^*\} \in (l_{3/2})^*$ diperoleh

$$\langle A(e_n), d_m^* \rangle = \langle e_n, A^*(d_m^*) \rangle$$

Jika $\{e_n\}, \{f_n\} \subset l_3$ basis pada l_3 dan $\{d_m\} \subset l_{3/2}$ basis pada $l_{3/2}$, maka untuk setiap $A \in \mathcal{L}_c(l_3, l_{3/2})$ berlaku

$$A^*(d_m^*) = \sum_{n=1}^{\infty} \langle e_n, A^*(d_m^*) \rangle e_n^* = \sum_{n=1}^{\infty} \langle A(e_n), d_m^* \rangle e_n^* \quad (a)$$

Dan

$$A^*(d_m^*) = \sum_{n=1}^{\infty} \langle e_n, A^*(d_m^*) \rangle f_n^* = \sum_{n=1}^{\infty} \langle A(e_n), d_m^* \rangle f_n^* \quad (b)$$

Berdasarkan persamaan (a) dan (b) diperoleh

$$\sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(f_k), d_m^* \rangle f_k^* \right\| \quad (c)$$

Berdasarkan persamaan (a), (b) dan (c) didefinisikan pengertian operator dari ruang barisan l_3 ke ruang barisan $l_{3/2}$ sebagai berikut

Definisi 1.1 Operator $A \in \mathcal{L}_c(l_3, l_{3/2})$ merupakan operator-SM jika

$$\sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| < \infty$$

Dengan $\{e_n\} \subset l_3$ basis pada l_3 dan $\{d_m\} \subset l_{3/2}$ basis pada $l_{3/2}$.

Dapat dipahami bahwa bilangan $\|A\|$ dengan

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| < \infty$$

Tidak bergantung pada pemilihan basis $\{e_n\}$ pada l_3 . Selanjutnya, notasi $SM(l_3, l_{3/2})$ menyatakan koleksi semua operator-SM dari ruang barisan l_3 ke ruang barisan $l_{3/2}$.

Teorema 1.2 Untuk setiap $A \in SM(l_3, l_{3/2})$ berlaku

- i. $\|A\| \leq \|A\|_{SM}$
- ii. $SM(l_3, l_{3/2})$ merupakan ruang Banach terhadap norma $\|\cdot\|_{SM}$
- iii. Jika $A \in SM(l_3, l_{3/2})$ maka A operator kompak.

Bukti :

- i. Diambil sebarang basis $\{e_n\} \subset l_3$ dan $x \in l_3$ dan $\{d_m\} \subset l_{3/2}$ basis pada $l_{3/2}$, maka berdasarkan (a), (b) dan (c) diperoleh

$$\begin{aligned} \|A(x)\| &= \left\| \sum_{m=1}^{\infty} \langle A(x), d_m^* \rangle d_m \right\| \\ &= \left\| \sum_{m=1}^{\infty} \langle x, A^*(d_m^*) \rangle d_m \right\| \\ &\leq \|x\| \sum_{m=1}^{\infty} \|A^*(d_m^*)\| \\ &= \|x\| \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= \|x\| \|A\|_{SM} \end{aligned}$$

Yang berakibat $\|A\| \leq \|A\|_{SM}$.

- ii. Pertama ditunjukkan bahwa $SM(l_3, l_{3/2})$ merupakan ruang bernorma terhadap norma $\|\cdot\|_{SM}$ sebab :

a. Untuk setiap $A \in SM(l_3, l_{3/2})$

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \geq 0$$

dan

$$\|A\|_{SM} = 0 \Leftrightarrow \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| = 0$$

$$\Leftrightarrow \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* = A^*(d_m^*) = \theta \ (\forall m)$$

$$\Leftrightarrow A^* = 0 \text{ (operator nol)}$$

$$\Leftrightarrow A = 0 \text{ (operator nol)}$$

b. Untuk setiap $A \in SM(l_3, l_{3/2})$ dan skalar α , diperoleh

$$\begin{aligned} \|\alpha A\|_{SM} &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= |\alpha| \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= |\alpha| \|A\|_{SM} \end{aligned}$$

c. Jika diberikan $A_1, A_2 \in SM(l_3, l_{3/2})$ maka

$$\begin{aligned} \|A_1 + A_2\|_{SM} &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle (A_1 + A_2)(e_k), d_m^* \rangle e_k^* \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle (A_1)(e_k) + (A_2)(e_k), d_m^* \rangle e_k^* \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A_1(e_k), d_m^* \rangle e_k^* + \sum_{k=1}^{\infty} \langle A_2(e_k), d_m^* \rangle e_k^* \right\| \\ &\leq \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A_1(e_k), d_m^* \rangle e_k^* \right\| + \left\| \sum_{k=1}^{\infty} \langle A_2(e_k), d_m^* \rangle e_k^* \right\| \end{aligned}$$

Dengan kata lain,

$$\|A_1 + A_2\|_{SM} \leq \|A_1\|_{SM} + \|A_2\|_{SM}$$

Selanjutnya menunjukkan kelengkapan ruang $SM(l_3, l_{3/2})$ sebagai berikut :

Diambil sebarang barisan Cauchy $\{A_i\} \subset SM(l_3, l_{3/2})$. Untuk setiap bilangan $\varepsilon > 0$, terdapat bilangan bulat positif n_0 sehingga untuk setiap bilangan bulat positif $i, j \geq n_0$, berlaku

$$\|A_i - A_j\|_{SM} \leq \frac{\varepsilon}{3}$$

Akan dibuktikan bahwa terdapat $A \in SM(l_3, l_{3/2})$ sehingga

$$\lim_{i \rightarrow \infty} \|A_i - A\|_{SM} = 0$$

Karena

$$\|A_i - A_j\|_{\mathcal{L}_c(l_3, l_{3/2})} \leq \|A_i - A_j\|_{SM} < \frac{\varepsilon}{3}$$

Untuk setiap $A_i, A_j \in SM(l_3, l_{3/2})$ dengan $i, j \geq n_0$, maka barisan $\{A_i\}$ juga merupakan barisan Cauchy di dalam $\mathcal{L}_c(l_3, l_{3/2})$. Karena $\mathcal{L}_c(l_3, l_{3/2})$ ruang lengkap maka terdapat $A \in \mathcal{L}_c(l_3, l_{3/2})$ sehingga $\lim_{j \rightarrow \infty} A_j = A$. Oleh karena itu,

$$\begin{aligned} & \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} ((A_i - A)(e_k), d_m^*) e_k^* \right\| \\ &= \lim_{j \rightarrow \infty} \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} ((A_i - A)(e_k), d_m^*) e_k^* \right\| \\ &= \lim_{j \rightarrow \infty} \|A_i - A_j\|_{SM} < \frac{\varepsilon}{3} \end{aligned}$$

Untuk sebarang bilangan bulat $i \geq n_0$. Dengan kata lain, $A - A_i \in SM(l_3, l_{3/2})$, untuk $i \geq n_0$. Oleh karena itu, $A - A_{n_0} + A_{n_0} = A \in SM(l_3, l_{3/2})$ dan terbukti bahwa barisan $\{A_i\}$ konvergen ke suatu $A \in SM(l_3, l_{3/2})$.

Jadi, $SM(l_3, l_{3/2})$ merupakan ruang bernorma yang lengkap atau ruang Banach.

iii. Jika $A \in SM(l_3, l_{3/2})$ dan $x \in l_3$, maka

$$A(x) = \sum_{m=1}^{\infty} \langle A(x), d_m^* \rangle d_m$$

Oleh karena itu, untuk setiap bilangan bulat positif n , dapat didefinisikan operator $A_n: l_3 \rightarrow l_{3/2}$ dengan

$$A_n(x) = \sum_{m=1}^n \langle A(x), d_m^* \rangle d_m$$

Jelas bahwa $A_n \in \mathcal{L}_c(l_3, l_{3/2})$ dan A_n merupakan operator berhingga. Dengan kata lain, A_n operator kompak. Karena $\{A_n\}$ konvergen ke K maka K operator kompak.

Berdasarkan Teorema 1.2 diperoleh :

Akibat 1.3 $SM(l_3, l_{3/2}) \subset K(l_3, l_{3/2}) \subset \mathcal{L}_c(l_3, l_{3/2})$ dengan $K(l_3, l_{3/2})$ koleksi operator kompak dari l_3 ke $l_{3/2}$. Operator $A \in SM(l_3, l_{3/2})$ dapat diwakili oleh matriks tak hingga $A = A_{\infty \times \infty}$. Oleh karena itu, dalam bentuk matriks tak hingga karakteristik operator-SM tersebut dapat diuraikan sebagai berikut:

Teorema 1.4 Suatu operator linear kontinu $A: l_3 \rightarrow l_{3/2}$ merupakan operator-SM jika dan hanya jika terdapat suatu matriks (a_{ij}) yang memenuhi :

- i. $Ax = \{\sum_{j=1}^{\infty} a_{ij}x_j\} \in l_{3/2}$ untuk setiap $x = (x_i) \in l_3$
- ii. $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^3 < \infty$
- iii. $\sum_{i=1}^{\infty} |\sum_{j=1}^{\infty} a_{ij}| < \infty$

Bukti :

(Syarat perlu) karena $A: l_3 \rightarrow l_{3/2}$ linear dan kontinu maka dengan sendirinya berlaku i) dan ii). Operator A dalam bentuk matriks (a_{ij}) dikerjakan pada basis standard $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$ berbentuk

$$Ae_k = (a_{jk})_{j=1}^{\infty}$$

Karena $A = (a_{ij})$ operator-SM maka

$$\begin{aligned} \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} (Ae_k, d_m^*) d_m \right\| &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} d_m \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} \right\| < \infty \end{aligned}$$

(Syarat cukup) berdasarkan i) dan ii) maka $A = (a_{ij}): l_3 \rightarrow l_{3/2}$ linear dan kontinu. Selanjutnya, berdasarkan iii) diperoleh

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} (Ae_k, d_m^*) d_m \right\| = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} \right\| < \infty$$

Terbukti $A = (a_{ij})$ merupakan operator-SM

Contoh 1.5 Matriks $A = (a_{ij})$ dengan

$$a_{ij} = \begin{cases} \frac{1}{i^{3/2}} & i = j \\ 0 & i \neq j \end{cases}$$

Merepresentasikan operator-SM $A: l_3 \rightarrow l_{3/2}$ sebab :

i. Untuk setiap $x = (x_j) \in l_3$ berlaku

$$\|Ax\| = \left\| \sum_{j=1}^{\infty} a_{ij}x_j \right\| = \left(\sum_{j=1}^{\infty} \left| \frac{x_j}{j^{3/2}} \right|^{3/2} \right)^{\frac{2}{3}} < \left(\sum_{j=1}^{\infty} |x_j|^{3/2} \right)^{\frac{2}{3}} < \infty$$

Jadi, $Ax \in l_{3/2}$

ii. Bagian kedua terpenuhi sebab

$$\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^{3/2} = \sum_{i=1}^{\infty} \left| \frac{1}{i^{3/2}} \right|^{3/2} < \infty$$

iii. Bagian ketiga terpenuhi sebab

$$\sum_{i=1}^{\infty} \left| \sum_{j=1}^{\infty} a_{ij} \right| = \sum_{i=1}^{\infty} \frac{1}{i^{3/2}} < \infty$$

5. SIMPULAN

Operator linear dan kontinu $A : l_3 \rightarrow l_{3/2}$ merupakan operator-SM jika dan hanya jika terdapat suatu matriks

$A = (a_{ij})$ yang memenuhi :

1. $Ax = \{\sum_{j=1}^{\infty} a_{ij}x_j\} \in l_{3/2}$ untuk setiap $x = (x_i) \in l_3$
2. $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^{3/2} < \infty$
3. $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}| < \infty$

Koleksi semua operator SM $A : l_3 \rightarrow l_{3/2}$ yang dinotasikan dengan SM $(l_3, l_{3/2})$ membentuk ruang Banach.

6. UCAPAN TERIMA KASIH

Penulis menyampaikan ucapan terimakasih kepada Bapak Muslim Ansori, Bapak Suharsono dan Bapak Agus Sutrisno karena telah membimbing dan memberi arahan kepada penulis sehingga penulis dapat menyelesaikan penelitian ini serta kepada Penyelenggara Seminar Nasional Metode Kuantitatif 2017 yang telah memfasilitasi kegiatan Seminar Nasional ini.

7. KEPUSTAKAAN

- [1] Darmawijaya, S. (2007). *Pengantar Analisis Abstrak*. Universitas Gajah Mada, Yogyakarta.
- [2] Kreyszig, E. (1989). *Introductory Function Analysis with Application*. Wiley Classic Library, New York.
- [3] Maddox, I.J. (1970). *Element of Functional Analysis*. Cambridge University Press, London.

ANALISIS RANGKAIAN RESISTOR, INDUKTOR DAN KAPASITOR (RLC) DENGAN METODE RUNGE-KUTTA DAN ADAMS BASHFORTH MOULTON

Yudandi Kuputra Aji¹, Agus Sutrisno², Amanto³, Dorrah Aziz⁴

^{1,2,3,4} Universitas Lampung Jalan Prof. Dr. Soemantri Brodjonegoro No.1 Bandar Lampung 35145, Telp 0721-704625 - Fax. 0721-704625

Website : <http://fmipa.unila.ac.id/web/> , Email : yudandikuputra@gmail.com

ABSTRAK

Rangkaian RLC dapat berupa rangkaian dengan persamaan diferensial homogen. Model pada rangkaian RLC seri yaitu $L \frac{d^2 I}{dt^2} + R \frac{dI}{dt} + \frac{1}{C} I = 0$, dimana di analisis menggunakan dua metode numerik yaitu metode Runge-Kutta orde empat dan Adams Bashforth Moulton orde 3 sebagai prediktor, orde empat sebagai korektor. Penyelesaian secara analitik digunakan sebagai pembandingan dalam mencari solusi terbaik dari kedua metode numerik yang digunakan. Visualisasi grafik menggunakan aplikasi Matlab R2013b. dengan menggunakan metode Adams Bashforth Moulton waktu komputasi lebih cepat, waktu iterasi lebih cepat dan galat lebih kecil jika dibandingkan dengan menggunakan metode Runge-Kutta orde empat. Metode terbaik dalam penyelesaian model rangkaian resistor, induktor dan kapasitor (RLC) adalah metode Adams Bashforth Moulton.

Kata kunci : Rangkaian RLC, Metode numerik, Runge-Kutta, Adams Bashforth Moulton.

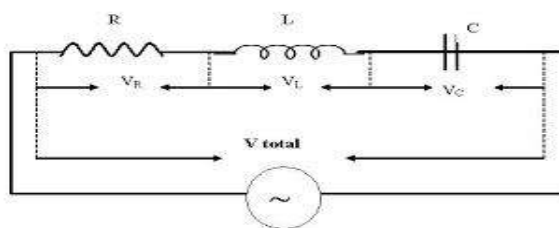
1. PENDAHULUAN

Rangkaian RLC merupakan rangkaian yang dihubungkan secara paralel ataupun seri. Rangkaian tersebut harus terdiri dari kapasitor, induktor dan resistor. Sesuai dengan namanya, susunan seri RLC merupakan susunan yang terdiri dari sebuah resistor (R), induktor (L), dan kapasitor (C) yang disusun secara seri dan dihubungkan dengan sumber tegangan. Karena terdiri dari tiga komponen, maka besar hambatan juga berasal dari ketiga komponen tersebut. Hambatan yang dihasilkan resistor disebut sebagai resistansi, hambatan yang dihasilkan oleh induktor biasa disebut reaktansi induktif yang disimbolkan dengan X_L , sedangkan hambatan yang dihasilkan oleh kapasitor disebut reaktansi kapasitif yang sering disimbolkan dengan X_C . Besar hambatan gabungan yang dihasilkan dalam rangkaian seri RLC disebut hambatan total atau impedansi.

2. LANDASAN TEORI

A. Rangkaian RLC

Pada rangkaian RLC seri arus listrik akan mendapat hambatan dari R, L dan C. Hambatan tersebut dinamakan Impedansi (Z). Impedansi merupakan gabungan secara vektor dari X_L , X_C dan R yang besarnya dilihat dari satuan Z.



Gambar 1 Rangkaian RLC seri dihubungkan dengan sumber tegangan arus bolak balik

Model pada rangkaian RLC seri dapat di tunjukkan pada persamaan berikut:

$$L \frac{d^2 I}{dt^2} + R \frac{dI}{dt} + \frac{1}{C} I = 0 \quad (1)$$

[1]

B. Metode Runge-Kutta

Metode Runge-Kutta orde 4 memiliki bentuk sebagai berikut:

$$x_{i+1} = x_i + \frac{1}{6}(k_1 + 2k_2 + 2k_3 + k_4) \quad (2)$$

dengan,

$$\begin{aligned} k_1 &= f(t_i, x_i) \\ k_2 &= f\left(t_i + \frac{1}{2}h, x_i + \frac{1}{2}hk_1\right) \\ k_3 &= f\left(t_i + \frac{1}{2}h, x_i + \frac{1}{2}hk_2\right) \\ k_4 &= f\left(t_i + h, x_i + hk_3\right) \end{aligned}$$

[2].

C. Metode Adams Bashforth Moulton

Metode Adams Bashfourth Moulton orde tiga sebagai prediktor memiliki persamaan sebagai berikut:

$$x_{k-1} = x_k + \frac{h}{12}(23f(t_k, x_k) - 16f(t_{k-1}, x_{k-1}) + 5f(t_{k-2}, x_{k-2})) \quad (3)$$

Untuk

$$k = 2, 3, 4, \dots$$

Pada metode ini galat hampiran adalah $O(h^3)$. Untuk menggunakan metode ini diperlukan tiga nilai awal x_0, x_1 , dan x_2 . Metode Adams Bashforth Moulton orde empat sebagai korektor memiliki persamaan sebagai berikut:

$$x_{k+1} = x_k + \frac{h}{24}(9f(t_{k+1}, x_{k+1}) + 19f(t_k, x_k) - 5f(t_{k-1}, x_{k-1}) + f(t_{k-2}, x_{k-2})) \quad (4)$$

Untuk

$$k = 2, 3, 4, \dots$$

Galat hampiran di dalam metode ini adalah $O(h^4)$, untuk hampiran ke- k . Metode ini juga merupakan metode implisit yang memerlukan tiga buah nilai awal x_0, x_1 , dan x_2 [3].

3. METODOLOGI PENELITIAN

Dalam melakukan penelitian ini langkah-langkah yang dilakukan adalah sebagai berikut:

1. Menyelesaikan persamaan model rangkaian RLC seri secara analitik, metode Runge-Kutta orde empat dan metode Adams Basfourth Moulton.
2. Membuat program dari rangkaian RLC menggunakan software Matlab R2013b.
3. Melihat metode terbaik untuk menyelesaikan persamaan rangkaian RLC dengan metode Runge-Kutta orde empat dan metode Adams Bashforth Moulton.

4. HASIL DAN PEMBAHASAN

A. Penyelesaian Secara Analitik

Persamaan RLC merupakan persamaan diferensial orde dua. Agar dapat diselesaikan dengan menggunakan metode Runge-Kutta orde empat dan Adam Basfourth Moulton, persamaan RLC harus diubah menjadi sistem persamaan diferensial orde satu [1].

$$\text{Misalkan } \frac{dI}{dt} = u, \frac{d^2I}{dt^2} = \frac{du}{dt}$$

$$\text{Maka, } \frac{du}{dt} + \frac{R}{L}u + \frac{1}{LC}I = 0$$

$$\frac{du}{dt} = -\frac{R}{L}u - \frac{1}{LC}I$$

Dengan demikian diperoleh sistem persamaan diferensial orde satu sebagai berikut :

$$\frac{dI}{dt} = u$$

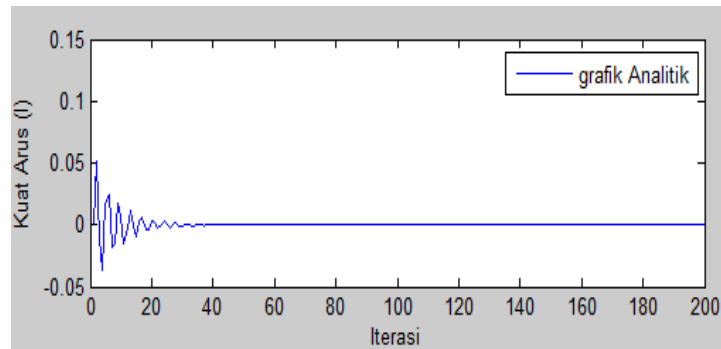
$$\frac{du}{dt} = -\frac{R}{L}u - \frac{1}{LC}I,$$

Dengan mencari akar- akar persamaannya, solusi umum dari persamaan RLC secara analitik adalah

$$I = c_1 e^{m_1 t} + c_2 e^{m_2 t}$$

dengan,

$$c_1 = c_2 = \frac{1}{m_1 - m_2}$$



Gambar 2 Grafik Solusi Analitik

B. Penyelesaian Menggunakan Metode Runge-Kutta^[2]

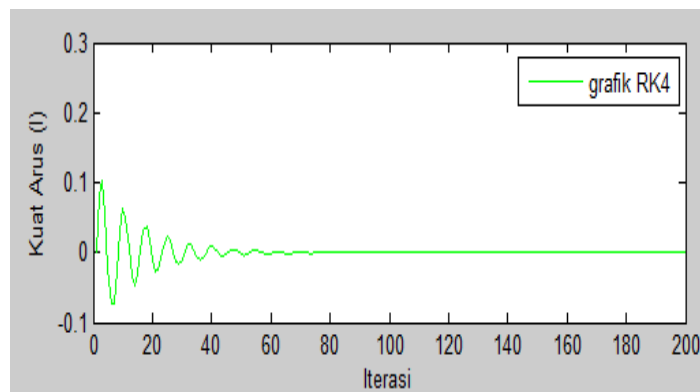
Berdasarkan rumus umumnya diperoleh ^[2]:

$$I_{i+1} = I_i + \frac{1}{6}h(k_1 + 2k_2 + 2k_3 + k_4)$$

$$u_{i+1} = u_i + \frac{1}{6}h(l_1 + 2l_2 + 2l_3 + l_4)$$

Dengan

- $k_1 = f(t, I, u) = u$
- $l_1 = g(t, I, u) = -\frac{R}{L}u - \frac{1}{LC}I$
- $k_2 = f\left(t + \frac{h}{2}, I + k_1\frac{h}{2}, u + l_1\frac{h}{2}\right) = u + l_1\frac{h}{2}$
- $l_2 = g\left(t + \frac{h}{2}, I + k_1\frac{h}{2}, u + l_1\frac{h}{2}\right) = -\frac{R}{L}\left(u + l_1\frac{h}{2}\right) - \frac{1}{LC}\left(x + k_1\frac{h}{2}\right)$
- $k_3 = f\left(t + \frac{h}{2}, I + k_2\frac{h}{2}, u + l_2\frac{h}{2}\right) = u + l_2\frac{h}{2}$
- $l_3 = g\left(t + \frac{h}{2}, I + k_2\frac{h}{2}, u + l_2\frac{h}{2}\right) = -\frac{R}{L}\left(u + l_2\frac{h}{2}\right) - \frac{1}{LC}\left(x + k_2\frac{h}{2}\right)$
- $k_4 = f(t + h, I + k_3h, u + l_3h) = u + l_3h$
- $l_4 = g(t + h, I + k_3h, u + l_3h) = -\frac{R}{L}(u + l_3h) - \frac{1}{LC}(x + k_3h)$



Gambar 3 Grafik Solusi Runge-Kutta Orde Empat

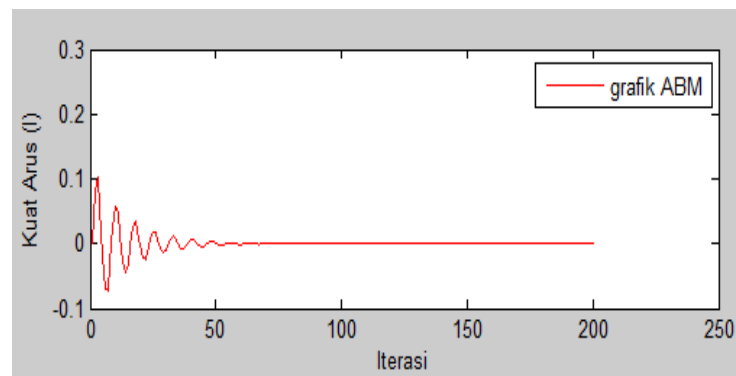
C. Penyelesaian Menggunakan Metode Adams Bashforth Moulton [3]

1. Melakukan tahap prediksi sehingga diperoleh:

- $I_{j+1} = I_j + \frac{1}{12}h(23f(t_j, I_j, u_j) - 16f(t_{j-1}, I_{j-1}, u_{j-1}) + 5f(t_{j-2}, I_{j-2}, u_{j-2}))$
- $u_{j+1} = u_j + \frac{1}{12}h(23g(t_j, I_j, u_j) - 16g(t_{j-1}, I_{j-1}, u_{j-1}) + 5g(t_{j-2}, I_{j-2}, u_{j-2}))$

2. Melakukan koreksi dengan skema Adams Moulton orde empat. Pada tahap koreksi,. Sehingga diperoleh:

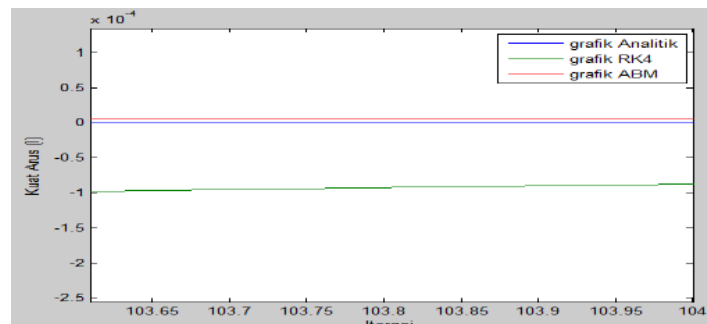
- $I_{j+1} = I_j + \frac{1}{24}h(9f(t_{j+1}, I_{j+1}, u_{j+1}) + 19f(t_j, I_j, u_j) - 5f(t_{j-1}, I_{j-1}, u_{j-1}) + f(t_{j-2}, I_{j-2}, u_{j-2}))$
- $u_{j+1} =$
 $u_j + \frac{1}{12}h(9g(t_{j+1}, I_{j+1}, u_{j+1}) + 19g(t_j, I_j, u_j) - 5g(t_{j-1}, I_{j-1}, u_{j-1}) + g(t_{j-2}, I_{j-2}, u_{j-2}))$



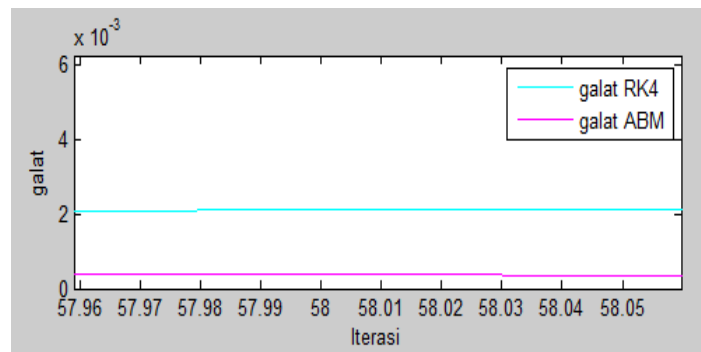
Gambar 4 Grafik Solusi Adams Bashforth Moulton

D. Perbandingan Metode

Model rangkaian RLC dengan metode Runge-Kutta dan Adams Bashforth Moulton diselesaikan dengan *Software* Matlab R2013b. Pada penelitian ini parameter yang digunakan adalah parameter simulasi yaitu dengan nilai resistansi (R) sebesar 8 *ohm*, nilai Induktansi (L) sebesar 6 *henry*, nilai Kapasitansi sebesar 23×10^{-4} *farad*, kuat arus awal ($I(0)$) sebesar 0 *ampere*, nilai perubahan arus (h) sebesar 10, dan nilai iterasi (n) sebesar 200, dengan tingkat ketelitian 10^{-5} . Berikut grafik hasil komputasi secara numerik maupun analitik :



Gambar 5 Grafik Perbandingan Solusi



Gambar 6 Grafik Perbandingan Galat

5. SIMPULAN

Dapat dilihat pada gambar 6 galat yang dihasilkan oleh metode Adams Bashforth Moulton lebih kecil dibandingkan metode Runge-Kutta orde empat. Pada gambar 5 jumlah iterasi dengan metode Adams Bashforth Moulton untuk mencapai nilai yang mendekati solusi analitik lebih cepat dibandingkan metode Runge-Kutta orde empat. Waktu komputasi yang dibutuhkan oleh metode Runge-Kutta orde empat adalah 0,044583. Sedangkan waktu komputasi yang dibutuhkan oleh metode Adams Bashforth Moulton adalah 0,034363. Maka dari itu metode Adams Bashforth Moulton merupakan metode terbaik untuk penyelesaian model rangkaian reduktor, induktor dan kapasitor (RLC).

KEPUSTAKAAN

- [1] Sutrisno. (1986). *Elektronika dan Aplikasinya*. ITB, Bandung.
- [2] Triatmodjo. (2002). *Metode Numerik Dilengkapi dengan Program Komputer*. Beta Offset, Yogyakarta.
- [3] Sahid. (2006). *Pengantar Komputasi Numerik dengan MATLAB*. Andi, Yogyakarta.

REPRESENTASI OPERATOR LINEAR DARI RUANG BARISAN l_{13} KE RUANG BARISAN $l_{13/12}$

Amanda Yona Ningtyas¹⁾, Muslim Ansori²⁾, Subian Saidi³⁾, Amanto⁴⁾

Jurusan Matematika, FMIPA, Universitas Lampung

Jalan Prof. Dr. Soemantri Brodjonegoro No. 1 Bandar Lampung 35145

E-Mail : yonamanda26@gmail.com

¹⁾Mahasiswa Jurusan Matematika FMIPA Universitas Lampung

^{2,3,4)}Dosen Jurusan Matematika FMIPA Universitas Lampung

ABSTRAK

Suatu pemetaan pada ruang vektor khususnya ruang bernorma disebut operator. Salah satu kajian tentang operator, dalam hal ini operator linear, merupakan suatu operator yang bekerja pada ruang barisan. Banyak kasus pada operator linear dari ruang barisan ke ruang barisan dapat diwakili oleh suatu matriks tak hingga. Matriks tak hingga yaitu suatu matriks berukuran tak hingga kali tak hingga. Sebagai contoh, suatu matriks $A : l_{13} \rightarrow l_{13/12}$, dengan $A = \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \ddots \end{bmatrix}$, $l_{13} = \{x = (x_i) \mid (\sum_{i=1}^{\infty} |x_i|^{13})^{\frac{1}{13}} < \infty\}$ dan $l_{13/12} = \{x = (x_i) \mid (\sum_{i=1}^{\infty} |x_i|^{\frac{13}{12}})^{\frac{12}{13}} < \infty\}$ merupakan barisan bilangan real. Selanjutnya dikonstruksikan operator A dari ruang barisan l_{13} ke ruang barisan $l_{13/12}$ dengan basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$. dan ditunjukkan bahwa koleksi semua operator membentuk ruang Banach.

Kata Kunci : Operator, Ruang Barisan Terbatas

1. PENDAHULUAN

Salah satu kajian tentang operator, dalam hal ini operator linear, merupakan suatu operator yang bekerja pada ruang barisan. Banyak kasus pada operator linear dari ruang barisan ke ruang barisan dapat diwakili oleh suatu matriks tak hingga. Matriks tak hingga yaitu suatu matriks berukuran tak hingga kali tak hingga.

Untuk setiap bilangan real p dengan $1 \leq p < \infty$ didefinisikan

$$l^p = \left\{ x \in \{x_j\} \in \omega : \sum_{j=1}^{\infty} |x_j|^p < \infty \right\}$$

Sebagai contoh, suatu matriks $A : l_{13} \rightarrow l_{13/12}$ dengan $A = \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \ddots \end{bmatrix}$,

$l_{13} = \{x = (x_i) \mid (\sum_{i=1}^{\infty} |x_i|^{13})^{\frac{1}{13}} < \infty\}$ dan

$l_{13/12} = \left\{ x = (x_i) \mid \left(\sum_{i=1}^{\infty} |x_i|^{\frac{13}{12}} \right)^{\frac{12}{13}} < \infty \right\}$ merupakan barisan bilangan real.

Jika $x = (x_i) \in l_{13}$ maka

$$\begin{aligned} A(x) = Ax &= \begin{bmatrix} a_{11} & a_{12} & \dots \\ a_{21} & a_{22} & \dots \\ \vdots & \vdots & \ddots \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix} \\ &= \begin{bmatrix} a_{11}x_1 + a_{12}x_2 + \dots \\ a_{21}x_1 + a_{22}x_2 + \dots \\ \vdots \end{bmatrix} \end{aligned}$$

$$= \begin{bmatrix} \sum_{j=1}^{\infty} a_{1j}x_j \\ \sum_{j=1}^{\infty} a_{2j}x_j \\ \vdots \end{bmatrix}$$

Sehingga timbul suatu permasalahan, syarat apa yang harus dipenuhi supaya $A(x) \in l_{13/12}$. Oleh karena itu, penelitian akan difokuskan pada permasalahan tersebut.

2. LANDASAN TEORI

OPERATOR

Definisi 2.1.1

Suatu pemetaan pada ruang vektor khususnya ruang bernorma disebut operator [1].

Definisi 2.1.2

Diberikan ruang Bernorm X dan Y atas *field* yang sama.

- Pemetaan dari X dan Y disebut operator.
- Operator $A : X \rightarrow Y$ dikatakan linear jika untuk setiap $x, y \in X$ dan setiap skalar α berlaku $A(\alpha x) = \alpha Ax$ dan $A(x + y) = Ax + Ay$.

BARISAN

Definisi 2.2.1

Barisan adalah suatu fungsi yang domainnya adalah himpunan bilangan bulat positif. Misal terdapat bilangan bulat positif 1, 2, 3,...,n,... yang bersesuaian dengan bilangan real x_n tertentu, maka $x_1, x_2, \dots, x_n, \dots$ dikatakan barisan.

Teorema 2.2.3

Setiap barisan bilangan real yang konvergen selalu terbatas [2].

Definisi 2.2.4

Diberikan ω yaitu koleksi semua barisan bilangan *real*, jadi :

$$\omega = \{\bar{x} = \{x_k\} : x_k \in \mathbb{R}\}$$

- Untuk setiap bilangan *real* p dengan $1 \leq p < \infty$ didefinisikan

$$l^p = \left\{ x \in \{x_j\} \in \omega : \sum_{j=1}^{\infty} |x_j|^p < \infty \right\}$$

dan norm pada l^p yaitu

$$\|x\|_p = \left(\sum_{j=1}^{\infty} |x_j|^p \right)^{\frac{1}{p}}$$

b. Untuk $p = \infty$ didefinisikan

$$l_{\infty} = \left\{ \bar{x} = \{x_k\} \in \omega : \sup_{k \geq 1} |x_k| < \infty \right\}$$

dan norm pada l_{∞} yaitu

$$\|x\|_{\infty} = \sup_{k \geq 1} |x_k|.$$

Definisi 2.2.5

Misal $p, q \in (1, \infty)$ dengan $\frac{1}{p} + \frac{1}{q} = 1$ (q konjugat p), untuk $x \in l^p$ dan $y \in l^p$ [3].

$$(x_j y_j)_{j \in N} \in l^{\infty} \text{ dan } \sum_{j=1}^{\infty} |x_j y_j| \leq \|x\|^p \|y\|^q \quad (3)$$

RUANG VEKTOR

Definisi 2.3.1

Ruang vektor adalah suatu himpunan tak kosong X yang dilengkapi dengan fungsi penjumlahan $(+): X \times X \rightarrow X$ dan fungsi perkalian skalar $(\cdot): F \times X \rightarrow X$ sehingga untuk setiap skalar λ, μ dengan elemen $x, y, z \in X$ berlaku:

- i. $x + y = y + x$
- ii. $(x + y) + z = x + (y + z)$
- iii. ada $\theta \in X$ sehingga $x + \theta = x$
- iv. ada $-x \in X$ sehingga $x + (-x) = \theta$
- v. $1 \cdot x = x$
- vi. $\lambda(x + y) = \lambda x + \lambda y$
- vii. $(\lambda + \mu)x = \lambda x + \mu x$
- viii. $\lambda(\mu x) = (\lambda \mu)x$

[3]

BASIS

Definisi 2.4.1

Ruang vektor V dikatakan terbangkitkan secara hingga (*finitely generated*) jika ada vektor-vektor $x_1, x_2, \dots, x_n \in V$ sehingga $V = [x_1, x_2, \dots, x_n]$. Dalam keadaan seperti itu $\{x_1, x_2, \dots, x_n\}$ disebut pembangkit (*generator*) ruang vektor V .

Menurut definisi di atas, ruang vektor V terbangkitkan secara hingga jika dan hanya jika ada vektor-vektor $x_1, x_2, \dots, x_n \in V$ sehingga untuk setiap vektor $x \in V$ ada skalar-skalar $\alpha_1, \alpha_2, \dots, \alpha_n$ sehingga

$$x = \alpha_1 x_1 + \alpha_2 x_2 + \cdots + \alpha_n x_n \quad (4)$$

Secara umum, jika $B \subset V$ dan V terbangkitkan oleh B , jadi $|B| = V$ atau B pembangkit V , maka untuk setiap $x \in V$ terdapat vektor-vektor $x_1, x_2, \dots, x_n \in B$ dan skalar $\alpha_1, \alpha_2, \dots, \alpha_n$ sehingga

$$x = \sum_{k=1}^n \alpha_k x_k$$

Definisi 2.4.2

Diberikan ruang vektor V . Himpunan $B \subset V$ dikatakan bebas linear jika setiap himpunan bagian hingga di dalam B bebas linear [3].

Definisi 2.4.3

Diberikan ruang vektor V atas lapangan $\mathcal{F}$. Himpunan $B \subset V$ disebut basis (*base*) V jika B bebas linear dan $|V| = B$.

Contoh :

Himpunan $\{\check{e}_1, \check{e}_2, \dots, \check{e}_n\}$, dengan $\check{e}_k$ vektor di dalam R^n yang komponen ke- k sama dengan 1 dan semua komponen lainnya sama dengan 0, merupakan basis ruang vektor R^n [3].

RUANG BERNORMA

Definisi 2.5.1

Diberikan ruang linear X . Fungsi $x \in X \rightarrow \|x\| \in R$ yang mempunyai sifat-sifat :

- i. $\|x\| \geq 0$ untuk setiap $x \in X$
- ii. $\|x\| = 0$, jika dan hanya jika $x = 0$, (0 vektor nol)
- iii. $\|\alpha x\| = \|\alpha\| \cdot \|x\|$ untuk setiap skalar α dan $x \in X$
- iv. $\|x + y\| \leq \|x\| + \|y\|$ untuk setiap $x, y \in X$

disebut norma (*norm*) pada X dan bilangan nonnegatif $\|x\|$ disebut norma vektor x . Ruang linear X yang dilengkapi dengan suatu norma $\|\cdot\|$ disebut ruang bernorma (*norm space*) dan dituliskan singkat dengan $X, \|\cdot\|$ atau X saja asalkan normanya telah diketahui [3].

RUANG BANACH

Definisi 2.6.1

Ruang Banach (*Banach space*) adalah ruang bernorma yang lengkap (sebagai ruang metrik yang lengkap) jika dalam suatu ruang bernorm X berlaku kondisi bahwa setiap barisan Cauchy di X adalah konvergen [3].

3. METODOLOGI PENELITIAN

Langkah-langkah yang digunakan dalam penelitian ini antara lain :

1. Mengkonstruksikan operator A dari ruang barisan terbatas l_{13} ke ruang barisan l_{13} dengan basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$.
2. Mengkonstruksikan norma operator A
3. Menyelidiki apakah koleksi semua operator A membentuk ruang Banach
4. Merepresentasikan operator A sebagai matriks takhingga yang dikerjakan dari ruang barisan terbatas l_{13} ke ruang barisan l_{13} dengan basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$.

4. PEMBAHASAN

Berdasarkan pada pembahasan sebelumnya, yang dimaksud dengan ruang barisan terbatas l_{13} dapat didefinisikan sebagai :

$$l_{13} = \left\{ x = (x_i) \left| \left(\sum_{i=1}^{\infty} |x_i|^{13} \right)^{\frac{1}{13}} < \infty \right. \right\}$$

dengan $(x_i) = (x_1, x_2, \dots)$ merupakan barisan bilangan real $\mathbb{R}$. Ruang barisan l_{13} merupakan ruang Banach dengan ruang dual $(l_{13})^* = \{x^*: l_{13} \rightarrow \mathbb{R}\}$ yaitu koleksi semua fungsional linear dan kontinu pada l_{13} .

Untuk sebarang $x^* \in (l_{13})^*$ dan $x \in l_{13}$, penulisan $\langle x, x^* \rangle$ dimaksudkan sebagai fungsional x^* pada x atau $x^*(x)$. Barisan vektor $\{e_n\} \subset l_{13}$ dinamakan basis pada l_{13} jika untuk setiap vektor $x \in l_{13}$ terdapat barisan skalar yang tunggal $\{a_n\}$ sehingga

$$x = \sum_{n=1}^{\infty} a_n e_n$$

Barisan $\{e_n^*\} \in (l_{13})^*$ dengan $\|e_n^*\| = 1$ untuk setiap n dikatakan biortonormal terhadap basis $\{e_n\} \subset l_{13}$ jika

$$\langle e_m, e_n^* \rangle = \delta_{mn}$$

dengan $\delta_{mn} = 1$ untuk $m = n$ dan $\delta_{mn} = 0$ untuk $m \neq n$. Selanjutnya, pasangan $\{\{e_n\}, \{e_n^*\}\}$ disebut sistem biortonormal pada l_{13} maka

$$x = \sum_{n=1}^{\infty} a_n e_n$$

dengan $\langle x, e_n^* \rangle = a_n$.

Jika $A \in \mathcal{L}_c(l_{13}, l_{13/12})$ maka operator $A^* \in \mathcal{L}_c((l_{13})^*, (l_{13/12})^*)$ disebut operator pendamping A jika dan hanya jika untuk setiap $x \in l_{13}$ dan $y^* \in (l_{13})^*$, berlaku

$$\langle A(x), y^* \rangle = \langle x, A^*(y^*) \rangle$$

Jadi, jika $\{e_n\} \subset l_{13}$ dan $\{d_m^*\} \in (l_{13/12})^*$ diperoleh

$$\langle A(e_n), d_m^* \rangle = \langle e_n^*, A^*(d_m^*) \rangle$$

Jika $\{e_n\}, \{f_n\} \subset l_{13}$ basis pada l_{13} dan $\{d_m\} \subset l_{13/12}$, maka untuk setiap $A \in \mathcal{L}_c(l_{13}, l_{13/12})$ berlaku

$$A^*(d_m^*) = \sum_{n=1}^{\infty} \langle e_n, A^*(d_m^*) \rangle e_n^* = \sum_{n=1}^{\infty} \langle A(e_n), d_m^* \rangle e_n^* \quad (a)$$

Dan

$$A^*(d_m^*) = \sum_{n=1}^{\infty} \langle e_n, A^*(d_m^*) \rangle f_n^* = \sum_{n=1}^{\infty} \langle A(e_n), d_m^* \rangle f_n^* \quad (b)$$

Berdasarkan persamaan (a) dan (b) diperoleh

$$\sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(f_k), d_m^* \rangle f_k^* \right\| \quad (c)$$

Berdasarkan persamaan (a), (b) dan (c) didefinisikan pengertian operator dari ruang barisan l_{13} ke ruang barisan $l_{13/12}$ sebagai berikut

Definisi 1.1 Operator $A \in \mathcal{L}_c(l_{13}, l_{13/12})$ merupakan operator-SM jika

$$\sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| < \infty$$

Dengan $\{e_n\} \subset l_{13}$ basis pada l_{13} . Dan $\{d_m\} \subset l_{13/12}$ basis pada $l_{13/12}$.

Dapat dipahami bahwa bilangan $\|A\|$ dengan

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| < \infty$$

Tidak bergantung pada pemilihan basis $\{e_n\}$ pada l_{13} .

Selanjutnya, notasi $SM(l_{13}, l_{13/12})$ menyatakan koleksi semua operator-SM dari ruang barisan l_{13} ke ruang barisan $l_{13/12}$.

Teorema 1.2 Untuk setiap $A \in SM(l_{13}, l_{13/12})$ berlaku

- i. $\|A\| \leq \|A\|_{SM}$
- ii. $SM(l_{13}, l_{13/12})$ merupakan ruang Banach terhadap norma $\|\cdot\|_{SM}$
- iii. Jika $A \in SM(l_{13}, l_{13/12})$ maka A operator kompak.

Bukti :

- i. Diambil sebarang $\{e_n\} \subset l_{13}$ basis pada l_{13} , $\{d_m\} \subset l_{13/12}$ basis pada $l_{13/12}$ dan $x \in l_{13}$, maka berdasarkan (a), (b) dan (c) diperoleh

$$\begin{aligned} \|A(x)\| &= \left\| \sum_{m=1}^{\infty} \langle A(x), d_m^* \rangle d_m \right\| \\ &= \left\| \sum_{m=1}^{\infty} \langle x, A^*(d_m^*) \rangle d_m \right\| \end{aligned}$$

$$\begin{aligned} &\leq \|x\| \sum_{m=1}^{\infty} \|A^*(d_m^*)\| \\ &\|x\| \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= \|x\| \|A\|_{SM} \end{aligned}$$

Yang berakibat $\|A\| \leq \|A\|_{SM}$.

- ii. Pertama ditunjukkan bahwa $SM(l_{13}, l_{13/12})$ merupakan ruang bernorma terhadap norma $\|\cdot\|_{SM}$ sebab :

- a.) Untuk setiap $A \in SM(l_{13}, l_{13/12})$

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \geq 0$$

dan

$$\|A\|_{SM} = 0 \Leftrightarrow \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| = 0$$

$$\Leftrightarrow \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* = A^*(d_m^*) = \theta \text{ (untuk setiap } m)$$

$$\Leftrightarrow A^* = 0 \text{ (operator nol)}$$

$$\Leftrightarrow A = 0 \text{ (operator nol)}$$

- b.) Untuk setiap $A \in SM(l_{13}, l_{13/12})$ dan skalar α , diperoleh

$$\begin{aligned} \|\alpha A\|_{SM} &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= |\alpha| \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A(e_k), d_m^* \rangle e_k^* \right\| \\ &= |\alpha| \|A\|_{SM} \end{aligned}$$

- c.) Jika diberikan $A_1, A_2 \in SM(l_{13}, l_{13/12})$ maka

$$\begin{aligned} \|A_1 + A_2\|_{SM} &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle (A_1 + A_2)(e_k), d_m^* \rangle e_k^* \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle (A_1)(e_k) + (A_2)(e_k), d_m^* \rangle e_k^* \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A_1(e_k), d_m^* \rangle e_k^* + \sum_{k=1}^{\infty} \langle A_2(e_k), d_m^* \rangle e_k^* \right\| \\ &\leq \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} \langle A_1(e_k), d_m^* \rangle e_k^* \right\| + \left\| \sum_{k=1}^{\infty} \langle A_2(e_k), d_m^* \rangle e_k^* \right\| \end{aligned}$$

Dengan kata lain,

$$\|A_1 + A_2\|_{SM} \leq \|A_1\|_{SM} + \|A_2\|_{SM}$$

Selanjutnya menunjukkan kelengkapan ruang $SM(l_{13}, l_{13/12})$ sebagai berikut :

Diambil sebarang barisan Cauchy $\{A_i\} \subset SM(l_{13}, l_{13/12})$ Untuk setiap bilangan $\varepsilon > 0$, terdapat bilangan bulat positif n_0 sehingga untuk setiap bilangan bulat positif $i, j \geq n_0$, berlaku

$$\|A_i - A_j\|_{SM} \leq \frac{\varepsilon}{13}$$

Akan dibuktikan bahwa terdapat $A \in SM(l_{13}, l_{13/12})$ sehingga

$$\lim_{i \rightarrow \infty} \|A_i - A\|_{SM} = 0$$

Karena

$$\|A_i - A_j\|_{\mathcal{L}_c(l_{13}, l_{13/12})} \leq \|A_i - A_j\|_{SM} < \frac{\varepsilon}{13}$$

Untuk setiap $A_i, A_j \in SM(l_{13}, l_{13/12})$ dengan $i, j \geq n_0$, maka barisan $\{A_i\}$ juga merupakan barisan Cauchy di dalam $\mathcal{L}_c(l_{13}, l_{13/12})$. Karena $\mathcal{L}_c(l_{13}, l_{13/12})$ ruang lengkap maka terdapat $A \in \mathcal{L}_c(l_{13}, l_{13/12})$ sehingga $\lim_{j \rightarrow \infty} A_j = A$. Oleh karena itu,

$$\begin{aligned} & \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} ((A_i - A)(e_k), d_m^*) e_k^* \right\| \\ &= \lim_{j \rightarrow \infty} \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} ((A_i - A)(e_k), d_m^*) e_k^* \right\| \\ &= \lim_{j \rightarrow \infty} \|A_i - A_j\|_{SM} < \frac{\varepsilon}{13} \end{aligned}$$

Untuk sebarang bilangan bulat $i \geq n_0$. Dengan kata lain, $A - A_i \in SM(l_{13}, l_{13/12})$, untuk $i \geq n_0$. Oleh karena itu, $A - A_{n_0} + A_{n_0} = A \in SM(l_{13}, l_{13/12})$ dan terbukti bahwa barisan $\{A_i\}$ konvergen ke suatu $A \in SM(l_{13}, l_{13/12})$. Jadi, $SM(l_{13}, l_{13/12})$ merupakan ruang bernorma yang lengkap atau ruang Banach.

iii. Jika $A \in SM(l_{13}, l_{13/12})$ dan $x \in l_{13}$, maka

$$A(x) = \sum_{m=1}^{\infty} \langle A(x), d_m^* \rangle d_m$$

Oleh karena itu, untuk setiap bilangan bulat positif n , dapat didefinisikan operator $A_n: l_{13} \rightarrow l_{13/12}$ dengan

$$A_n(x) = \sum_{m=1}^n \langle A(x), d_m^* \rangle d_m$$

Jelas bahwa $A_n \in \mathcal{L}_c(l_{13}, l_{13/12})$ dan A_n merupakan operator berhingga. Dengan kata lain, A_n operator kompak. Karena $\{A_n\}$ konvergen ke K maka K operator kompak.

Berdasarkan Teorema 1.2 diperoleh

Akibat 1.3 $SM(l_{13}, l_{13/12}) \subset K(l_{13}, l_{13/12}) \subset \mathcal{L}_c(l_{13}, l_{13/12})$ dengan $K(l_{13}, l_{13/12})$ koleksi operator kompak dari l_{13} ke $l_{13/12}$.

Operator $A \in SM(l_{13}, l_{13/12})$ dapat diwakili oleh matriks tak hingga $A = A_{\infty \times \infty}$. Oleh karena itu, dalam bentuk matriks tak hingga karakteristik operator-SM tersebut dapat diuraikan sebagai berikut:
Teorema 1.4 Suatu operator linear kontinu $A: l_{13} \rightarrow l_{13/12}$ merupakan operator-SM jika dan hanya jika terdapat suatu matriks (a_{ij}) yang memenuhi :

- i.) $Ax = \{\sum_{j=1}^{\infty} a_{ij}x_j\} \in l_{13/12}$ untuk setiap $x = (x_i) \in l_{13}$
- ii.) $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^{13} < \infty$
- iii.) $\sum_{i=1}^{\infty} |\sum_{j=1}^{\infty} a_{ij}| < \infty$

Bukti :

(Syarat perlu) karena $A: l_{13} \rightarrow l_{13/12}$ linear dan kontinu maka dengan sendirinya berlaku i) dan ii).

Operator A dalam bentuk matriks (a_{ij}) dikerjakan pada basis standar $\{e_k\}$ dengan $e_k = (0, 0, \dots, 1_{(k)}, \dots)$ berbentuk

$$Ae_k = (a_{jk})_{j=1}^{\infty}$$

Karena $A = (a_{ij})$ operator-SM maka

$$\begin{aligned} \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} (Ae_k, d_m^*) d_m \right\| &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} d_m \right\| \\ &= \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} \right\| < \infty \end{aligned}$$

(Syarat cukup) berdasarkan i) dan ii) maka $A = (a_{ij}): l_{13} \rightarrow l_{13/12}$ linear dan kontinu.

Selanjutnya, berdasarkan iii) diperoleh

$$\|A\|_{SM} = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} (Ae_k, d_m^*) d_m \right\| = \sum_{m=1}^{\infty} \left\| \sum_{k=1}^{\infty} a_{mk} \right\| < \infty$$

Terbukti $A = (a_{ij})$ merupakan operator-SM

Contoh 1.5 Matriks $A = (a_{ij})$ dengan

$$a_{ij} = \begin{cases} \frac{1}{i^{13/12}} & i = j \\ 0 & i \neq j \end{cases}$$

Merepresentasikan operator-SM $A: l_{13} \rightarrow l_{13/12}$ sebab :

- i.) Untuk setiap $x = (x_j) \in l_{13}$ berlaku

$$\|Ax\| = \left\| \sum_{j=1}^{\infty} a_{ij}x_j \right\| = \left(\sum_{j=1}^{\infty} \left| \frac{x_j}{j^{13/12}} \right|^{13/12} \right)^{\frac{12}{13}} < \left(\sum_{j=1}^{\infty} |x_j|^{13/12} \right)^{\frac{12}{13}} < \infty$$

Jadi, $Ax \in l_{13/12}$

ii.) Bagian kedua terpenuhi sebab

$$\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^{13/12} = \sum_{i=1}^{\infty} \left| \frac{1}{i^{13/12}} \right|^{13/12} < \infty$$

iii.) Bagian ketiga terpenuhi sebab

$$\sum_{i=1}^{\infty} \left| \sum_{j=1}^{\infty} a_{ij} \right| = \sum_{i=1}^{\infty} \frac{1}{i^{13/12}} < \infty$$

5. SIMPULAN

Operator linear dan kontinu $A : l_{13} \rightarrow l_{13/12}$ merupakan operator-SM jika dan hanya jika terdapat suatu matriks $A = (a_{ij})$ yang memenuhi :

1. $Ax = \{\sum_{j=1}^{\infty} a_{ij}x_j\} \in l_{13/12}$ untuk setiap $x = (x_i) \in l_{13}$
2. $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}|^{13/12} < \infty$
3. $\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} |a_{ij}| < \infty$

Koleksi semua operator SM $A : l_{13} \rightarrow l_{13/12}$ yang dinotasikan dengan SM $(l_{13}, l_{13/12})$ membentuk ruang Banach.

KEPUSTAKAAN

- [1] Kreyszig, E. (1989). *Introductory Function Analysis with Application*. Willey Classic Library, New York.
- [2] Martono, K. 1984. *Kalkulus dan Ilmu Ukur Analitik 2*. Angkasa, Bandung.
- [3] Darmawijaya, S. (2007). *Pengantar Analisis Abstrak*. Universitas Gajah Mada, Yogyakarta.

DESAIN KONTROL MODEL SUHU RUANGAN

¹Zulfikar Fakhri Bismar dan ²Aang Nuryaman

^{1,2} Jurusan Matematika FMIPA Universitas Lampung

Email: ¹zulfikar.fb95@gmail.com dan ²aangnuryaman@gmail.com

ABSTRAK

Pada paper ini, akan dikaji model pengendalian suhu di sebuah ruangan dengan variabel terkait adalah suhu ruangan, suhu kaca termometer, dan suhu air raksa. Dengan menggunakan model hasil linierisasi di sekitar titik kesetimbangan, dikonstruksi kontrol input yang mungkin agar suhu ruangan sesuai dengan yang diinginkan. Profil dari dinamik model hasil linierisasi akan disimulasikan secara numerik untuk beberapa alternatif kontrol dan nilai parameter berbeda.

Kata Kunci : Suhu, Ruangan, Kontrol Input, Linierisasi.

1. PENDAHULUAN

Ruangan yang nyaman diperlukan untuk melakukan berbagai aktivitas sehari-hari. Untuk membuat keadaan ruangan menjadi nyaman, salah satu acuannya untuk kenyamanan suatu ruangan adalah suhu pada ruangan tersebut. Karena suhu ruangnya terlalu panas ataupun terlalu dingin, mempersulit seseorang untuk melakukan aktivitas di dalam ruangan tersebut sebagaimana mestinya. Untuk itu, diperlukan model pengendalian suhu ruangan yang baik agar suhu pada ruangan tersebut dapat dinikmati sesuai yang diinginkan. Alat yang mengendalikan suhu ruangan terletak pada katup udara pada ruangan. Model yang akan dibuat merupakan model untuk mencari pengaturan pada katup udara yang terbaik untuk membuat suhu ruangan dapat dinikmati untuk beraktivitas sebagaimana mestinya.

2. LANDASAN TEORI

Pemodelan sistem dinamik didefinisikan sebagai pembuatan model yang bergantung terhadap waktu untuk sistem fisis. model sistem dinamik merupakan salah satu tipe model yang sangat vital untuk merepresentasikan bagaimana suatu sistem bekerja dan berubah keadaannya dari waktu ke waktu. Gagasan dari pemodelan sistem dinamik yaitu untuk digunakan sebagai suatu hipotesis yang perlu dibuktikan pada konteks fenomena yang terjadi di dunia [1].

Sistem kontrol merupakan sistem dinamik yang menggabungkan *input* kontrol yang didesain untuk mencapai tujuan pengontrolan. Sistem kontrol merupakan dimensi yang terbatas jika ruang fasenya (contohnya ruang vektor atau semacamnya) merupakan dimensi yang terbatas. Waktu yang kontinu pada sistem kontrol dirumuskan sebagai berikut.

$$\frac{dx}{dt}(t) = f(x(t); u(t)), \quad x \in X, u \in U, \text{ dan } t \in \mathbb{R} \quad (1)$$

Dengan $x \in X$, melambangkan keadaan pada sistem, sedangkan $u \in U$, X dan U merupakan himpunan terbuka dengan $X \subset \mathbb{R}^n$ dan $U \subset \mathbb{R}^m$ melambangkan *input* pada sistem. pemetaan fungsi $f : X \times U \rightarrow U$ merupakan fungsi nonlinier yang mendekati solusi analitik dari suatu sistem kontrol. $\frac{dx}{dt} = \dot{x}$ sebagai laju perubahansuhu ruangan. [2].

Model ruang keadaan atau *state-space model* ditetapkan oleh tiga matriks A, B dan C. ketiga matriks ini dinamakan matriks sistem, ditambah dengan matriks *co-variant* dari vektor gangguan dari luar. Didalam kerangka gambaran Ruang Keadaan terdapat banyak variasinya. Salah satu bentuk umum Ruang Keadaan yang sering digunakan adalah.

$$\dot{\mathbf{x}} = \mathbf{A}\mathbf{x} + \mathbf{B}\mathbf{u} \quad (2)$$

$$\dot{\mathbf{y}} = \mathbf{C}\mathbf{x} + \mathbf{D}\mathbf{u} \quad (3)$$

Dengan vektor *input* $\mathbf{u}$, vektor $\mathbf{x}$ disebut vektor keadaan. Persamaan yang pertama menentukan bagaimana vektor-vektor tersebut terhubung secara halus dengan data $\dot{\mathbf{y}}$ yang berubah dari waktu ke waktu. Matriks A sendiri merupakan matriks dinamik dan nilai eigen dari A sangat penting dan menentukan bagaimana sistem bekerja. Persamaan kedua adalah persamaan observasi yang menentukan hubungan antara data dengan vektor keadaan $\mathbf{x}$ tersebut [3].

Suatu sistem dikatakan terkontrol dengan baik pada waktu t_0 jika untuk setiap keadaan $\mathbf{x}(t_0)$ pada suatu model ruang keadaan dan setiap $\mathbf{x}(t_1)$ pada model tersebut, terdapat waktu dengan $t_1 > t_0$ serta *input* $\mathbf{u}_{[t_0, t_1]}$ sehingga akan memindahkan keadaan $\mathbf{x}(t_0)$ menjadi keadaan $\mathbf{x}(t_1)$ pada waktu t_1 . Lain dari itu, maka sistem dapat dikatakan tidak terkontrol pada waktu t_0 . Sifat keterkontrolan merupakan sifat yang menghubungkan antara *input* dan keadaan benda. Oleh karena itu, sifat ini akan melibatkan matriks A dan B pada sistem pengendalian suhu ruangan. Dari definisi keterkontrolan, dapat diamati jika sistem dapat terkontrol, berarti suhu ruangan dapat diperlakukan atau dikendalikan sesuai keinginan. Akan tetapi, jika suatu sistem bukan sistem yg terkontrol, maka suhu ruangan tidak akan dapat dikembalikan menuju kesetimbangan sesuai yang diinginkan. Sistem pengendalian suhu ruangan dikatakan terobservasi pada waktu t_0 jika untuk setiap $\mathbf{x}(t_0)$ pada waktu t_0 pada suatu model ruang keadaan, terdapat waktu t_1 dengan $t_1 > t_0$ sehingga *input* $\mathbf{u}_{[t_0, t_1]}$ dan *output* $\mathbf{y}_{[t_0, t_1]}$ cukup untuk menentukan keadaan $\mathbf{x}(t_0)$. Selain dari itu, sistem tersebut dikatakan tidak terobservasi pada waktu t_0 [4].

[5] Suatu sistem kontrol didefinisikan dengan

$$\frac{d}{dt} \mathbf{x} = \mathbf{A}\mathbf{x} + \mathbf{B}\mathbf{u} \quad (4)$$

pernyataan diatas ekuivalen dengan :

1. Sistem (2.9) terkontrol.
2. $\text{Rank}(\mathbf{B}, \mathbf{A}\mathbf{B}, \mathbf{A}^2\mathbf{B}, \dots, \mathbf{A}^{n-1}\mathbf{B}) = n$ [5].

Pada persamaan (2.2) dan (2.3) dengan $\dim(\mathbf{x}) = n$, sistem tersebut dapat dikatakan sistem terkontrol (terkontrol) jika $\text{rank}(\mathbf{B}, \mathbf{A}\mathbf{B}, \mathbf{A}^2\mathbf{B}, \dots, \mathbf{A}^{n-1}\mathbf{B}) = n$, serta dapat dikatakan terobservasi jika $\text{rank}(\mathbf{C}^T, \mathbf{A}^T\mathbf{C}^T, (\mathbf{A}^T)^2\mathbf{C}^T, \dots, (\mathbf{A}^T)^{n-1}\mathbf{C}^T) = n$.

Permasalahan mengenai keterkendalian suatu sistem pada umumnya secara matematis jika diberikan keadaan awal $\mathbf{x}_0 \in X$, keadaan terakhir $\mathbf{x}_T \in X$, dan waktu $T > 0$ digambarkan sebagai berikut.[1]

$$\begin{cases} \dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t), \mathbf{u}(t)) \\ \mathbf{x}(0) = \mathbf{x}_0 \end{cases} \quad (5)$$

Memenuhi $\mathbf{x}(T) = \mathbf{x}_T$.

Posisi titik *equilibrium* untuk sistem diatas adalah titik $\bar{x} \in X$ sehingga terdapat nilai $\bar{u} \in U$ (khususnya bernilai nol) sedemikian sehingga $f(\bar{x}, \bar{u}) = 0$. Hukum umpan balik kestabilan asimtotik adalah fungsi $k: X \rightarrow U$ dengan $k(\bar{x}) = \bar{u}$, sehingga sistem (2.10) berubah menjadi.

$$\dot{\mathbf{x}}(t) = f(\mathbf{x}(t), \mathbf{k}(\mathbf{x})) \quad (6)$$

Yang diperoleh dengan memasukkan *input* kontrol (2.12) pada persamaan (2.10).

$$\mathbf{u}(t) = \mathbf{k}(\mathbf{x}(t)) \quad (7)$$

Jika $\mathbf{x}$ bergerak mendekati nol sebagaimana $t \rightarrow \infty$. Keadaan $\mathbf{x}(t)$ yang terbentuk dari hukum umpan balik ini merupakan fungsi kontinu, yang berarti solusi dari persamaan (2.11) untuk setiap fungsi $\mathbf{x}(t)$ yang berlaku untuk semua t [2].

Bentuk umum dari model ruang keadaan memiliki bentuk lain sebagai berikut.

$$\begin{cases} \dot{\mathbf{x}}(t) = \mathbf{A}\mathbf{x} \\ \mathbf{x}(0) = \mathbf{x}_0 \end{cases} \quad (8)$$

Terdapat kriteria kestabilan sistem sederhana yang hanya bergantung pada nilai eigen $\lambda_{1,2}$ pada matriks $\mathbf{A}$. Jika semua nilai eigen memiliki bagian nilai real negatif, maka setiap solusi yang terbentuk akan bergerak mendekati nol sebagaimana $t \rightarrow \infty$, dan solusi trivial tersebut menunjukkan bahwa sistem dapat dikatakan stabil asimtotik. Jika satu saja nilai eigen memiliki bagian nilai real positif, maka beberapa solusi akan bergerak menjauhi titik *equilibrium* membentuk kurva eksponensial, sehingga menimbulkan sistem yang tidak stabil [6].

3. METODOLOGI PENELITIAN

Waktu dan Tempat Penelitian

Penelitian dilakukan pada Semester Ganjil Tahun Ajaran 2017/2018 di Jurusan Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lampung.

Metode Penelitian

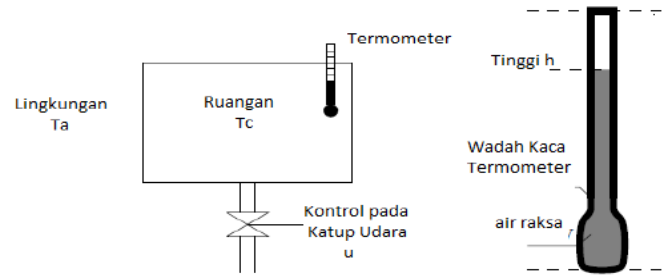
Penelitian ini dilakukan secara simulasi dengan menggunakan *software* MATLAB[®]. Adapun langkah-langkah yang dilakukan dalam penelitian ini adalah sebagai berikut :

1. Membuat model matematika sistem dinamik dari suhu ruangan dengan menggunakan relasi dinamik berdasarkan asumsi yang ditetapkan.
2. Menentukan titik *equilibrium* suhu ruangan beserta titik *equilibrium* untuk tinggi air raksa.
3. Mengkaji respon step untuk melihat respon yang terjadi tanpa kontrol input.
4. Menentukan beberapa alternatif kontrol dengan memvariasikan parameter yang terkait dengan input kontrol tersebut.
5. Mensimulasikan sistem yang sudah diberi kontrol *input* serta memvariasikan parameter yang terkait dengan kontrol *input* untuk melihat apa yang terjadi pada sistem.
6. Menginterpretasikan hasil yang didapat dan kemudian mengambil kesimpulan.

4. HASIL DAN PEMBAHASAN

Model Dinamik

Pada pembahasan ini dirancang suatu model kontrol suhu ruangan, yang bertujuan untuk menganalisa kontrol pada katup udara yang berfungsi mengendalikan laju aliran panas pada ruangan tersebut. Adapun pengaturan pada katup udara tersebut ditinjau berdasarkan pengukuran dari termometer. Suhu dari ruangan tersebut juga dipengaruhi oleh suhu lingkungan. Termometer yang digunakan yaitu termometer raksa. Pengukuran suhu tidak diketahui langsung dari suhu ruangnya, melainkan dengan memperhatikan tinggi air raksa pada termometer. Adapun gambaran kontrol suhu ruangan yang terjadi antara lain seperti pada gambar berikut.



Gambar 1. Kontrol suhu udara pada ruangan.

Berdasarkan gambar diatas, akan dibentuk model dinamik yang terjadi pada ruangan tersebut. Berikut merupakan variabel nyata, variabel-variabel pokok yang akan dimodelkan.

- u : Input kontrol pada katup udara (variabel kontrol).
- T_a : Suhu lingkungan (sebagai factor gangguan dari luar).
- T_c : Suhu ruangan (variabel yang dikontrol atau sebagai output).
- h : Tinggi raksa pada termometer (variabel yang diukur).

Untuk mendapatkan model yang sesuai, diperhatikan relasi antara variabel-variabel pokok, untuk mengetahuinya, diperlukan beberapa variabel-variabel bantu. Antara lain yaitu.

- q : Nilai kalor yang berpindah dari katup udara menuju ruangan.
- q_a : Nilai kalor yang berpindah dari lingkungan menuju ruangan.
- q_g : Nilai laju perpindahan panas dari wadah termometer menuju ruangan.
- q_{Hg} : Nilai laju perpindahan panas dari air raksa menuju wadah termometer.
- T_g : Suhu wadah termometer.
- T_{Hg} : Suhu pada air raksa.
- A_g : Luas penampang wadah termometer.
- V_g : Volume bagian dalam wadah termometer.
- V_{Hg} : Volume air raksa.

Sehingga didapatkan relasi antarvariabel sebagai berikut.

$$\begin{aligned}
 q &= a_0 u, \\
 q_a &= a_1(T_a - T_c), \quad q_g = a_2(T_c - T_g), \quad q_{Hg} = a_3(T_g - T_{Hg}), \\
 \frac{dT_c}{dt} &= b_1(q_a + q), \quad \frac{dT_g}{dt} = b_2(q_g + q_{Hg}), \quad \frac{dT_{Hg}}{dt} = b_3 q_{Hg}
 \end{aligned} \tag{9}$$

Dengan a_0, a_1, a_2, a_3 merupakan konstanta satuan termal yang bernilai positif, b_1, b_2 , dan b_3 merupakan konstanta satuan laju perpindahan panas per detik yang bernilai positif. Bentuk relasi diatas mengasumsikan bahwa nilai laju perpindahan panas dari katup udara menuju ruangan berbanding lurus dengan kontrol pada katup udara. Laju perubahan suhu pada ruangan berkaitan dengan suhu lingkungan dan suhu ruangan. Serta laju perubahan suhu pada wadah termometer berkaitan dengan suhu ruangan dan suhu pada wadah termometer. serta raksa didalamnya memiliki kaitan dengan suhu air raksa.

Selanjutnya akan diperkirakan persamaan yang menunjukkan tinggi air raksa pada termometer. Dengan asumsi tinggi raksa memiliki kaitan dengan suhu wadah termometer beserta suhu raksa itu sendiri. Oleh karena itu, A_g sebanding dengan T_g^2 , V_g sebanding dengan T_g^3 , dan V_{Hg} sebanding dengan T_{Hg}^3 . Sehingga diperoleh model berikut.

$$A_g = c_1 T_g^2, V_g = c_2 T_g^3, V_{Hg} = c_3 T_{Hg}^3, h = \frac{V_{Hg} - V_g}{A_g} \quad (10)$$

Dengan c_1 sebagai konstanta satuan muai penampang termometer terhadap suhu, dan c_2 serta c_3 sebagai konstanta satuan muai volume terhadap suhu. Baik c_1, c_2 , maupun c_3 bernilai positif.

4. Hasil dan Pembahasan

Model laju perubahan suhu pada persamaan (9) dan tinggi raksa pada persamaan (10) dapat dibentuk menjadi model ruang keadaan. Dengan u dan T_a sebagai variabel input, serta h dan T_c merupakan variabel output. Namun, perlu diketahui letak titik kesetimbangan untuk tinggi air raksa beserta suhu ruangan agar dapat ditentukan pada di suhu berapakah ruangan dapat digunakan dengan nyaman sebagaimana mestinya.

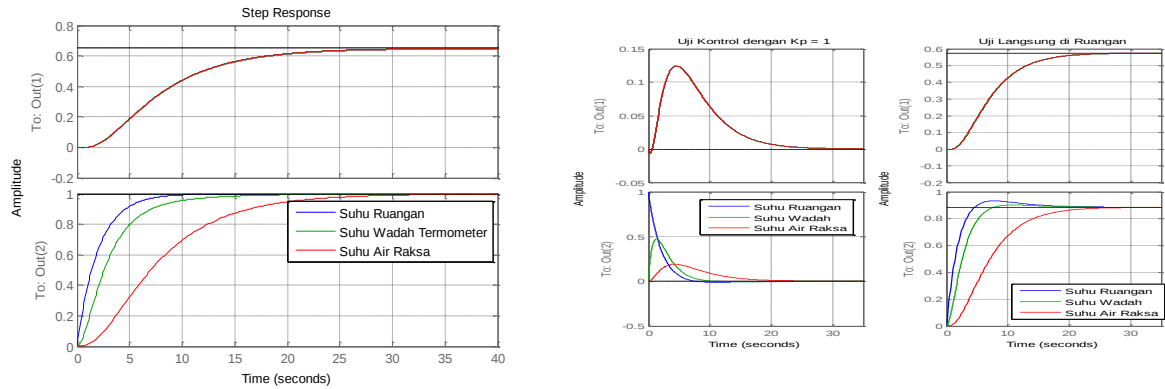
Letak titik equilibrium dapat diketahui dengan mempertimbangkan salah satu teori tentang kesetimbangan, yaitu terdapat nilai input u khususnya bernilai nol. Yang berarti dengan asumsi tanpa gangguan dari suhu lingkungan serta kontrol pada katup sedemikian sehingga $f(x,u) = 0$. Dengan permisalan tersebut, diperoleh titik equilibrium suhu ruangan T_c , suhu wadah termometer T_g , dan suhu air raksa T_{Hg} bernilai nol. Dengan kata lain, tujuan dari desain pengendalian suhu ruangan ini adalah dengan mengembalikan suhu ruangan kembali ke titik kesetimbangan suhu ruangan, yaitu nol. Dari hasil linearisasi disekitar titik equilibrium serta penyatuan hasil kali variabel bebas diperoleh model umum ruang keadaan berikut

$$\begin{aligned} \frac{dT_c}{dt} &= \alpha_1(T_a - T_c) + \beta_1 u, \\ \frac{dT_g}{dt} &= \alpha_2(T_c - T_g) + \alpha_3(T_{Hg} - T_g), \end{aligned} \quad (11)$$

$$\begin{aligned} \frac{dT_{Hg}}{dt} &= \alpha_4(T_g - T_{Hg}) \\ h &= \gamma_1 T_{Hg} - \gamma_2 T_g \end{aligned} \quad (12)$$

Dengan nilai $\beta_1 = 0.1\alpha_1 = 0.5$, $\alpha_2 = 1$, $\alpha_3 = 0.1$, $\alpha_4 = 0.2$, $\gamma_1 = 0.7$, $\gamma_2 = 0.05$.

Pertama-tama, disimulasikan respon step yang terjadi untuk melihat apa yang akan terjadi bila suhu ruangan belum dikendalikan. Dengan suhu lingkungan T_a merupakan fungsi unit step. diperoleh gambaran sebagai berikut.



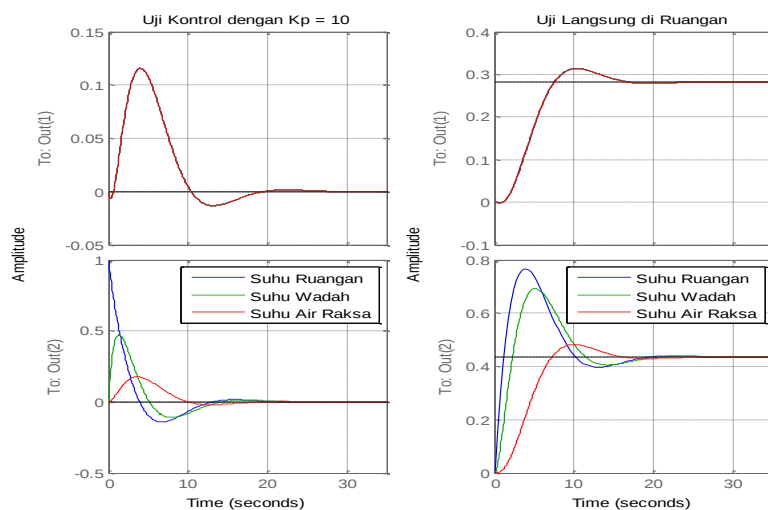
Gambar 2. Respon tanpa kontrol

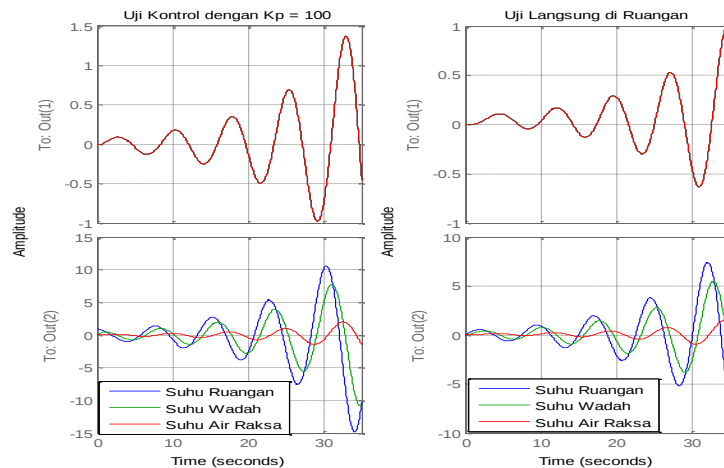
Gambar menunjukkan bahwa suhu akan bergerak menuju titik bernilai satu. Hal ini berarti bahwa ruangan akan menjadi panas bila tidak dikontrol. Oleh karena itu, akan dicoba beberapa alternatif kontrol pada katup udara agar suhu ruangan kembali lagi menuju titik kesetimbangan.

Selanjutnya untuk dicoba kontrol dengan asumsi suhu ruangan akan sebanding dengan tinggi raksa pada termometer. Sehingga didapat alternatif kontrol u sebagai berikut.

$$u = -K_p h \quad (13)$$

Dengan K_p merupakan nilai konstanta positif. Akan dicoba beberapa nilai K_p bernilai 1, 10, dan 100. Dari alternatif kontrol tersebut didapatkan gambaran sebagai berikut.





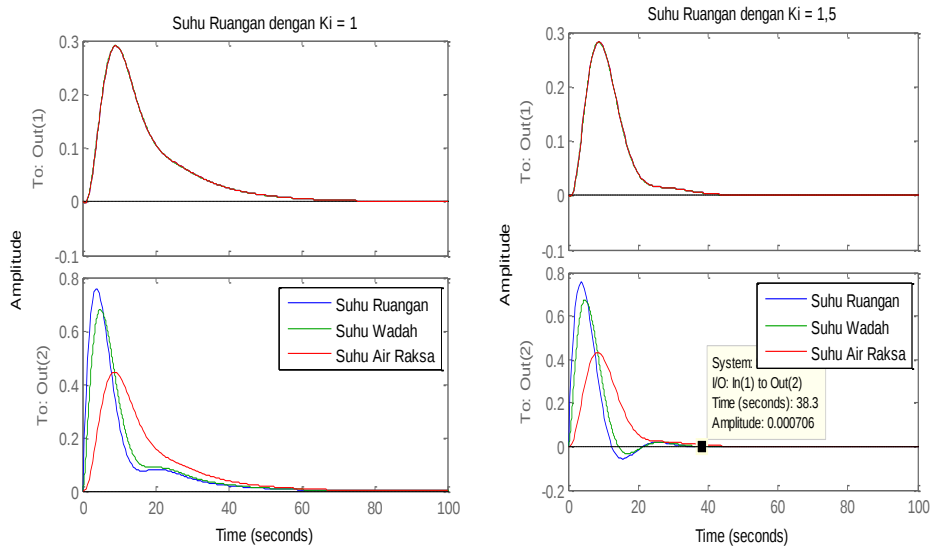
Gambar 3. Skema simulasi kontrol suhu ruangan dengan alternatif kontrol pertama.

Kurva atas merupakan kurva tinggi raksa, dan kurva bawah merupakan kurva suhu ruangan, suhu wadah thermometer, beserta suhu air raksa. Dari gambar berikut dapat dilihat bahwa dengan nilai $K_p = 1$ dan $K_p = 10$, terlihat bahwa pada uji kontrol, waktu yang dibutuhkan cukup lama untuk mencapai equilibrium, hal ini dikarenakan matriks dinamik dari alternatif kontrol dengan nilai konstanta K_p diatas bernilai real negatif bernilai kecil. Namun, jika dicoba langsung di ruangan yang memiliki suhu lingkungan sebagai gangguan, suhu ruangan masih belum dapat kembali ke titik kesetimbangan sehingga ruangan masih terasa panas dan belum nyaman sesuai yang diinginkan. Namun lain halnya dengan suhu ruangan dengan menggunakan alternatif kontrol pertama dengan nilai $K_p = 100$. Nilai eigen pada matriks dinamik memiliki nilai real negatif dan 2 bilangan kompleks dengan bagian nilai real positif. Sehingga bila disimulasikan, maka suhu ruangan yang semula baik-baik saja menjadi kacau seperti pada gambar. Dari gambar tersebut dapat disimpulkan bahwa alternatif kontrol pertama belum cocok untuk digunakan pada ruangan agar ruangan dapat digunakan dengan nyaman.

Selanjutnya akan digunakan alternatif kontrol lainnya dengan memanfaatkan fungsi integral dari tinggi air raksa. Alternatif kontrol kedua ini menyisipkan alternatif kontrol pada persamaan (12) dan (13) dengan fungsi integral dari tinggi air raksa sehingga diperoleh alternatif kontrol berikut.

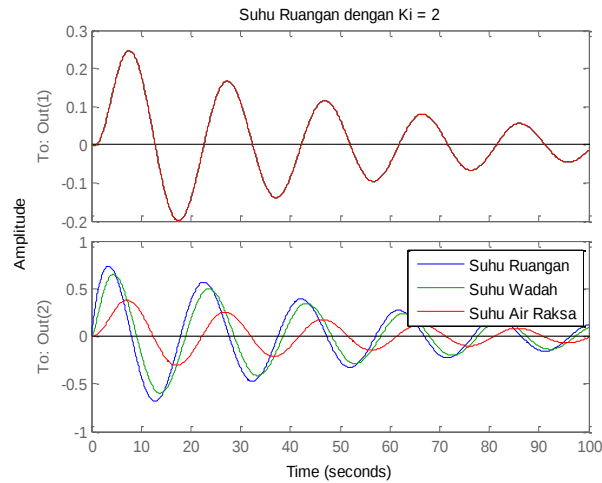
$$u = -K_p h - K_i z \quad (14)$$

Dengan $\frac{dz}{dt} = h$ atau z merupakan integral dari tinggi air raksa terhadap waktu. Serta K_i merupakan konstanta positif. Selanjutnya akan dicoba K_i dengan nilai yang berbeda-beda. Dengan nilai $K_p = 10$ diperoleh dari alternatif kontrol sebelumnya, $K_i = 1$, $K_i = 1,5$, $K_i = 2$. Didapat grafik suhu ruangan antara lain sebagai berikut.



Gambar 4. Skema suhu ruangan dengan $K_i = 1$, dan $K_i = 1,5$.

Dari gambar diatas, waktu yang dibutuhkan untuk mencapai kesetimbangan memang cukup lama, namun tidak seperti sebelumnya. Suhu ruangan dapat kembali menuju titik kesetimbangan dengan nilai $K_i = 1$ membutuhkan waktu 80 detik, dan $K_i = 1,5$ membutuhkan waktu 35 detik. Adapun respon yang terjadi dengan nilai $K_i = 2$ adalah sebagai berikut.



Gambar 4. Skema suhu ruangan dengan nilai $K_i = 2$

Dapat dilihat bahwa dalam jangka waktu 100 detik, suhu ruangan belum kembali menuju kesetimbangan. Sehingga dapat disimpulkan model pengendalian terbaik yaitu alternatif kontrol kedua dengan pengaturan pada katup u berbentuk.

$$u = -10h - 1,5z \quad (15)$$

Dengan pengaturan pada katup udara diatas, suhu ruangan dapat kembali menuju titik kesetimbangan dalam waktu 35 detik.

5. KESIMPULAN

Adapun simpulan yang didapat dari paper ini adalah.

1. Sistem yang bekerja pada pengendalian suhu ruangan merupakan sistem yang terkontrol dan terobservasi dengan baik.
2. Model kontrol terbaik pada katup udara yaitu dengan persamaan $u = -10h - 1,5z$, dengan h sebagai tinggi air raksa dan z merupakan fungsi integral dari fungsi raksa terhadap waktu.
3. Waktu yang diperlukan untuk mencapai titik kesetimbangan adalah 35 detik.

KEPUSTAKAAN

- [1] Fishwick, P.A. 2007. *Handbook of Dynamic System Modeling*. Chapman & Hall/CRC. New York
- [2] Meyers, R.A. 2011. *Mathematics of Complexity and Dynamical Systems*. Springer-Verlag, Inc., New York
- [3] Aoki, M. 2013. *State Space Modeling of Time Series*. Springer-Verlag, Inc., New York.
- [4] Gopal, M. 1993. *Modern Control Systems Theory*. New Age International, New Delhi.
- [5] Lagarrigue, F.L. and Loria, A. 2005. *Advanced Topics in Control Systems Theory*. Springer-Verlag Paris Inc., Paris.
- [6] Murdock, J.A. 1999. *Perturbations: Theory and Methods*. Siam, New York.

PENERAPAN LOGIKA FUZZY PADA CONTROL SUARA TV SEBAGAI ALTERNATIVE MENGHEMAT DAYA LISTRIK

Agus Wantoro

Sistem Informasi, Universitas Teknokrat Indonesia

Jl. H.ZA Pagaralam, No 9-11, Labuhanratu, Bandarlampung

Handphone : 085279078098, e-Mail : aguswantoro.ilkom@gmail.com

ABSTRAK

Di era sekarang ini, ilmu pengetahuan dan teknologi dikembangkan dan dimanfaatkan untuk membantu pekerjaan manusia agar lebih mudah. Penerapan teknologi banyak dilakukan diberbagai bidang, salahsatunya pada media elektronik. Televisi merupakan media elektronik yang banyak digunakan masyarakat sebagai media hiburan dirumah. Televisi juga sebagai media pandang sekaligus media pendengar (*audio-visual*)[3]. Banyak pengoprasian televisi salahsatunya pengaturan suara [2]. Saat ini, untuk memperbesar suara televisi menggunakan *remote control*. Hal ini dirasa kurang efektif karena akan sering diubah apabila penonton dalam jumlah banyak dan harus mengurangi suara televisi apabila kondisinya hening seperti saat tidur. Penggunaan suara televisi yang cukup besar ≥ 75 dB dapat mengkonsumsi daya listrik cukup besar yaitu 120 watt. Standar penggunaan suara pada ruangan/linkungan yang tenang sebesar 65-75 DbA. Pengurangan volume sebanyak 20% dapat mengurangi pemakaian daya listrik [7]. Berdasarkan hal tersebut, perlu teknologi pada televisi yang dapat menambah dan mengurangi suara televisi secara otomatis, untuk itu diperlukan logika fuzzy yang mengatur suara televisi dengan mengambil input dari suara disekitar. Satuan input suara yang digunakan yaitu DbA dan frekuensi yang diperoleh dari sensor suara yang selanjutnya akan diproses dan menghasilkan output berupa pengaturan volume televisi secara otomatis. Prototype teknologi ini diharapkan memberikan kemudahan penonton dalam mengatur volume televisi dan mengurangi konsumsi penggunaan listrik serta mengurangi konsumsi baterai remote

Katakunci : *Suara, Logika Fuzzy, Volume Televisi, Daya Listrik*

1. PENDAHULUAN

Di era globalisasi sekarang ini, semakin pesatnya perkembangan ilmu pengetahuan dan teknologi di dunia. Ilmu pengetahuan dan teknologi ini dimanfaatkan dan dikembangkan oleh manusia untuk dapat membantu pekerjaan mereka sehingga dapat menyelesaikan pekerjaan dengan lebih mudah dan efisien. TV merupakan peralatan listrik yang banyak digunakan masyarakat sebagai elektronik yang harus ada dirumah sebagai media hiburan. Televisi juga sebagai media pandang sekaligus media pendengar (*audio-visual*), dimana orang tidak hanya memandang

gambar, tetapi sekaligus mendengar atau mencerna narasi dari gambar tersebut [3]. Terdapat tiga macam karakteristik televisi, yaitu: (1) *Audiovisual* yakni dapat didengar sekaligus dilihat. (2) Berpikir dalam gambar, ada dua tahap yang dilakukan (a) visualisasi yaitu menerjemahkan kata-kata yang mengandung gagasan yang menjadi gambar secara individual. (b) penggambaran yakni kegiatan merangkai gambar-gambar individual sedemikian rupa sehingga kontinuitasnya mengandung makna tertentu (3) Pengoprasian lebih kompleks, dibandingkan dengan radio. Pengoprasian televisi siaran jauh lebih kompleks dan lebih banyak melibatkan orang [2]. Salah satu pengaturan pada televisi adalah suara / volume. Saat ini untuk memperbesar suara televisi masih menggunakan remote control. Hal ini dirasa kurang efektif karna akan sering dilakukan apabila penonton dalam jumlah banyak dan harus mengurangi suara televisi apabila jumlah penonton sedikit atau dalam keadaan hening seperti saat tidur. Penggunaan volume / suara televisi yang cukup besar atau ≥ 75 dB dapat mengkonsumsi daya listrik yang cukup besar yaitu 120 watt. Standar penggunaan suara pada ruangan dan lingkungan yang tenang sebesar 65-75 DbA. Pengurangan volume sebanyak 20% dapat mengurangi daya listrik [7]. Berdasarkan hal tersebut, perlu (*prototype*) teknologi pada televisi yang dapat menambah dan mengurangi suara televisi secara otomatis sesuai kebutuhan, untuk itu diperlukan logika salah satunya logika fuzzy yang mengatur suara televisi dengan mengambil input dari suara disekitar televisi. Logika Fuzzy merupakan satu komponen pembentuk *soft computing*. Suatu nilai dapat bernilai besar atau salah secara bersamaan [12]. Logika fuzzy banyak diterapkan pada beberapa peralatan listrik seperti vacuum cleaner, mesin cuci, AC, transmisi automatic pada kendaraan dan lain-lain [11]

2. LANDASAN TEORI

2.1. Televisi

Televisi adalah media pandang sekaligus media pendengar (*audio-visual*), yang dimana orang tidak hanya memandang gambar yang ditayangkan televisi, tetapi sekaligus mendengar atau mencerna narasi dari gambar tersebut [3]. Terdapat tiga macam karakteristik televisi, yaitu: (1) *Audiovisual* yakni dapat didengar sekaligus dilihat. (2) Berpikir dalam gambar, ada dua tahap yang dilakukan (a) visualisasi yaitu menerjemahkan kata-kata yang mengandung gagasan yang menjadi gambar secara individual. (b) penggambaran yakni kegiatan merangkai gambar-gambar individual sedemikian rupa sehingga kontinuitasnya mengandung makna tertentu (3) Pengoprasian lebih kompleks, dibandingkan dengan radio. Pengoprasian televisi siaran jauh lebih kompleks, dan lebih banyak melibatkan orang [2]



Gambar 1 Televisi dan Remote

2.2. Suara Multimedia

Dalam multimedia, salah satu elemen yang ada didalamnya adalah audio atau suara. Pakar multimedia mendefinisikan suara adalah sesuatu yang disebabkan akibat dari perubahan tekanan udara yang menjangkau gendang telinga manusia [10]. Alex Firth dkk, telah menjelaskan dalam bukunya, *100 Things to Know About Science* Saat kita menonton TV, normalnya suara TV tersebut berukuran 60-70 dB.

2.3. Satuan Suara

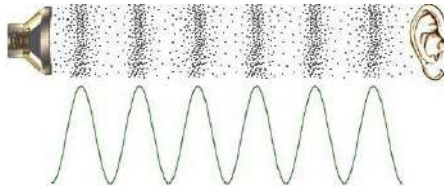
Suara diukur berdasarkan tiga hal yaitu periode, frekuensi dan amplitudo. Periode berarti lamanya suatu suara yang didengar. Berikut satuan suara yang dijadikan sebagai input :

a. Satuan Desibel (Db)

Amplitudo atau kekuatan suara diukur dalam satuan desibel (dB). *Desibel* adalah satuan yang digunakan untuk menyatakan kuantitas elektrik dari perubahan kuat-lemahnya amplitudo gelombang sinyal suara yang didengar oleh telinga manusia [5]. Percakapan biasa memiliki tingkat suara kira-kira 60 desibel. Para audiolog (ahli ilmu pendengaran) mengatakan bahwa semakin lama Anda mendengar suara berkekuatan di atas 85 desibel, semakin parah kerusakan pada indera pendengaran Anda. Majalah *Newsweek* mengatakan, "Telinga Anda sanggup bertahan mendengar suara mesin bor (100 dB) selama dua jam dengan aman, tetapi tidak akan tahan terhadap bisingnya areal permainan *video game* (110 dB) selama 30 menit. Volume audio televisi memiliki volume berkisar 0 -100 Db

b. Frekuensi (Hertz)

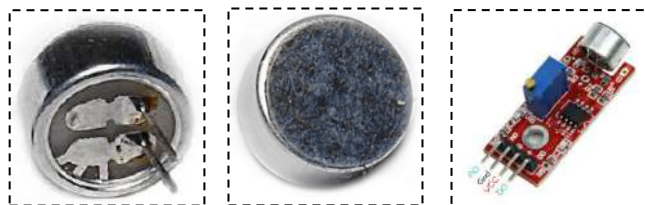
Frekuensi atau tinggi-rendah suara, dihitung dalam getaran per detik (hertz). Jangkauan frekuensi yang terdengar cukup jelas dan aman bagi pendengaran adalah antara 20 dan 20.000 getaran per detik. Tekanan udara berada pada rentang frekuensi 20 Hz sampai 20.000 Hz, maka telinga manusia akan mengidentifikasi tekanan udara sebagai suara. Pada manusia, telinga merupakan organ untuk pendengaran dan menjaga keseimbangan tubuh yang terdiri atas telinga luar, telinga tengah, dan telinga dalam [1]. Berdasarkan frekuensi dapat dibedakan menjadi 3 daerah: (1) Infrasonik yaitu 0 – (20 Hz), seperti getaran tanah, gempa bumi (2) Daerah Sonik, 20 – 20.000 Hz : yaitu daerah yang dapat didengar oleh manusia (audio frekuensi). (3): Daerah Ultrasonik >20.000 Hz akan membahayakan telinga manusia [9]



Gambar 2 Suara pada Satuan Desibel dan Frekuensi [Gabriel, 1996]

2.4. Sensor Suara

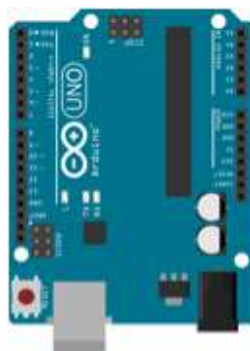
Sensor suara merupakan sensor yang mensensing besaran suara untuk diubah menjadi besaran listrik. Sensor ini bekerja berdasarkan besar kecilnya kekuatan gelombang suara yang diterima. Dimana gelombang suara tersebut mengenai membran sensor, yang menyebabkan bergeraknya membran sensor yang memiliki kumparan kecil sehingga menghasilkan besaran listrik. Kecepatan bergeraknya kumparan kecil tersebut menentukan kuat lemahnya gelombang listrik yang akan dihasilkan. Salah satu contoh komponen yang termasuk dalam sensor ini adalah condenser microphone atau mic [14]



Gambar 3 Condesor dan Modul Sensor

2.5. Arduino Uno

Arduino ini merupakan sebuah *board* mikrokontroler yang didasarkan pada ATmega328. *Board* ini dapat terhubung ke 14 sinyal digital I/O dan 6 sinyal analog *input* lalu *board* ini bersifat *open-source* [4]



Gambar 4 Board Mikrokontroler Arduino Uno

2.5. Logika Fuzzy

Logika fuzzy adalah suatu cara yang tepat untuk memetakan suatu ruang input dalam suatu ruang output dan memiliki nilai yang berlanjut. Kelebihan logika fuzzy ada pada kemampuan penalaran secara bahasa. Sehingga, dalam perancangannya tidak memerlukan persamaan matematis yang kompleks dari objek yang akan dikendalikan [12]. Logika fuzzy sering digunakan untuk mengekspresikan suatu nilai yang diterjemahkan dalam bahasa (*linguistic*), semisal untuk mengekspresikan suhu dalam ruangan apakah ruangan tersebut dingin, hangat, atau panas. Logika fuzzy banyak diterapkan pada beberapa peralatan listrik seperti vacuum cleaner, mesin cuci, AC, transmisi automatic pada kendaraan dan lain-lain [11]

a. Variabel Fuzzy

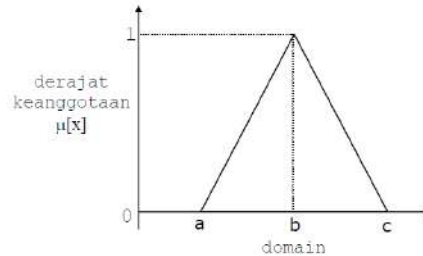
Variabel fuzzy merupakan kriteria yang hendak digunakan dalam perhitungan sistem fuzzy [13]. Pada penelitian ini variable input yang digunakan yaitu suara frekuensi dan desibel dan variable output adalah volume televisi

b. Himpunan Fuzzy

Pada himpunan tegas (crisp), nilai keanggotaan suatu item x dalam suatu himpunan A , yang sering ditulis dengan $\mu_A[x]$, memiliki 2 kemungkinan, yaitu Satu (1), yang berarti bahwa suatu item menjadi anggota dalam suatu himpunan. Nol (0), yang berarti bahwa suatu item tidak menjadi anggota dalam suatu himpunan [13]

c. Fungsi Keanggotaan Kurva Segitiga

Fungsi Keanggotaan (*membership function*) adalah suatu kurva yang menunjukkan pemetaan titik-titik input data ke dalam nilai keanggotaannya (sering juga disebut dengan derajat keanggotaan) yang memiliki interval antara 0 sampai 1. Salah satu cara yang dapat digunakan untuk mendapatkan nilai keanggotaan adalah dengan melalui pendekatan fungsi. Ada beberapa fungsi yang bisa digunakan [13]



Fungsi Keanggotaan:

$$\mu[x] = \begin{cases} 0; & x \leq a \text{ atau } x \geq c \\ (x - a)/(b - a); & a \leq x \leq b \\ (b - x)/(c - b); & b \leq x \leq c \end{cases}$$

Gambar 4. Kurva Segitiga

d. Semesta Pembicaraan

Semesta pembicaraan adalah keseluruhan nilai yang diperbolehkan untuk dioperasikan dalam suatu variabel fuzzy. Semesta pembicaraan merupakan himpunan bilangan real yang senantiasa naik (bertambah) secara monoton dari kiri ke kanan. Nilai semesta pembicaraan dapat berupa bilangan positif maupun negatif. Adakalanya nilai semesta pembicaraan ini tidak dibatasi batas atasnya. Contoh: Semesta pembicaraan untuk variabel temperatur: $[0, 40]$ [13]

e. Domain

Domain himpunan fuzzy adalah keseluruhan nilai yang diijinkan dalam semesta pembicaraan dan boleh dioperasikan dalam suatu himpunan fuzzy. Seperti halnya semesta pembicaraan, domain merupakan himpunan bilangan real yang senantiasa naik (bertambah) secara monoton dari kiri ke kanan. Nilai domain dapat berupa bilangan positif maupun negatif. Contoh domain himpunan fuzzy: MUDA = $[0, 45]$ [13]

f. Operator Dasar

Nilai keanggotaan sebagai hasil dari operasi 2 himpunan sering dikenal dengan nama *fire strength* atau α -predikat. Ada 3 operator dasar yaitu: Operator AND dan OR. Operator ini berhubungan dengan operasi interseksi pada himpunan. α -predikat sebagai hasil operasi dengan operator AND diperoleh dengan mengambil nilai keanggotaan terkecil antar elemen pada himpunan-himpunan yang bersangkutan. $\mu A \cap B = \min(\mu A[x], \mu B[y])$ [13]

g. Fungsi Implikasi

Tiap-tiap aturan (proposisi) pada basis pengetahuan fuzzy berhubungan dengan suatu relasi fuzzy. Bentuk fungsi implikasi: IF x is A THEN y is B dengan x dan y adalah skalar, dan A dan B adalah himpunan fuzzy. Proposisi yang mengikuti IF disebut sebagai anteseden, sedangkan proposisi yang mengikuti THEN disebut sebagai konsekuen. Proposisi ini dapat diperluas dengan menggunakan operator fuzzy, seperti: IF $(x_1 \text{ is } A_1) \cdot (x_2 \text{ is } A_2) \cdot (x_3 \text{ is } A_3) \cdot \dots \cdot (x_N \text{ is } A_N)$ THEN y is B [13]

2.6. Fuzzy Inferensi System (FIS) Mamdani

Metode Mamdani sering juga dikenal dengan nama Metode Max-Min. Pada Metode Mamdani, baik variabel input maupun variabel output dibagi menjadi satu atau lebih himpunan fuzzy. fungsi implikasi yang digunakan adalah Min [12]

2.7. Daya Listrik TV

Daya listrik diartikan sebagai besar energi listrik yang dihasilkan setiap detik. Pada setiap alat listrik selalu tercantum besarnya daya listrik tersebut. Satuan SI daya listrik adalah watt yang menyatakan banyaknya tenaga listrik yang mengalir per satuan waktu (joule/detik) [15]. Berikut table penggunaan daya listrik (Watt) pada televisi 32" dengan rata-rata penggunaan televisi 12 jam/hari berbagai merk:

Tabel 2. Daya Listrik pada Televisi

No	Merk	Ukuran	Watt	Biaya
1	LG LED	32"	160	Rp. 49.700
2	LG Plasma	32"	160	Rp. 49.700
3	Samsung LED	32"	100	Rp. 31.100
4	Samsung LCD	32"	50	Rp. 15.550
5	LG LCD	32"	100	Rp. 31.100
6	Samsung Plasma	32"	200	Rp. 62.200
7	Panasonic LCD	32"	95	Rp. 29.500
8	Panasonic Plasma	32"	230	Rp. 71.500

3. METODOLOGI PENELITIAN

3.1. Metode Penelitian

Penelitian ini menggunakan metode logika fuzzy dengan FIS Mamdani sebagai kontrol volume televisi secara otomatis. Data input penelitian ini diambil menggunakan sensor suara. Sensor ini akan mengambil suara

kebisingan dengan satuan decibel (Db) dan hertz (Hz) yang akan diolah menggunakan logika fuzzy menghasilkan volume televisi secara otomatis

3.2. Tahapan Penelitian

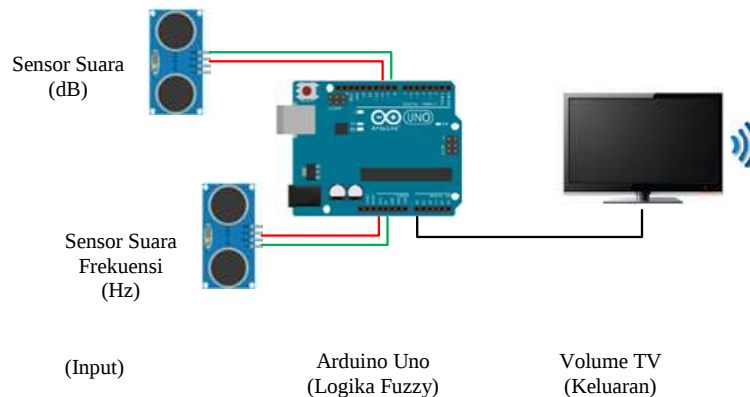
Tahapan penelitian yaitu penjelasan uraian atau tentang tahapan penelitian pemecahan masalah yang telah diidentifikasi atau dirumuskan. Tahapan penelitian dalam sebuah penelitian kuantitatif, sangat menentukan kejelasan pada proses penelitian secara keseluruhan. Tahapan penelitian dapat dilihat pada gambar 5 :



Gambar 5 Tahapan Penelitian

3.3. Prototype Control Volume TV

Tahap ini menggambarkan model prototype system yang akan dikembangkan. Sebagai data input menggunakan sensor suara dengan satuan decibel (Db) dan satuan suara frekuensi (Hz) yang akan diterima oleh Arduino sebagai control dengan logika fuzzy dan menjadi keluaran pada televisi

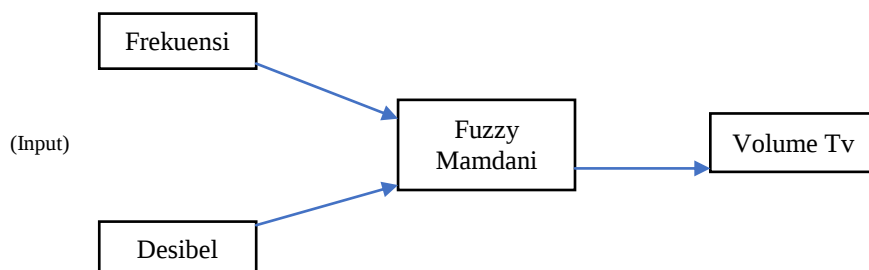


Gambar 6 Prototype

3.4. Penerapan Logika Fuzzy Mamdani

a. Variabel

Pada penelitian ini variable input yang digunakan yaitu suara yang diambil dari sensor suara dengan satuan Hz dan dB dengan variable keluaran berupa volume televisi



Gambar 7 Variabel Input

b. Domain

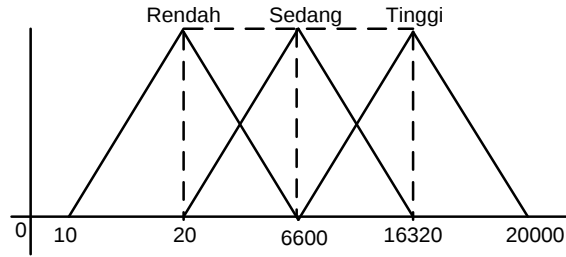
Penentuan domain diambil berdasarkan variable input. Variabel frekuensi (Hz) dengan himpunan rendah [10-6600], sedang [20 – 13320] dan tinggi [6600 – 20000]. Variabel decibel [Db] memiliki 3 himpunan yaitu, rendah dengan domain [0-60], sedang [30-85] dan tinggi [60-100]

Tabel 3 Domain Variabel

Variabel Input	Himpunan	Domain	Satuan
Frekuensi	Rendah	[10 – 6.600]	Hz
	Sedang	[20 – 13.320]	
	Tinggi	[6600 – 20.000]	
Desibel	Rendah	[0 – 60]	dB
	Sedang	[30 – 85]	
	Tinggi	[60 – 100]	
Volume TV	Rendah	[0 – 65]	dB
	Sedang	[35 – 90]	
	Tinggi	[65 – 100]	

c. Kurva dan Fungsi Keanggotaan

1. Kurva Frekuensi (Input)



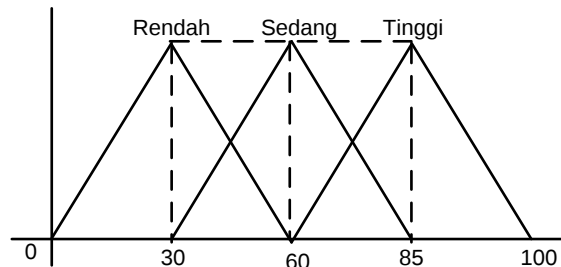
Gambar 8 Kurva Keanggotaan Frekuensi

$$Rendah[f] = \begin{cases} 0 & f \leq 10 \text{ atau } f \geq 6600 \\ \frac{f - 10}{20 - 10} & 10 \leq f \leq 20 \\ \frac{6600 - f}{6600 - 20} & 20 \leq f \leq 6600 \end{cases}$$

$$Sedang[f] = \begin{cases} 0 & f \leq 20 \text{ atau } f \geq 16320 \\ \frac{f - 20}{6600 - 20} & 20 \leq f \leq 6600 \\ \frac{16320 - f}{16320 - 6600} & 6600 \leq f \leq 16320 \end{cases}$$

$$Tinggi[f] = \begin{cases} 0 & f \leq 6600 \text{ atau } f \geq 20000 \\ \frac{f - 6600}{16320 - 6600} & 6600 \leq f \leq 16320 \\ \frac{20000 - f}{20000 - 16320} & 16320 \leq f \leq 20000 \end{cases}$$

2. Kurva Desibel (Input)



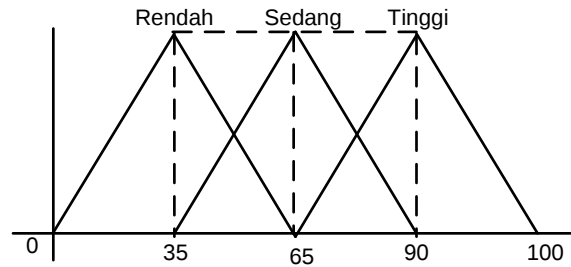
Gambar 9 Kurva Keanggotaan Desibel

$$Rendah[d] = \begin{cases} 0 & d \leq 0 \text{ atau } d \geq 60 \\ \frac{d-0}{30-0} & 0 \leq d \leq 30 \\ \frac{60-d}{60-30} & 30 \leq d \leq 60 \end{cases}$$

$$Sedang[d] = \begin{cases} 0 & d \leq 30 \text{ atau } d \geq 85 \\ \frac{d-30}{60-30} & 30 \leq d \leq 60 \\ \frac{85-d}{85-60} & 60 \leq d \leq 85 \end{cases}$$

$$Tinggi[d] = \begin{cases} 0 & d \leq 60 \text{ atau } d \geq 100 \\ \frac{d-60}{85-60} & 60 \leq d \leq 85 \\ \frac{100-d}{100-85} & 85 \leq d \leq 100 \end{cases}$$

3. Kurva Volume (Output)



Gambar 10 Kurva Keanggotaan Volume TV (Output)

$$Rendah[v] = \begin{cases} 0 & v \leq 0 \text{ atau } v \geq 65 \\ \frac{v-0}{35-0} & 0 \leq v \leq 35 \\ \frac{65-v}{65-35} & 35 \leq v \leq 65 \end{cases}$$

$$Sedang[v] = \begin{cases} 0 & v \leq 35 \text{ atau } v \geq 90 \\ \frac{v-35}{65-35} & 35 \leq v \leq 65 \\ \frac{90-v}{90-65} & 65 \leq v \leq 90 \end{cases}$$

$$Tinggi[v] = \begin{cases} 0 & v \leq 65 \text{ atau } v \geq 100 \\ \frac{v-65}{90-65} & 65 \leq v \leq 90 \\ \frac{100-v}{100-90} & 90 \leq v \leq 100 \end{cases}$$

d. Fungsi Implikasi(Rule Base)

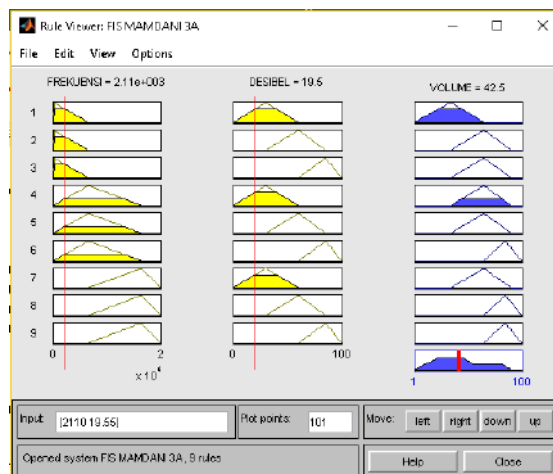
Pada control suara televisi menggunakan 2variabel input dengan 3 himpunan(rendah, sedang dan tinggi) dan 1 variabel *output* dengan 3 himpunan yaitu rendah, sedang dan tinggi, maka terdapat 9 aturan [R1-R9] aturan implikasi fuzzy yaitu:

- [R1] IF $\mu_{\text{Frekuensi}}=\text{Rendah}$ dan $\mu_{\text{Desibel}}=\text{Rendah}$ Then $\mu_{\text{Volume}}=\text{Rendah}$
- [R2] IF $\mu_{\text{Frekuensi}}=\text{Rendah}$ dan $\mu_{\text{Desibel}}=\text{Sedang}$ Then $\mu_{\text{Volume}}=\text{Sedang}$
- [R3] IF $\mu_{\text{Frekuensi}}=\text{Rendah}$ dan $\mu_{\text{Desibel}}=\text{Tinggi}$ Then $\mu_{\text{Volume}}=\text{Sedang}$
- [R4] IF $\mu_{\text{Frekuensi}}=\text{Sedang}$ dan $\mu_{\text{Desibel}}=\text{Rendah}$ Then $\mu_{\text{Volume}}=\text{Sedang}$
- [R5] IF $\mu_{\text{Frekuensi}}=\text{Sedang}$ dan $\mu_{\text{Desibel}}=\text{Sedang}$ Then $\mu_{\text{Volume}}=\text{Sedang}$
- [R6] IF $\mu_{\text{Frekuensi}}=\text{Sedang}$ dan $\mu_{\text{Desibel}}=\text{Tinggi}$ Then $\mu_{\text{Volume}}=\text{Tinggi}$
- [R7] IF $\mu_{\text{Frekuensi}}=\text{Tinggi}$ dan $\mu_{\text{Desibel}}=\text{Rendah}$ Then $\mu_{\text{Volume}}=\text{Sedang}$
- [R8] IF $\mu_{\text{Frekuensi}}=\text{Tinggi}$ dan $\mu_{\text{Desibel}}=\text{Sedang}$ Then $\mu_{\text{Volume}}=\text{Tinggi}$
- [R9] IF $\mu_{\text{Frekuensi}}=\text{Tinggi}$ dan $\mu_{\text{Desibel}}=\text{Tinggi}$ Then $\mu_{\text{Volume}}=\text{Tinggi}$

4. HASIL DAN PEMBAHASAN

4.1. Pengujian Kontrol Volume

Setelah pembuatan rule base pada implikasi, maka akan dilakukan pengujian menggunakan rule viewer matlab dengan FIS Mamdani menggunakan input suara pertama yaitu frekuensi yang dibuat dalam 3 himpunan (rendah, sedang dan tinggi). Input suara kedua yaitu Desibel (Db) dibuat dengan 3 himpunan (rendah, sedang dan tinggi) yang mengontrol volume tv

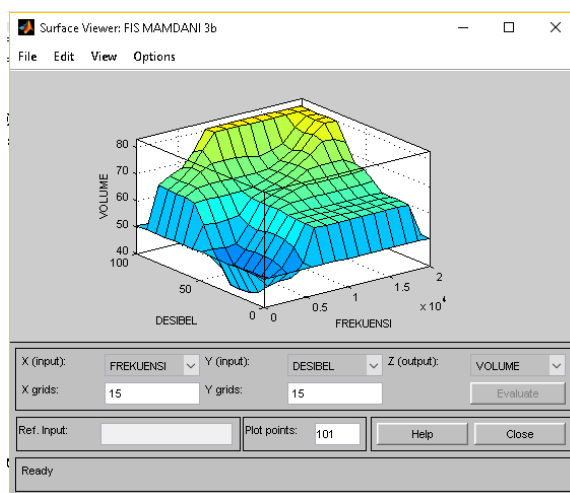


Gambar 11 Control Volume Berdasarkan Suara Frekuensi & Disibel

Berdasarkan hasil pengujian menggunakan rule viewer fuzzy, semakin rendah frekuensi dan decibel, maka akan semakin rendah volume Televisi, apabila frekuensi dan decibel tinggi, maka volume televisi akan naik atau membesar. Berikut hasil pengujian :

Tabel 4 Hasil Prediksi Pengujian Kontrol Volume TV

Frekuensi (Hz)	Desibel (Db)	Volume (Db)
826	25	37.5
2400	40	43
9200	55	66
13000	65	73
16000	80	83



Gambar 12. Grafik Perubahan Volume berdasarkan input frekuensi (Hz) dan decibel (Db)

4.2. Hipotesa

Berdasarkan hasil pengujian, maka dapat disimpulkan bahwa dengan menggunakan teknologi ini, maka volume televisi dapat diharapkan dapat menyesuaikan kebutuhan penonton dan memiliki kelebihan sebagai berikut :

- Volume televisi dapat mengatur sesuai kebutuhan, maka dapat menghemat daya listrik
- Pengurangan volume televisi sebesar 20% dapat menghemat daya listrik [7]
- Penggunaan remote control berkurang, maka dapat menghemat baterai remote
- Control volume televisi secara otomatis berdasarkan inputan sensor suara lingkungan, maka memudahkan pengaturan volume televisi
- Dapat menyesuaikan kebutuhan volume untuk penonton
- Volume televisse akan menyesuaikan kebutuhan, maka dapat mengurangi kebisingan akibat suara televisse yang tidak dkecilkan
- Kebisingan yang ditimbulkan televisi dapat berpengaruh terhadap dampak kesehatan [9]

4.3. Kontribusi

Hasil dari penelitian ini adalah berupa rancangan (*prototype*) dapat digunakan industry elektronik untuk menerapkan pada media elektronik sebagaipengatur volume suara secara otomatis agar lebih efektif dan efisien

5. SIMPULAN

Berdasarkan hasil dan pembahasan diatas, maka dapat disimpulkan sebagai berikut :

- a. Permasalahan pada penelitian ini diambil dari kondisi sehari-hari pada control volume televisi
- b. Penelitian ini menggunakan logika fuzzy dengan FIS Mamdani
- c. Input yang digunakan adalah data yang diambil dari sensor suara dengan satuan frekuensi (Hz) dan decibel (Db) yang diproses menggunakan logika fuzzy dan menghasilkan control berupa volume televisi
- d. Berdasarkan hasil pengujian menggunakan rule *viewer fuzzy*, semakin rendah frekuensi dan decibel, maka akan semakin rendah volume Televisi dan sebaliknya
- e. Penelitian ini hanya sebatas *prototype*, maka diharapkan dapat dilanjutkan ketahap pembuatan system dan penerapan

KEPUSTAKAAN

- [1] Andleigh, Prabhat K; Thakhar, Kiran. (1995). *Multimedia System Design*. Prentice Hall, Inc, New Jersey
- [2] Ardianto, Elvinaro. 2007. *Komunikasi Massa Suatu Pengantar*. Bandung : Simbosa Rekatama Media, 13 - 15
- [3] Badjuri, Adi. 2010. *Jurnalistik Televisi*. Yogyakarta: Graha Ilmu, 10-23
- [4] Banzi, M. “*Getting Started with Arduino*” O'Reilly. 2008
- [5] Bell A. 1996. *Noise : An Occupational Hazard and Public Nuisance*. WHO. Geneva. Switzerland
- [6] Darwinson, R. Wahyudi. Kontrol Kecepatan Robot Hexapod Pemadam Api Menggunakan Metoda Logika Fuzzy, Jurnal Nasional Teknik Elektro. 4 (2015), p. 227-234
- [7] Depkes RI. (2002). Keputusan Menkes RI No. 228/MENKES/SK/XI/2002 tentang *Pedoman Penyusunan Standar Pelayanan Minimal Perkantoran yang Wajib Dilaksanakan Daerah*
- [8] Frith, A., Lacey, M., Martin J, Jonathan “*100 Things to Know About Science*”, Newyork, 2015
- [9] Gabriel. 1996. *Fisika Kedokteran Buku Kedokteran*. Jakarta : EGC
- [10] G. Lu, Teknologi Multimedia, ArtechHouse Publishers, 1999. ISBN 0890063427
- [11] Hellendoorn, Hans, Palm & Rainer, 1994, Fuzzy System Technology at siemens R & D, Fuzzy and system 63 North Holland

- [12] Kusumadewi, Sri dan Purnomo Hari. 2010, “Aplikasi Logika Fuzzy”, Cetakan Pertama, Graham Ilmu, Yogyakarta
- [13] Lotfi A. Zadeh. *Fuzzy Set*. “Fuzzy Sets”. *Information and Control*, 8:338-353, 1965
- [14] Petruzella, Frank D. 2001. *Elektronik Industri*. Terjemahan sumanto. Edisi kedua. Yogyakarta: Andi
- [15] Thales, 1745, Listrik dan Magnet, Yunani

CLUSTERING WILAYAH LAMPUNG BERDASARKAN TINGKAT KESEJAHTERAAN

Henida Widyatama¹⁾

¹⁾BPS Kabupaten Pesawaran, 085658978008

E-mail: henida.widyatama@bps.go.id

ABSTRAK

Hingga saat ini, kesejahteraan masyarakat masih menjadi tujuan utama dari pembangunan di Indonesia, tidak terkecuali di Provinsi Lampung. Pengelompokan tingkat kesejahteraan wilayah di Lampung bertujuan untuk melihat wilayah mana yang menjadi prioritas perbaikan pembangunan ekonomi oleh Pemerintah. Keberhasilan pembangunan ekonomi dapat terlihat dari beberapa indikator. Indikator tersebut diantaranya yaitu Indeks Pembangunan Manusia, Pendapatan per Kapita, dan Persentase Penduduk Miskin. Metode analisis yang digunakan untuk pengelompokan ini adalah analisis cluster (Hierarchical Cluster). Pengelompokan kabupaten/kota di Lampung berdasarkan metode klaster Average Linkage. Dari hasil analisis cluster tersebut terbentuk empat kelompok dengan karakteristik yang berbeda. Kelompok pertama beranggotakan Kabupaten Lampung Barat, Tanggamus, Way Kanan, Pringsewu, dan Pesisir Barat. Kelompok kedua terdiri dari Lampung Selatan, Lampung Timur, Lampung Utara, dan Pesawaran. Kelompok ketiga terdiri dari Lampung Tengah, Tulang Bawang, Mesuji, Tulang Bawang Barat. Terakhir, kelompok keempat terdiri dari Kota Bandar Lampung dan Metro. Dengan demikian, pemerintah daerah dapat mengambil kebijakan yang sesuai dengan permasalahan di daerahnya.

Kata kunci: Average Linkage, Cluster, Kesejahteraan Rakyat

1. PENDAHULUAN

Salah satu indikator keberhasilan pembangunan suatu negara adalah laju pertumbuhan ekonomi yang dapat mencerminkan kemampuan pertambahan pendapatan nasional dari waktu ke waktu. Ada tiga komponen yang berpengaruh terhadap pertumbuhan ekonomi, yaitu akumulasi modal, pertumbuhan jumlah penduduk yang pada akhirnya menyebabkan pertumbuhan angkatan kerja, dan kemajuan teknologi. Modal manusia (*human capital*) merupakan salah satu komponen penting terkait sumber daya manusia. Untuk meningkatkan produktivitas, modal manusia merupakan terminologi yang mengacu kepada pengembangan kapasitas manusia di bidang pendidikan, kesehatan, dan pengembangan potensi lainnya. Pengembangan sumberdaya manusia dinilai menjadi penggerak kemajuan ekonomi suatu negara.

Tujuan pertama yang tercantum dalam *Sustainable Development Goals (SDGs)* adalah mengentaskan kemiskinan. Diangkatnya kemiskinan sebagai tujuan pertama bukan tanpa alasan. Peningkatan kesejahteraan yang diukur dari penurunan tingkat kemiskinan merupakan cerminan keberhasilan pembangunan yang diharapkan oleh setiap negara, termasuk Indonesia. Oleh sebab itu, pemerintah menargetkan tingkat kemiskinan turun 7-8 persen dalam Rencana Pembangunan Jangka Menengah Nasional (RPJMN) 2015-2019.

Berdasarkan Rencana Pembangunan Jangka Menengah Daerah (RPJMD) Provinsi Lampung, Program pembangunan Pemerintah Daerah Provinsi Lampung bertujuan untuk mengurangi tingkat kemiskinan dan pengangguran melalui peningkatan kualitas SDM, pengembangan teknologi, peningkatan pertumbuhan ekonomi, pemerataan pembangunan sesuai dengan kondisi, potensi dan permasalahannya. Prioritas pembangunan daerah didasarkan pada tiga dimensi pembangunan dan NAWACITA JOKOWI-JK 2015-2019 yang terdiri dari pembangunan manusia, pembangunan sektor unggulan, dan pemerataan pembangunan antar wilayah.

Secara statistik, tahun 2015, Lampung merupakan provinsi dengan tiga terbesar penduduk miskin yang mencapai angka 14,31 persen dari total penduduk. Selain itu, Lampung juga memiliki Indeks Pembangunan Manusia (IPM) terendah di Pulau Sumatera.

Tujuan dari penelitian ini adalah untuk mengelompokkan tingkat kesejahteraan di Lampung sehingga dapat terlihat wilayah mana yang menjadi prioritas perbaikan pembangunan ekonomi oleh Pemerintah Daerah.

2. LANDASAN TEORI

Pembangunan ekonomi dikatakan berhasil jika tingkat kesejahteraan masyarakat semakin baik. Tingkat kesejahteraan merupakan representasi yang bersifat kompleks karena multidimensi. Kesejahteraan hidup seseorang memiliki banyak indikator keberhasilan yang dapat diukur.

Kesejahteraan masyarakat menunjukkan ukuran hasil pembangunan masyarakat dalam mencapai kehidupan yang lebih baik yang meliputi: (i) peningkatan kemampuan dan pemerataan distribusi kebutuhan dasar seperti makanan, perumahan, kesehatan, dan perlindungan; (ii) peningkatan tingkat kehidupan, tingkat pendapatan, pendidikan yang lebih baik, dan peningkatan atensi terhadap budaya dan nilai-nilai kemanusiaan; (iii) memperluas skala ekonomi dan ketersediaan pilihan sosial dari individu dan bangsa.

Umumnya, pertumbuhan pendapatan per kapita dari waktu ke waktu membawa perubahan terhadap kesejahteraan masyarakat dengan arah yang sama. Pertimbangan menggunakan pendapatan per kapita sebagai indikator kesejahteraan masyarakat karena data tersebut mudah diperoleh.

Menurut [1], kemiskinan dipandang sebagai ketidakmampuan dari sisi ekonomi untuk memenuhi kebutuhan dasar makanan dan bukan makanan yang diukur dari sisi pengeluaran. Penduduk miskin adalah penduduk yang memiliki rata-rata pengeluaran perkapita perbulan di bawah garis kemiskinan. Pemahaman terhadap karakteristik penduduk miskin penting untuk dicermati agar paket kebijakan dan terobosan baru yang diciptakan terkait kemiskinan dapat tepat sasaran. Pengentasan kemiskinan dalam semua bentuk dan dimensi menjadi prasyarat untuk pembangunan berkelanjutan.

Pada tahun 1990, UNDP (*United Nations Development Programme*) memperkenalkan suatu ukuran yang dikenal dengan Indeks Pembangunan Manusia (IPM) dan diterbitkan secara berkala dalam laporan tahunan HDR (*Human Development Report*). IPM merupakan indikator penting untuk mengukur keberhasilan dalam upaya pembangunan kualitas hidup manusia/masyarakat/penduduk. Tujuan utama pembangunan adalah untuk menciptakan lingkungan yang memungkinkan masyarakat untuk menikmati umur panjang, sehat, dan menjalankan kehidupan yang produktif. Dengan demikian, semakin maju kualitas kesejahteraan manusia maka akan semakin banyak pilihan yang dimiliki.

Ada 18 indikator kesejahteraan masyarakat yang digunakan oleh *United Nations Research Institute for Social Development*. Apabila indikator tersebut digunakan maka perbedaan tingkat pembangunan antara negara maju dan negara sedang berkembang tidak terlalu besar. Kedelapan belas indikator tersebut antara lain: tingkat harapan hidup; konsumsi protein hewani per kapita; persentase anak-anak yang belajar di sekolah dasar dan menengah; persentase anak-anak yang belajar di sekolah kejuruan; jumlah surat kabar; jumlah telepon; jumlah radio; jumlah penduduk di kota-kota yang mempunyai 20.000 penduduk atau lebih; persentase laki-laki dewasa di sektor pertanian; persentase tenaga kerja yang bekerja di sektor listrik, gas, air, kesehatan, pengangkutan, pergudangan, dan transportasi; persentase tenaga kerja yang memperoleh gaji; persentase PDB yang berasal dari industri pengolahan; konsumsi energi per kapita; konsumsi listrik per kapita; konsumsi baja per kapita; nilai per kapita perdagangan luar negeri; produk pertanian rata-rata dari pekerja laki-laki di sektor pertanian; dan pendapatan per kapita Produk Nasional Bruto.

3. METODOLOGI PENELITIAN

Metode yang digunakan dalam penelitian ini adalah analisis *cluster* (analisis gerombol). Analisis kluster adalah teknik yang digunakan untuk mengklasifikasikan objek ke dalam kelompok yang relatif homogen yang disebut kluster. Objek dalam tiap cluster cenderung memiliki kemiripan satu dengan lainnya, sedangkan antar cluster mempunyai sifat yang berbeda. Analisis kluster juga disebut analisis klasifikasi atau taksonomi *numeric* (*numerical taxonomy*). Analisis kluster pada

prinsipnya digunakan untuk mereduksi data, yaitu meringkas sejumlah variabel menjadi lebih sedikit dan menamakannya sebagai kluster.

Analisis kluster dapat dibagi menjadi dua jenis, yaitu *Hierarchical Cluster* dan *K-Means cluster (nonhierarchical cluster)*. Pengelompokan secara hierarki biasanya digunakan untuk jumlah sampel yang relatif sedikit. Sedangkan untuk data yang banyak dapat digunakan *K-Means Cluster*.

Tujuan pengklasteran adalah untuk mengelompokkan objek yang mirip dalam kluster yang sama, maka diperlukan beberapa ukuran untuk mengakses seberapa mirip atau berbeda objek-objek tersebut. Pendekatan yang digunakan untuk mengukur kemiripan dinyatakan dalam jarak (distance) antara pasangan objek. Objek dengan jarak yang lebih pendek antara mereka akan lebih mirip satu sama lain dibandingkan dengan pasangan dengan jarak yang lebih panjang. Ada beberapa cara untuk mengukur jarak antara dua objek (kasus). Untuk mengelompokkan, digunakan pengukuran “kedekatan” atau “kemiripan”. Namun terkadang kita menggunakan subjektivitas dalam menentukan pengukuran kemiripan. Hal penting untuk dipertimbangkan adalah sifat dasar variabel (diskret, kontinyu, biner), skala pengukuran (nominal, ordinal, interval, rasio), dan pengetahuan *subject matter*.

Ukuran kemiripan yang dapat digunakan antara lain: *euclidean distance* atau nilai kuadrat jarak. Euclidean distance merupakan akar dari jumlah kuadrat perbedaan/deviasi di dalam nilai untuk setiap variabel. Selain itu, ada juga ukuran jarak lainnya, yaitu *the city-block* atau *manhattan distance*, yaitu jarak antara dua objek yang merupakan jumlah perbedaan mutlak/absolut di dalam nilai untuk setiap variabel. Sedangkan *chebyshev distance* merupakan jarak antara dua objek yang merupakan perbedaan mutlak/absolut yang maksimum di dalam nilai untuk setiap variabel.

Jarak Euclidean di antara p-dimensi observasi

$$\begin{aligned}x' &= [x_1, x_2, \dots, x_p] \text{ dan } y' = [y_1, y_2, \dots, y_n] \\d(x, y) &= \sqrt{(x_1 - y_1)^2 + (x_2 - y_2)^2 + \dots + (x_p - y_p)^2} \\&= \sqrt{(x - y)'(x - y)}\end{aligned}$$

Jarak statistik di antara dua observasi yang sama, yaitu

$$d(x, y) = \sqrt{(x - y)'A(x - y)}$$

Tata cara pengelompokan secara hierarki seperti struktur organisasi sehingga objek atau elemen yang berada di kluster tertentu dapat ditelusuri kenapa objek tersebut berada pada kluster yang bersangkutan. Sementara itu, penelusuran elemen dalam kluster tertentu tidak dapat dilakukan pada metode kluster nonhierarki. Pengelompokan secara hierarki biasanya digunakan ketika jumlah sampel relatif sedikit, sedangkan pengelompokan nonhierarki digunakan ketika jumlah sampel banyak. Selain itu, pada metode kluster hierarki, objek atau elemen yang sudah masuk dalam satu kluster, tidak dapat masuk lagi ke kluster yang lain, sedangkan pada metode kluster nonhierarki bisa karena tidak bersifat unik. Dengan pertimbangan-pertimbangan tersebut, maka metode pengelompokan yang digunakan dalam penelitian ini adalah metode pengelompokan secara hierarki.

Berdasarkan prosesnya, metode kluster hierarki dapat dikategorikan menjadi dua, yaitu *agglomerative hierarchical methods* (penggabungan/pemusatan) dan *divisive hierarchical methods* (pembagian/penyebaran). Dalam metode *agglomerative*, setiap obyek atau observasi dianggap sebagai sebuah kluster tersendiri. Kemudian pada tahap selanjutnya, dua cluster yang mempunyai kemiripan digabungkan menjadi sebuah klasster baru, demikian seterusnya. Sebaliknya, dalam metode *divisive*, dimulai dari sebuah kluster besar yang terdiri dari semua obyek atau observasi. Selanjutnya, obyek atau observasi yang paling tinggi nilai ketidakmiripannya dipisahkan, demikian seterusnya.

4. HASIL DAN PEMBAHASAN

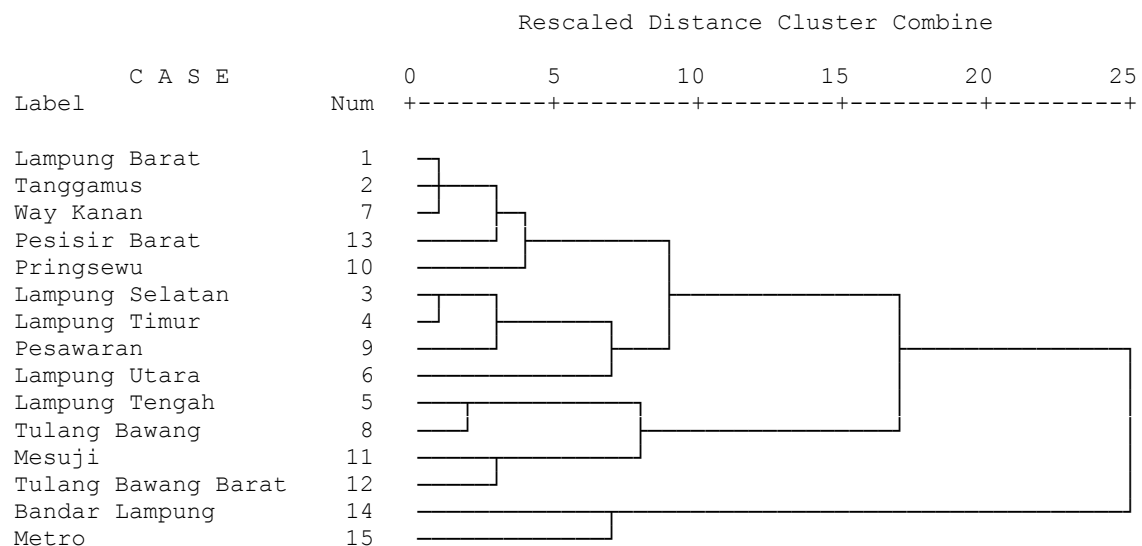
Jumlah penduduk Provinsi Lampung pada tahun 2016 sebanyak 8,2 juta jiwa. Sementara itu, jumlah penduduk miskin di Lampung ada sebanyak 1,2 juta jiwa atau 14,29 persen. Adapun pendapatan per kapita penduduk Provinsi Lampung sebesar 25,57 juta. Pada tahun yang sama, indeks pembangunan manusia Provinsi Lampung sebesar 67,65.

Karakteristik tingkat kesejahteraan di Provinsi Lampung dapat dilihat dari berbagai indikator, antara lain: IPM, pendapatan per kapita, dan persentase penduduk miskin. Berdasarkan hasil analisis kluster, 15 kabupaten/kota di Provinsi Lampung dikelompokkan berdasarkan karakteristik yang sama. Untuk mengetahui jarak dari ketidakmiripan antar wilayah kabupaten/kota berdasarkan karakteristik tingkat kesejahteraan yang digunakan dapat dilihat pada tabel berikut.

Tabel 1 Jarak antar Kluster

Case	1.Lampung Barat	2.Tanggamus	3.Lampung Selatan	4.Lampung Timur	5.Lampung Tengah	6.Lampung Utara	7.Way Kanan	8.Tulang Bawang	9.Pesawaran	10.Pringsewu	11.Mesuji	12.Tulang Bawang Barat	13.Pesisir Barat	14.Bandar Lampung	15.Metro
1.Lampung Barat															
2.Tanggamus	.000														
3.Lampung Selatan	.195	.000													
4.Lampung Timur	.360	.325	.000												
5.Lampung Tengah	.369	.347	.215	.000											
6.Lampung Utara	.946	.818	.184	.216	.000										
7.Way Kanan	.588	.632	.312	.257	.818	.000									
8.Tulang Bawang	.508	.307	.185	.182	.518	.490	.000								
9.Pesawaran	.678	.709	.362	.384	.812	.117	.523	.000							
10.Pringsewu	.234	.184	.837	.338	.495	.213	.191	.540	.000						
11.Mesuji	.121	.119	.372	.389	.831	.908	.942	.541	.415	.000					
12.Tulang Bawang Barat	.114	.911	.629	.827	.518	.182	.273	.225	.739	.828	.000				
13.Pesisir Barat	.588	.413	.412	.554	.420	.135	.385	.182	.563	.321	.135	.000			
14.Bandar Lampung	.109	.746	.317	.413	.918	.532	.118	.840	.167	.367	.838	.532	.000		
15.Metro	.160	.151	.811	.743	.336	.174	.169	.421	.135	.822	.121	.938	.161	.000	

Dari gambar dendrogram di bawah ini, dapat diputuskan berapa banyak kluster yang akan dibentuk. Dari gambar tersebut, sebaiknya dibuat empat kluster. Empat kluster tersebut sudah mampu menggambarkan perbedaan karakteristik pendidikan dari satu kluster dengan kluster yang lain.



Gambar 1 Hasil Output Dendrogram

Tabel 2 berikut ini memperjelas pengklusteran yang ditunjukkan oleh gambar dendrogram. Tabel 2 menunjukkan keanggotaan di setiap kluster yang terbentuk, yakni sebanyak empat kluster.

Tabel 2 Keanggotaan Tiap Klaster

No.	Kabupaten/Kota	Klaster
(1)	(2)	(3)
1	Lampung Barat	1
2	Tanggamus	1
3	Lampung Selatan	2
4	Lampung Timur	2
5	Lampung Tengah	3
6	Lampung Utara	2
7	Way Kanan	1
8	Tulang Bawang	3
9	Pesawaran	2
10	Pringsewu	1
11	Mesuji	3
12	Tulang Bawang Barat	3
13	Pesisir Barat	1
14	Bandar Lampung	4
15	Metro	4

Dari pembentukan empat klaster di atas, diperoleh masing-masing anggota klaster sebagai berikut.

- Klaster pertama terdiri dari Kabupaten Lampung Barat, Tanggamus, Way Kanan, Pringsewu, dan Pesisir Barat.
- Klaster kedua beranggotakan Kabupaten Lampung Selatan, Lampung Timur, Lampung Utara, dan Pesawaran.
- Klaster ketiga, yaitu Kabupaten Lampung Tengah, Tulang Bawang, Mesuji, dan Tulang Bawang Barat.
- Klaster keempat, terdiri dari Kota Bandar Lampung dan Metro.

Berdasarkan hasil pengelompokan tersebut, maka pencirian klaster dari masing-masing klaster tersebut dapat dilihat pada Tabel 3.

Tabel 3 Pencirian Klaster

Klaster		IPM	Pendapatan per Kapita (000 Rupiah)	Persentase Penduduk Miskin (Persen)
(1)		(2)	(3)	(4)
1	Lampung Barat	65,07 (cukup rendah)	16.827 (rendah)	14,27 (cukup tinggi)
	Tanggamus			
	Way Kanan			
	Pringsewu			
	Pesisir Barat			
2	Lampung Selatan	65,87 (cukup tinggi)	24.543 (cukup rendah)	18,34 (tinggi)
	Lampung Timur			
	Lampung Utara			
	Pesawaran			
3	Lampung Tengah	64,89 (rendah)	29.853 (tinggi)	9,97 (rendah)
	Tulang Bawang			
	Mesuji			
	Tulang Bawang Barat			
4	Kota Bandar Lampung	75,40 (tinggi)	27.846 (cukup tinggi)	10,15 (cukup rendah)
	Kota Metro			

Pencirian klaster dilakukan dengan membandingkan nilai antar klaster dalam suatu variabel. Perbandingan tersebut menggunakan empat skala, yaitu tinggi, cukup tinggi, cukup rendah, dan rendah. Adapun pencirian dari empat klaster tersebut adalah sebagai berikut.

- Klaster pertama terdiri dari Kabupaten Lampung Barat, Tanggamus, Way Kanan, Pringsewu, dan Pesisir Barat. Klaster ini dicirikan dengan pendapatan per kapita yang rendah dan persentase penduduk miskin yang masih cukup tinggi. Selain itu, IPM pada klaster ini juga masih cukup rendah. Pendapatan per kapita yang rendah dapat meningkatkan persentase penduduk miskin. Hal ini disebabkan daya beli masyarakat rendah.
- Klaster kedua terdiri dari Kabupaten Lampung Selatan, Lampung Timur, Lampung Utara, dan Pesawaran. Klaster ini dicirikan dengan persentase penduduk miskin yang tinggi dan pendapatan per kapita yang cukup rendah. Namun pada klaster kedua ini, IPM nya sudah cukup tinggi.
- Klaster ketiga terdiri dari Kabupaten Lampung Tengah, Tulang Bawang, Mesuji, dan Tulang Bawang Barat. Klaster ini dicirikan dengan nilai IPM yang rendah. Namun, pendapatan per kapita dalam klaster ini sudah tinggi dan persentase penduduk miskinnya pun rendah. Pendapatan per kapita yang tinggi mampu menurunkan persentase penduduk miskin karena daya beli masyarakat tinggi.
- Klaster keempat terdiri dari Kota Bandar Lampung dan Metro. Dari keempat klaster yang ada, klaster ini menunjukkan tingkat kesejahteraan yang paling baik. Hal ini dicirikan dengan angka IPM yang tinggi, pendapatan per kapita yang cukup tinggi, dan persentase penduduk miskin yang cukup rendah. Penduduk kota dianggap lebih baik dari sisi SDM nya. Hal ini wajar saja karena fasilitas pendidikan di kota lebih banyak dan lengkap serta akses informasi di kota lebih mudah. Sehingga kualitas sumber daya manusianya pun semakin baik.

5. SIMPULAN

Dari hasil pembahasan dapat diketahui bahwa tingkat kesejahteraan di Lampung dapat diklasterkan/dikelompokkan ke dalam 4 klaster. Klaster pertama bercirikan pendapatan per kapita yang rendah. Klaster kedua bercirikan persentase penduduk miskin yang tinggi. Klaster ketiga bercirikan IPM yang rendah, dan klaster keempat bercirikan persentase penduduk miskin yang cukup rendah.

Dari kesimpulan tersebut, saran yang dapat diberikan adalah pemerintah hendaknya memperhatikan kelemahan-kelemahan pada setiap klaster tersebut untuk memperbaiki perekonomian maupun kesejahteraan masyarakat.

KEPUSTAKAAN

- [1] BPS. (2016). Indikator Kesejahteraan Rakyat 2016. BPS, Jakarta.

PEMANFAATAN SISTEM INFORMASI GEOGRAFIS UNTUK VALUASI JASA LINGKUNGAN MANGROVE DALAM PENYAKIT MALARIA DI PROVINSI LAMPUNG

Imawan Abdul Qohar¹⁾, Samsul Bakri²⁾, & Dyah W.S.R Wardani³⁾

¹⁾ Mahasiswa PS Kehutanan Fakultas Pertanian, ²⁾ Dosen PS Kehutanan Fakultas Pertanian,

³⁾ Dosen Fakultas Kedokteran, Universitas Lampung

Jl. Soemantri Brojonegoro No.1 Bandar Lampung

Email : imwnimwn@gmail.com Mobile: 08127214600

ABSTRAK

Perubahan tutupan ekosistem mangrove berdampak terhadap angka kesakitan malaria (annual parasite incidence). Tujuan dari penelitian ini yaitu mengetahui besaran manfaat hutan mangrove dengan pendekatan human capital atau biaya perawatan (Medical cost) dari penyakit malaria. Penelitian ini dilakukan mulai bulan Agustus—Oktober 2016 dengan lingkup wilayah penelitian Provinsi Lampung tahun 2000—2015. Dinamika perubahan tutupan lahan dan penggunaan lahan per Kabupaten/kota diidentifikasi melalui sistem informasi geografis serta interpretasi citra landsat 2000, 2009, dan 2015 dan menghasilkan persentase luas tutupan lahan dan penggunaan lahan. Dari hasil uji statistik diketahui ekosistem mangrove berpengaruh terhadap penurunan kejadian malaria dengan koefisien sebesar -0,07937 dan nilai P value = 0,001. Luas ekosistem mangrove sebesar 9401,62 Ha, apabila terjadi kenaikan 10 persen luas tutupan mangrove maka akan menurunkan kejadian malaria sebesar 0,007937 per 1000 penduduk atau 0,000007937 insidensi. Hal ini berarti setiap penambahan 10 persen tutupan mangrove akan mengakibatkan penurunan kejadian penyakit malaria sebesar 64,42676 insidensi dengan asumsi bahwa setiap faktor penularan malaria adalah positif. Model valuasi jasa lingkungan mangrove dilakukan dengan pendekatan biaya kesehatan. Nilai manfaat dari hutan mangrove dengan pendekatan human capital adalah Rp. 2.266.255.815,5,-/tahun.

Kata kunci: Malaria, Mangrove, Medical cost, Penggunaan lahan

1. PENDAHULUAN

Luas mangrove di Indonesia adalah sekitar 4,25 juta hektar, yang mempresentasikan 25 % dari mangrove dunia. Indonesia merupakan pusat dari sebagian biogeografi genus mangrove. Hutan mangrove diketahui memiliki manfaat ganda (*multiple use*) yang dapat dibedakan atas manfaat langsung dan manfaat tidak langsung. Manfaat langsung merupakan manfaat yang dapat dirasakan secara langsung kegunaannya dan nilainya dapat diukur secara kuantitatif dalam pemenuhan kebutuhan manusia dari hasil hutan berupa barang dan jasa. Manfaat tidak langsung yaitu manfaat yang nyata namun sulit dirasakan dan diukur secara kuantitatif nilainya[1]. Beberapa fungsi ekologi yang dimiliki hutan mangrove adalah sebagai daerah asuhan (*nursery ground*), daerah untuk mencari makan (*feeding ground*) dan daerah pemijahan (*spawning ground*) bagi berbagai biota laut, tempat bersarangnya burung, habitat alami bagi berbagai jenis biota, sumber plasma nutfah (hewan, tumbuhan dan mikroorganisme) dan pengontrol penyakit seperti malaria[2].

Luas kawasan mangrove di Indonesia yang bervegetasi adalah sekitar 3.244.018,46 ha. Akan tetapi luas hutan mangrove tersebut telah banyak mengalami penurunan kualitas dan kuantitas yang dikarenakan kegiatan konversi (tambak, pemukiman, persawahan), penebangan kayu yang tidak bertanggung jawab (kayu bakar, pembuatan arang), pencemaran dan lainnya[3]. Kecenderungan konversi hutan mangrove menjadi bentuk penggunaan lahan lain semakin meningkat, yang didasari semata-mata kepentingan ekonomi dan kurang memperhatikan keberlanjutan kepentingan ekologi dan sosial [4].

Hutan mangrove sebagai habitat nyamuk dapat memengaruhi kehidupan larva nyamuk karena kanopi tegakan mangrove dapat menghalangi sinar matahari yang masuk atau melindungi dari serangan makhluk hidup lain, sehingga larva tersebut dapat berkembang biak dengan baik di dalam hutan mangrove tersebut[5]. Salah satu fungsi ekologi hutan mangrove adalah sebagai habitat berbagai nyamuk termasuk nyamuk penyebab penyakit malaria (*Anopheles sp.*). Wabah penyakit malaria bisa meningkat akibat terdegradasinya hutan mangrove. Kondisi hutan mangrove yang buruk menstimulasi nyamuk *Anopheles sp.* untuk bermigrasi ke habitat lain seperti pemukiman, yang selanjutnya menjadi vektor penyakit malaria[6]. Fluktuatif insidensi malaria disamping disebabkan perubahan cuaca juga dikarenakan adanya perubahan lingkungan diantaranya tambak-tambak udang yang terlantar, pembukaan hutan, perkebunan, dan penebangan hutan bakau. Minimnya penilitan mengenai pengkajian antara peranan jasa ekosistem mangrove terhadap pengendalian penyakit malaria melatarbelakangi diadakannya penelitian ini[7].

Tujuan Penelitian

Tujuan dari penelitian ini yaitu mengetahui besaran manfaat hutan mangrove dengan pendekatan human capital / *Medical Cost* dari penyakit malaria.

Kerangka Pemikiran

Dinamika luasan tutupan atau ekosistem mangrove begitu pesat. Pada wilayah-wilayah pinggiran urban umumnya cepat berkurang atau terkonversi menjadi kawasan dengan intensitas pembangunan yang relatif cepat seperti tambak, pemukiman bahkan industri. Padahal fungsi ekosistem mangrove memiliki jasa lingkungan yang sangat besar termasuk sebagai pengendali berbagai jenis penyakit termasuk malaria. Deforestasi kawasan mangrove dapat mengubah kondisi biofisik utamanya penurunan ekosistem mangrove, terbentuknya kawasan lain, tambak atau penggunaan lain. Lebih lanjut konversi ini dapat memicu pertumbuhan dan kepadatan jumlah penduduk. Kecuali itu kedua jenis perubahan tersebut dapat menyebabkan degradasi kondisi ekosistem, dengan kata lain terjadi *ecological shock* (guncangan ekologis) seperti perubahan iklim, suhu, kelembaban udara maupun kondisi hidrologis seperti terbentuknya rawa-rawa, intrusi air laut, tingkat salinitas, dampak tersebut dapat merangsang berkembangnya habitat nyamuk *Anopheles* maupun dominasinya. Perubahan iklim mikro di lain pihak dapat menurunkan tingkat kenyamanan lingkungan hidup bagi manusia yang berarti juga pada gejala fisiologisnya. Gejala fisiologis ini dapat mempengaruhi ketahanan masyarakat terhadap malaria.

Argumentasi tentang hubungan antara perubahan tutupan mangrove sampai pada penurunan ketahanan masyarakat terhadap malaria tersebut dapat membawa imajinasi untuk ditetapkan parameter biofisik maupun sosial demografinya. Artinya perlu mengembangkan model hubungan kausalitas sederhana. Saat ini belum diketahui korbanan jasa lingkungan tersebut atas konversi mangrove menjadi areal penggunaan lainnya. Apabila penelitian ini dapat menghasilkan informasi besarnya korbanan tersebut maka hasil tersebut dapat digunakan sebagai dasar perhitungan kompensasi atas perubahan luas ekosistem mangrove. Selanjutnya dapat dijadikan landasan bagi penentu kebijakan publik untuk mengendalikan deforestasi ekosistem mangrove maupun pengendalian penyakit malaria.

2. LANDASAN TEORI

Ekosistem Hutan Mangrove

Vegetasi mangrove secara khas memperlihatkan adanya pola zonasi. Beberapa ahli menyatakan bahwa hal tersebut berkaitan erat dengan tipe tanah (lumpur, pasir, atau gambut), keterbukaan (terhadap hempasan gelombang), salinitas serta pengaruh pasang surut. Di Indonesia, substrat berlumpur ini sangat baik untuk tegakan *Rhizophora mucronata* dan *Avicennia marina*. Jenis-jenis lain seperti *Rhizophora stylosa* tumbuh dengan baik pada substrat berpasir, bahkan pada pulau karang yang memiliki substrat berupa pecahan karang, kerang, dan bagian-bagian dari *Halimeda*[8].

Secara sederhana, mangrove umumnya tumbuh dalam 4 zona, yaitu pada daerah terbuka, daerah tengah, daerah yang memiliki sungai berair payau sampai hampir tawar, serta daerah ke arah daratan yang memiliki air tawar. Mangrove terbuka merupakan mangrove yang berada pada bagian yang berhadapan dengan laut. Komposisi floristik dari komunitas di zona terbuka sangat bergantung pada substratnya. *Sonneratia alba* cenderung untuk mendominasi daerah berpasir, sementara *Avicennia marina* dan *Rhizophora mucronata* cenderung untuk mendominasi daerah yang lebih berlumpur [9]. Mangrove tengah merupakan mangrove di zona yang terletak di belakang mangrove terbuka. Di zona ini, biasanya didominasi oleh jenis *Rhizophora*. Mangrove payau merupakan mangrove yang berada di sepanjang sungai berair payau hingga hampir tawar. Di zona ini biasanya didominasi oleh komunitas *Nypa* atau *Sonneratia*. Sedangkan mangrove daratan merupakan mangrove yang berada di zona perairan payau atau hampir tawar di belakang jalur hijau mangrove yang sebenarnya. Jenis-jenis yang umum ditemukan pada zona ini termasuk *Ficus microcarpus*, *Intsia bijuga*, *Nypa fruticans*, *Lumnitera racemosa*, *Pandanus sp.* dan *Xylocarpus moluccensis* [10]. Zona ini memiliki kekayaan jenis yang lebih tinggi dibandingkan dengan zona lainnya.

Peranan Ekosistem Mangrove Terhadap Endemik Malaria

Hutan mangrove memiliki berbagai macam fungsi yaitu:

1. Fungsi fisik; menjaga garis pantai agar tetap stabil, melindungi pantai dari erosi (abrasi) dan intrusi air laut, peredam gelombang dan badai, penahan lumpur, penangkap sedimen, pengendali banjir, mengolah bahan limbah, penghasil detritus, memelihara kualitas air, penyerap CO₂ dan penghasil O₂ serta mengurangi resiko terhadap bahaya tsunami.
2. Fungsi biologis; merupakan daerah asuhan (*nursery ground*), daerah untuk mencari makan (*feeding ground*) dan daerah pemijahan (*spawning ground*) dari berbagai biota laut, tempat bersarangnya burung, habitat alami bagi berbagai jenis biota, sumber plasma nutfah (hewan, tumbuhan dan mikroorganisme) serta pengontrol penyakit malaria.
3. Fungsi sosial ekonomi; sumber mata pencarian, produksi berbagai hasil hutan (kayu, arang, obat dan makanan), sumber bahan bangunan, bahan kerajinan, tempat wisata alam, objek pendidikan dan penelitian, areal pertambakan, tempat pembuatan garam serta areal perkebunan.

Secara ekologis hutan mangrove memegang peranan kunci dalam perputaran nutrisi pada perairan pantai di sekitarnya. Fungsi hutan mangrove yaitu sebagai stabilisator tepian sungai/pesisir, memberikan dinamika pertumbuhan di kawasan pesisir seperti pengendalian erosi pantai, menjaga stabilitas sedimen, dan turut berperan dalam menambah perluasan lahan daratan (*land building*) [11].

Manfaat lain dari fungsi ekologisnya adalah sebagai habitat nyamuk, sehingga kerusakan hutan mangrove dapat berakibat pada peningkatan populasi nyamuk sebagai vektor penyakit malaria [12]. Hutan mangrove sebagai habitat nyamuk dapat mempengaruhi kehidupan larva nyamuk karena kanopi tegakan mangrove dapat menghalangi sinar matahari yang masuk atau melindungi dari serangan makhluk hidup lain, sehingga larva tersebut dapat berkembang biak dengan baik di dalam hutan mangrove tersebut [13]. Munculah asumsi bahwa dengan adanya hutan mangrove sebagai habitat nyamuk maka daerah jelajah nyamuk khususnya *Anopheles sp.* hanya di dalam dan sekitar hutan mangrove itu saja, sehingga kawasan penduduk akan aman dari serangan nyamuk tersebut pada radius jarak tertentu. Berbeda jika kualitas dan kuantitas hutan mangrove tersebut buruk, seperti terjadinya pembukaan areal hutan mangrove yang dapat menimbulkan masalah kesehatan [6].

Sistem Informasi Geografis (SIG)

Sistem Informasi Geografis (SIG) adalah sistem informasi yang didasarkan pada kerja komputer yang memasukkan, mengelola, memanipulasi dan menganalisa data serta memberi uraian [14]. SIG menurut Burrough (1986) merupakan alat yang bermanfaat untuk pengumpulan, penyimpanan, pengambilan kembali data yang diinginkan dan penayangan data keruangan yang berasal dari kenyataan dunia [15].

Menurut [16], subsistem-subsistem dari SIG adalah sebagai berikut:

1. Data *input* Subsistem ini bertugas untuk mengumpulkan dan mempersiapkan data spasial dan data atribut dari berbagai sumber. Subsistem ini pula yang bertanggungjawab dalam mengkonversi atau mentransformasi format-format data aslinya ke dalam format yang dapat digunakan SIG.
2. Data *output* Subsistem ini menampilkan atau menghasilkan keluaran seluruh atau sebagian basis data baik dalam bentuk *softcopy* maupun *hardcopy*.
3. Data manajemen Subsistem ini mengorganisasi data, baik data spasial maupun data atribut ke dalam sebuah data sedemikian rupa sehingga mudah untuk digunakan, diperbaharui, dan diolah.
4. Data *manipulation* dan *analysis* Subsistem ini menentukan informasi-informasi yang dapat dihasilkan oleh SIG.

Citra Landsat

Dari sekian banyak satelit penginderaan jauh yang sering digunakan untuk pemetaan penutupan lahan adalah Landsat (Land Satelit). Seri Landsat yang dikenal pertama kali adalah *Earth Resource Technology Satelit* (ERTS). Citra landsat merupakan satelit sumberdaya milik Amerika Serikat yang diluncurkan sejak tahun 1972. Jenis citra yang direkam landsat hingga saat ini adalah Landsat MSS dan Landsat TM/ETM+/OLI. Jenis citra Landsat yang sudah mengorbit saat ini adalah Landsat generasi ke Delapan (Landsat 8). *Landsat Data Continuity*

Mission atau yang lebih dikenal Landsat 8 menggunakan sensor OLI (*Onboard Operational Land Image*) dan TIRS (*Thermal Infrared Sensor*) yang diluncurkan pada 11 Februari 2013 yang pada setiap saluran/kanal (*band*) mempunyai karakteristik dan kemampuan aplikasi atau penggunaan yang berbeda.

Penyakit Malaria

Penyakit malaria adalah penyakit menular yang menyerang dalam bentuk infeksi akut ataupun kronis. Penyakit ini disebabkan oleh protozoa genus *plasmodium* bentuk aseksual, yang masuk ke dalam tubuh manusia dan ditularkan oleh nyamuk *Anopheles* betina. Istilah malaria diambil dari dua kata bahasa italia yaitu *mal* = buruk dan *area* = udara atau udara buruk karena dahulu banyak terdapat di daerah rawa – rawa yang mengeluarkan bau busuk. Penyakit ini juga mempunyai nama lain seperti demam roma, demam rawa, demam tropik, demam pantai, demam charges, demam kura dan paludisme. Di dunia ini hidup sekitar 400 spesies nyamuk *anopheles*, tetapi hanya 60 spesies berperan sebagai vektor malaria alami. Di Indonesia, ditemukan 80 spesies nyamuk *Anopheles* tetapi hanya 16 spesies sebagai vektor malaria.[17]

Malaria disebabkan oleh protozoa darah yang termasuk ke dalam genus *Plasmodium*. *Plasmodium* ini merupakan protozoa obligat intraseluler. Pada manusia terdapat 4 spesies yaitu *Plasmodium falciparum*, *Plasmodium vivax*, *Plasmodium malariae* dan *Plasmodium ovale*. Penularan pada manusia dilakukan oleh nyamuk betina *Anopheles* ataupun ditularkan langsung melalui transfuse darah atau jarum suntik yang tercemar serta dari ibu hamil kepada janinnya. *Malaria vivax* disebabkan oleh *P. vivax* yang juga disebut juga sebagai malaria tertiana. *P. malariae* merupakan penyebab malaria malariae atau malaria kuartana. *P. ovale* merupakan penyebab malaria ovale, sedangkan *P. falciparum* menyebabkan malaria falsiparum atau malaria tropika [18].

Hubungan Host, Agent dan Faktor Lingkungan

1. Host

a. Manusia (*Host Intermediate*)

Pada dasarnya setiap orang dapat terkena malaria, tetapi kekebalan yang ada pada manusia merupakan perlindungan terhadap infeksi *Plasmodium* malaria.

b. Nyamuk *Anopheles* spp (*Host Defenitive*)

Nyamuk *Anopheles* spp sebagai penular penyakit malaria yang menghisap darah hanya nyamuk betina yang diperlukan untuk pertumbuhan dan mematangkan telurnya. Jenis nyamuk *Anopheles* spp di Indonesia lebih dari 90 macam. Dari jenis yang ada hanya beberapa jenis yang mempunyai potensi untuk menularkan malaria (Vektor) yaitu berjumlah 18 spesies. Di Indonesia dijumpai beberapa jenis *Anopheles* spp sebagai vector Malaria, antara lain :*An. sudaicus* sp, *An. Maculates* sp, *An. Balabacensis* sp, *An. Barbnirostrip* sp [19]. Di setiap daerah dimana terjadi transmisi malaria biasanya hanya ada 1 atau paling banyak 3 spesies *Anopheles* yang menjadi vektor penting. Vector-vektor tersebut memiliki habitat mulai dari rawa-rawa, pegunungan, sawah, pantai dan lain-lain [13].

2. Agent

Agent atau penyebab penyakit adalah semua unsur atau elemen hidup ataupun tidak hidup dimana kehadirannya, bila diikuti dengan kontak efektif dengan manusia yang rentan akan terjadi stimulasi untuk memudahkan terjadi suatu proses penyakit. *Agent* penyebab penyakit malaria termasuk agent biologis yaitu protozoa. Sampai saat ini dikenal empat macam agent penyebab malaria yaitu: *Plasmodium Falciparum*, penyebab malaria tropika yang sering menyebabkan malaria berat/malaria otak yang fatal, gejala serangnya timbul berselang setiap dua hari (48 jam) sekali. *Plasmodium vivax*, penyebab penyakit malaria tertiana yang gejala serangnya timbul berselang setiap tiga hari (Sering Kambuh). *Plasmodium malariae*, penyebab penyakit malaria quartana yang gejala serangnya timbul berselang setiap empat hari sekali. Dan *Plasmodium ovale*, jenis ini jarang sekali dijumpai, umumnya banyak di Afrika dan Pasifik Barat [19].

3. Faktor Lingkungan

Nyamuk *Anopheles* hidup di iklim tropis dan subtropics, namun bias juga hidup di daerah yang beriklim sedang. *Anopheles* juga ditemukan pada daerah pada daerah dengan ketinggian lebih dari 2000-2500m. Nyamuk *Anopheles* betina membutuhkan minimal 1 kali memangsa darah agar telurnya dapat berkembang biak. *Anopheles* mulai menggigit sejak matahari terbenam (jam 18.00) hingga subuh dan puncaknya pukul 19.00-21.00. Jarak terbang *Anopheles* tidak lebih dari 0,5 – 3 km dari tempat perindukannya. Waktu yang dibutuhkan untuk pertumbuhan (sejak telur menjadi dewasa) bervariasi antara 2-5 minggu, tergantung pada spesies, makanan yang tersedia dan suhu udara.

Dalam mencapai tingkat kesehatan dalam masyarakat, tidak hanya factor individu yang berpengaruh, terdapat juga beberapa factor lain seperti factor lingkungan fisik, factor biologis, factor sosial-ekonomi, dan faktor lainnya. Pada infeksi yang disebabkan oleh transmisi nyamuk, terdapat dua factor yang berpengaruh, yaitu factor iklim dan factor non-iklim [20]. Faktor iklim diantaranya terdiri dari suhu, kelembaban dan curah hujan. Faktor non-iklim berpengaruh besar terhadap kejadian malaria, dapat berpengaruh pada tempat perindukan vektor, transmisi malaria, dengan penjelasan sebagai berikut:

- a. Faktor lingkungan fisik berpengaruh pada perkembangbiakan vektor malaria ditinjau dari perairan yang menjadi tempat perindukannya. Perubahan lingkungan dapat memengaruhi peningkatan transmisi vektor malaria. Larva *anopheles* muncul lebih sering pada genangan air yang bersifat sementara di area buatan dibandingkan dengan rawa-rawa alami ataupun area perhutanan. Sejak rawa buatan mendapatkan penyinaran matahari lebih daripada rawa alami, suhu udara di rawa buatan lebih tinggi dibandingkan rawa alami. Habitat perairan di area persawahan yang diutamakan sebagai parit dan genangan air sementara menjadi sebagian dari aktivitas antropogenik.
- b. Perindukan vektor malaria juga dapat dipengaruhi oleh faktor lingkungan kimiawi yaitu kadar garam yang terdapat dalam zona perindukan, contohnya beberapa jenis nyamuk *anopheles* yang dapat berkembangbiak pada air payau dengan kadar air garam 12-18% dan tidak dapat berkembang biak pada kadar garam diatas 40%.
- c. Orang yang memiliki imunitas yang baik memiliki toleransi dan kesempatan lebih besar untuk tidak terinfeksi malaria dibandingkan dengan yang memiliki imunitas lemah. Selain itu, factor resistensi obat malaria

juga berpengaruh dalam factor infeksi malaria, setelah diberikan obat berkali-kali dapat juga menyebabkan suatu resistensi obat malaria.

Konsep Metode Valuasi Ekonomi

Penetapan nilai ekonomi total maupun nilai ekonomi kerusakan lingkungan digunakan pendekatan harga pasar dan pendekatan non pasar. Pendekatan harga pasar dapat dilakukan melalui pendekatan produktivitas, pendekatan modal manusia (*human capital*) atau pendekatan nilai yang hilang (*foregone earning*), dan pendekatan biaya kesempatan (*opportunity cost*). Sedangkan pendekatan harga non pasar dapat digunakan melalui pendekatan preferensi masyarakat (*non-market method*). Beberapa pendekatan non pasar yang dapat digunakan antara lain adalah metode nilai hedonis (*hedonic pricing*), metode biaya perjalanan (*travel cost*), metode kesediaan membayar atau kesediaan menerima ganti rugi (*contingent valuation*), dan metode *benefit transfer*. [21]

Pendekatan Modal Manusia (*Human Capital*)

Pada pendekatan ini, valuasi yang dilakukan untuk memberikan harga modal manusia yang terkena dampak akibat perubahan kualitas SDALH. Pendekatan ini sedapat mungkin menggunakan harga pasar sesungguhnya ataupun dengan harga bayangan. Hal ini terutama dapat dilakukan untuk memperhitungkan efek kesehatan dan bahkan kematian dapat dikuantifikasi harganya di pasar. Pendekatan ini dapat dilakukan melalui teknik:

- 1) Pendekatan Pendapatan yang Hilang,
- 2) Biaya Pengobatan,
- 3) Keefektifan Biaya Penanggulangan.

Adapun Pendekatan Biaya Pengobatan (*Medical Cost/Cost of Illness*) dapat diketahui dari dampak perubahan kualitas lingkungan yang berakibat negatif pada kesehatan, yaitu menyebabkan sekelompok masyarakat menjadi sakit.

Tahapan pelaksanaannya:

- a) Mengetahui bahwa telah terjadi gangguan kesehatan yang berakibat perlunya biaya pengobatan dan atau kerugian akibat penurunan produktifitas kerja.
- b) Mengetahui biaya pengobatan yang dibutuhkan sampai sembuh.
- c) Mengetahui kerugian akibat penurunan produktifitas kerja, misal dengan pendekatan tingkat upah atau harga produk yang dihasilkan.
- d) Menghitung total biaya pengobatan dan penurunan produktifitas kerja.

3. METODOLOGI PENELITIAN

Tempat dan Waktu Penelitian

Penelitian dilaksanakan di Laboratorium Inventarisasi dan Pemetaan Hutan Jurusan Kehutanan Fakultas Pertanian Universitas Lampung. Waktu penelitian dilakukan pada bulan Agustus—Oktober 2016. Jenis data

yang digunakan dalam penelitian ini adalah data primer dan data sekunder. Data primer berupa citra *Landsat* Provinsi Lampung tahun perekaman 2000, 2009 dan 2015 dan data tutupan lahan dari Kementerian Lingkungan Hidup dan Kehutanan. Data sekunder dalam penelitian ini meliputi peta administrasi kabupaten/kota Provinsi Lampung, data sekunder pendukung (angka kejadian malaria, kepadatan penduduk, jumlah penduduk) dari instansi terkait kepadatan penduduk kabupaten/kota di Provinsi Lampung.

Prosedur Analisis Data

Prosedur pengolahan citra

Analisis perubahan tutupan mangrove di Provinsi Lampung antara tahun 2000, 2009 dan 2015 membutuhkan peta tutupan lahan untuk setiap tahun yang diteliti serta data sekunder lain. Peta klasifikasi tutupan lahan dihasilkan melalui beberapa tahapan, yaitu: pra pengolahan citra, pengolahan citra digital, dan analisis perubahan tutupan lahan.

Analisis Regresi Linier Berganda

Berikut model dari analisis regresi linier berganda:

$$[Y]_{it} = \beta_0 + \beta_1[Bair] + \beta_2[HUTAN] + \beta_3[BLKR] + \beta_4[PMKM] + \beta_5[LTERBK] + \beta_6[PLKR] + \beta_7[SWH] + \beta_8[MRV] + \beta_9[RW] + \beta_{10}[TMBK] + \beta_{11}[KP] + e_{it} \quad (1)$$

Hipotesis

$$H_0 : \beta_1 = \beta_2 = \beta_3 = \beta_4 \dots \beta_{11} = 0$$

$$H_1 : \beta_1 \neq \beta_2 \neq \beta_3 \neq \beta_4 \dots \beta_{11} \neq 0 \quad ; \text{minimal ada satu } \beta_i \neq 0 \quad (2)$$

Adapun variabel, simbol dalam model, satuan, sumber data variabel *response* dan *predictor* disajikan dalam Tabel 1.

Tabel 1 Variabel, simbol dalam model, satuan dan skor, sumber data

No	Variabel	Simbol	Satuan dan Skor	Sumber Data
1	Angka Kesakitan Malaria	[Y]	Per 1000 Penduduk	Dinas Kesehatan Provinsi Lampung
2	Badan Air	[BAIR]	%	Interpretasi Citra, Peta RBI, KLHK
3	Hutan	[HUTAN]	%	Interpretasi Citra, Peta RBI, KLHK
4	Belukar	[BLKR]	%	Interpretasi Citra, Peta RBI, KLHK
5	Pemukiman	[PMKM]	%	Interpretasi Citra, Peta RBI, KLHK
6	Lahan Terbuka	[LTERBK]	%	Interpretasi Citra, Peta RBI, KLHK

7	Pertanian Lahan Kering	[PLKR]	%	Interpretasi Citra, Peta RBI, KLHK
8	Sawah	[SWH]	%	Interpretasi Citra, Peta RBI, KLHK
9	Hutan Mangrove	[MRV]	%	Interpretasi Citra, Peta RBI, KLHK
10	Tambak	[TMBK]	%	Interpretasi Citra, Peta RBI, KLHK
11	Kepadatan Penduduk	[KPD]	Jiwa/Km ²	BPS Provinsi Lampung

Uji F dilakukan untuk mengetahui pengaruh variabel bebas secara simultan terhadap variabel terikat. Uji t digunakan untuk menguji apakah variabel independen secara parsial berpengaruh terhadap variabel dependen. Tingkat signifikansi yang digunakan dalam penelitian adalah 1%, 5% dan 10%.

Penetapan Nilai Jasa Lingkungan

Biaya Kesehatan yang harus dikeluarkan untuk mengoati penyakit malaria diambil berdasarkan Peraturan Menteri Kesehatan Republik Indonesia Nomor 52 Tahun 2016 Tentang Standar Tarif Pelayan Kesehatan dalam Penyelenggaraan Progam Jaminan Kesehatan. Besarnya biaya yang dikeluarkan akibat malaria dihitung dengan menggunakan rumus

$$\text{Jumlah insidensi} \times \text{Biaya pengobatan penyakit.} \quad (3)$$

Biaya pengobatan penyakit malaria ditentukan / disesuaikan berdasarkan biaya tarif Indonesian-Case Based Groups atau Tarif INA-CBG yang merupakan besaran pembayaran klaim oleh BPJS Kesehatan kepada fasilitas Kesehatan Rujukan Tingkat Lanjutan atas paket layanan yang didasarkan kepada pengelompokan diagnosis dan prosedur. Melalui simulasi model dari hasil regresi linear dan input biaya pengobatan dapat diketahui besaran pengaruh perubahan kualitas lingkungan pada kisaran harga tertentu. Pengaruh tersebut kemudian akan dilihat apakah dengan menambah luas tutupan lahan tertentu akan mempengaruhi biaya pengobatan dalam kisaran harga tertentu juga. [22]

4. HASIL DAN PEMBAHASAN

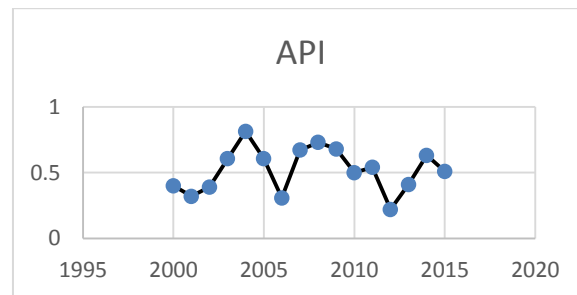
Hasil Penelitian

Angka Kesakitan Malaria di Provinsi Lampung

Malaria menduduki urutan kedelapan dari 10 besar penyakit penyebab utama kematian di Indonesia. Untuk mengetahui wilayah yang paling besar terinfeksi malaria adalah dengan cara mengetahui insiden kasus malaria tersebut. Oleh karena itu kasus malaria dapat distandarisasikan kedalam 2 indikator yaitu *Annual Malaria Incidence* (AMI) dan *Annual Parasite Incidence* (API). Menurut [22] indikator API menjadi sangat penting dalam menggambarkan besar masalah malaria disuatu wilayah. Karena indikator API menggambarkan kasus malaria secara lebih akurat dibandingkan dengan menggunakan indikator AMI dengan mensyaratkan bahwa

setiap kasus malaria harus dibuktikan dengan hasil pemeriksaan sediaan darah bukan hanya berdasarkan diagnosa gejala klinis.

Provinsi Lampung merupakan salah satu daerah endemis karena terdapat 10% Desa endemis dari seluruh jumlah Desa yang ada dengan angka kesakitan malaria per tahun 0,4 per 1000 penduduk. Indikator API juga digunakan untuk mengetahui tingkat endemisitas suatu wilayah, Sumatera khususnya Provinsi Lampung tergolong wilayah dengan tingkat endemisitas sedang dengan API berkisar antara 1 - <5 per 1000 penduduk. Rata-rata angka kesakitan API tahun 2000, 2009, dan 2015 sebesar 0.42 per 1000 penduduk dengan API berkisar antara 0.00—0.81 per 1000 penduduk.



Gambar 1 Trend angka kesakitan malaria Provinsi Lampung tahun 2000—2015

Penggunaan Lahan Ekosistem Mangrove

Data penggunaan lahan diidentifikasi melalui interpretasi citra *landsat* tahun 2000, 2009, dan 2015. Sehingga diperoleh persentase perubahan tutupan hutan mangrove. Total luas Lahan mangrove dari interpretasi sebesar 9401,62 Ha seperti dalam tabel berikut.

Tabel 2 Luas Mangrove, Tambak dan Rawa tahun 2000, 2009, 2015

No	Tahun	2000	2009	2015
1	Mangrove	4709,26	4693,29	9401,62
2	Tambak	32415,17	36999,89	31916,90
3	Rawa	46038,51	44592,33	49858,18

Sumber: Hasil Penelitian (2017)

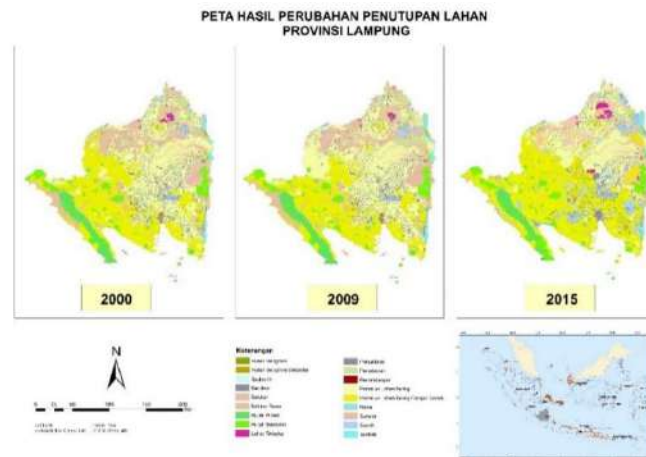
Data presentase tutupan mangrove di Provinsi Lampung disajikan dalam tabel berikut:

Tabel 3 Persentase (%) tutupan ekosistem mangrove terhadap luas wilayah kabupaten/kota di Provinsi Lampung tahun 2000 2009 dan 2015

no	Kabupaten/Kota	Mangrove		
		2000	2009	2015
1	Tulang Bawang	0,47	0,47	0,22
2	Lampung Timur	0,27	0,27	0,52
3	Lampung Selatan	0,15	0,13	0,15

Sumber: Hasil Penelitian (2017)

Hasil interpretasi perubahan tutupan lahan disajikan dalam gambar berikut.



Gambar 2 Perubahan tutupan lahan Provinsi Lampung tahun 2000, 2009, 2015

Kepadatan penduduk kabupaten/kota di Provinsi Lampung

Statistik kepadatan penduduk kabupaten/kota di Provinsi Lampung pada tahun 2002, 2009, dan 2014, ditampilkan pada Tabel 6, dengan satuan jiwa/km².

Tabel 4 Kepadatan penduduk (jiwa/km²) di kabupaten/kota Provinsi Lampung

No	Kabupaten	Tahun		
		2000	2009	2015
1	Lampung Barat	74	81	137
2	Tanggamus	238	255	190
3	Lampung Selatan	356	432	440
4	Lampung Timur	200	221	188
5	Lampung Tengah	218	250	326
6	Lampung Utara	195	210	222
7	Way Kanan	91	94	110
8	Tulang Bawang	89	103	124
9	Bandar Lampung	3.849	4.320	3308
10	Metro	1.917	2.205	2564

Sumber: Lampung Dalam Angka 2000—2015

Seluruh kabupaten/kota di Provinsi Lampung mengalami peningkatan kepadatan penduduk. Kota Bandar Lampung menjadi wilayah dengan kepadatan penduduk tertinggi di Provinsi Lampung, statusnya sebagai pusat pemerintahan dan perekonomian, merupakan daya tarik bagi masyarakat untuk menjadikan kota ini sebagai tempat tinggal dan tujuan urbanisasi dari wilayah lain di Provinsi Lampung. Selanjutnya Kota Metro memiliki kepadatan penduduk tertinggi kedua di Provinsi Lampung, dengan statusnya sebagai kota pendidikan dan letaknya yang relatif dekat dengan Kota Bandar Lampung, menjadikannya sebagai tujuan urbanisasi, baik untuk

bekerja maupun belajar, sehingga meningkatkan kepadatan penduduknya. Wilayah dengan kepadatan penduduk nomor tiga di Provinsi Lampung adalah Kabupaten Lampung Selatan, dan merupakan zona penyangga bagi Kota Bandar Lampung dan Metro. Kabupaten Lampung Tengah, Lampung Timur, Lampung Utara, dan Tanggamus adalah kawasan pertanian dan perkebunan, sehingga kepadatan penduduk di wilayah tersebut berkisar 199—283 jiwa/km². Letak yang jauh dari pusat pemerintahan dan ekonomi, aksesabilitas yang sulit, dan berada di perbatasan wilayah administrasi Provinsi Lampung menyebabkan Kabupaten Lampung Barat, Tulang Bawang, dan Way Kanan menjadi tiga wilayah dengan tingkat kepadatan penduduk yang lebih rendah dibanding wilayah lain di Provinsi Lampung [23].

Tarif Dasar Pelayanan Kesehatan Provinsi Lampung

Pengukuran valuasi hutan mangrove dapat diukur melalui perubahan dampak lingkungan yang diukur melalui perubahan lahan ekosistem mangrove serta pendekatan modal manusia (Human capital) sesuai dengan Permen LH tahun 2012 tentang Panduan Valuasi Ekonomi Ekosistem Hutan. Pada pendekatan ini, valuasi yang dilakukan untuk memberikan harga modal manusia yang terkena dampak akibat perubahan kualitas SDALH. Pendekatan ini sedapat mungkin menggunakan harga pasar sesungguhnya ataupun dengan harga bayangan. Hal ini terutama dapat dilakukan untuk memperhitungkan efek kesehatan dan bahkan kematian dapat dikuantifikasi harganya di pasar. Pendekatan ini dapat dilakukan melalui teknik biaya pengobatan yaitu dengan mengetahui biaya pengobatan yang dibutuhkan sampai sembuh. Berdasarkan Permenkes No 52 Tahun 2016 Tentang Standar Tarif Pelayanan Kesehatan Dalam Penyelenggaraan Program Jaminan Kesehatan, Penyakit malaria termasuk dalam Penyakit Infeksi Bakteri dan Parasit lain-lain dengan Code A-4-14-I.

Data statistik deskriptif biaya pengobatan penyakit malaria disajikan dalam table berikut:

Tabel 5 Statistik deskriptif biaya pengobatan penyakit malaria Provinsi Lampung

No	Variabel	N	Mean	SE	StDev	Min	Median	Max
Mean								
1	Tarif Kelas 3	24	2931308	160265	785136	2111300	2668750	4865100
2	Tarif Kelas 2	24	3517567	192320	942171	2533600	3202500	5838200
3	Tarif Kelas 1	24	4103833	224374	1099203	2955900	3736250	6811200

Sumber : Permenkes No 52 Tahun 2016 (2016)

Pembahasan

Hasil dari uji t dan R square atau koefisien determinasi disajikan dalam tabel berikut:

Tabel 6 Hasil uji t dan koefisien determinasi

Predictor	Symbol	Coef	SE Coef	T	P
Constant	[Co]	0,4961	0,1959	2,53	0,028
Badan Air	[BAIR]	-0,1386	0,2199	-0,63	0,541
Hutan	[HUTAN]	0,027416	0,005990	4,58	0,001
Belukar	[BLKR]	0,027176	0,007557	3,60	0,004

Pemukiman	[PMKM]	0,07306	0,02321	3,15	0,009
Lahan Terbuka	[LTERBK]	-0,2706	0,1141	-2,37	0,037
Pertanian Lahan Kering	[PLKR]	-0,025072	0,006150	-4,08	0,002
Sawah	[SWH]	-0,07231	0,01610	-4,49	0,001
Tambak	[TMBK]	0,21618	0,07340	2,95	0,013
Mangrove	[MRV]	-0,07937	0,01741	-4,56	0,001
Kepadatan Penduduk	[KPD]	-0,0003569	0,0002059	-1,73	0,111
S = 0,135332		R-Sq = 93,5%		R-Sq(adj) = 87,7%	

Sumber: Hasil Penelitian (2017)

R Square di dalam mintab ditunjukkan dengan nilai R-Sq di mana pada uji ini nilainya dapat dilihat di output session yaitu sebesar 95,3% artinya variabel independen yang digunakan mampu menjelaskan kejadian Malaria secara serentak atau simultan sebesar 95,3% sedangkan sisanya 6,5% dijelaskan oleh variabel lain di luar model yang tidak diteliti.

Berdasarkan hasil uji F, dapat disimpulkan bahwa secara keseluruhan variabel prediktor mempunyai pengaruh yang nyata. Insidensi malaria yang terjadi dapat dijelaskan menggunakan variabel X dengan kemungkinan meleset 0,000 atau 4 insidensi per 10000 penduduk.

Penggunaan Lahan Ekosistem Mangrove

Berdasarkan hasil uji t dan koefisien determinasi terdapat hubungan yang nyata antara luas tutupan mangrove terhadap kejadian malaria dengan p-value 0,001 pada taraf 10%. Tutupan mangrove memiliki nilai koefisien negatif sebesar -0,07937 artinya setiap penambahan 1 % tutupan mngrove akan mengurangi kejadian malaria sebesar 0,07937/1000 penduduk. Luas tutupan Tambak memiliki nilai koefisien 0,21618 dengan nilai p value 0,013, artinya setiap kenaikan 1 % tutupan tambak akan mengakibatkan penaikan kejadian malaria sebesar 0,21618 /1000 penduduk

Penggunaan Model Regresi sebagai Pendekatan Valuasi Jasa Lingkungan

Nilai koefisien persamaan regresi mangrove adalah -0,07937. Luas ekosistem mangrove sebesar 9401,62 Ha, apabila terjadi penaikan 10 persen luas tutupan mangrove atau 940,162 Ha maka akan menurunkan kejadian malaria sebesar $0,07937 \times 10\% = 0,007937$ per 1000 penduduk atau 0,000007937 insidensi. Hal ini berarti setiap penambahan 10 persen tutupan mangrove akan mengakibatkan penurunan kejadian penyakit malaria sebesar $0,000007937 \times 8117268 = 64,42676$ insidensi dengan asumsi bahwa setiap faktor penularan malaria adalah positif.

Menurut epidemiologi, penularan malaria secara alamiah terjadi akibat adanya interaksi antara tiga faktor yaitu *Host*, *Agent*, dan *Environment*. Manusia adalah host vertebrata dari *Human plasmodium*, nyamuk sebagai *Host invertebrate*, sementara *Plasmodium* sebagai parasit malaria sebagai agent penyebab penyakit yang

sesungguhnya, sedangkan faktor lingkungan dapat dikaitkan dalam beberapa aspek, seperti aspek fisik, biologi dan sosial ekonomi.

Tabel 7 Simulasi Valuasi Jasa lingkungan Mangrove dengan pendekatan human capital

No	Kelas Pengobatan	Insidensi / 10%	mangrove/10% (Ha)	Harga Pengobatan	Total (Rp)	Total manfaat/Ha (Rp)
1	Tarif Kelas 3	64,42676	940,162	2931308	188854665,6	200874,6
2	Tarif Kelas 2	64,42676	940,162	3517567	226625431,2	241049,3
3	Tarif Kelas 1	64,42676	940,162	4103833	264396647,8	281224,6
Rata-rata				3517569	226625581,5	241049,5

Sumber: Hasil Penelitian (2017)

Dari tabel simulasi diatas diketahui bahwa manfaat hutan mangrove adalah sebesar 241.049,5/Ha. Tabel 12 diatas juga menjelaskan bahwa dengan berkurangnya luas mangrove sebesar 10% akan mengakibatkan kerugian sebesar Rp. 226.625.581,-. Jika luas total mangrove adalah 9401,62Ha maka total manfaat dari hutan mangrove dengan pendekatan human capital adalah Rp. 2.266.255.815,5,-/tahun.

5. SIMPULAN

Penelitian ini membuktikan bahwa terdapat hubungan antara perubahan tutupan lahan dan kepadatan penduduk dengan insidensi malaria di Provinsi Lampung pada taraf 10 %. Kelas tutupan lahan yang berpengaruh nyata terhadap insidensi malaria pada balita adalah hutan mangrove dengan $p\text{-value} = 0,001$. Setiap kenaikan luas tutupan mangrove sebesar 10% akan mengakibatkan penurunan kejadian malaria sebesar 0,007937 per 1000 penduduk. Nilai manfaat jasa lingkungan mangrove di Provinsi Lampung dengan pendekatan *human capital* / *medical cost* malaria adalah Rp. 2.266.255.815,5,-.

KEPUSTAKAAN

- [1] Kustanti, A. (2011). *Manajemen Hutan Mangrove*. Institut Pertanian Bogor Press. Bogor.
- [2] Rahmawaty. (2006). *Upaya pelestarian mangrove berdasarkan pendekatan masyarakat*. Fakultas Pertanian Universitas Sumatera Utara. Medan.
- [3] Bakosurtanal. (2009). *Luas Kawasan Mangrove Indonesia*. Barkosurtanal. Bogor
- [4] Siregar, A, F. (2012). *Valuasi Ekonomi dan Analisis Konservasi Hutan Mangrove di Kabupaten Kubu Raya Provinsi Kalimantan Barat*. Thesis. Institut Pertanian Bogor. Bogor.
- [5] Ahmadi, S. (2008). *Faktor risiko kejadian malaria di Desa Lubuk Nipis Kecamatan Tanjung Agung Kabupaten Muara Enim*. Tesis. Program Studi Magister Kesehatan Lingkungan Universitas Diponegoro. Semarang.
- [6] Putra, A.K. (2015). *Peranan Ekosistem Hutan Mangrove Pada Imunitas Terhadap Malaria: Studi di Kecamatan Labuhan Maringgai Kabupaten Lampung Timur*. Jurnal Sylva Lestari. Vol 3 No 2 Mei 2015. 67-78.

- [7] Dinas Kesehatan Provinsi Lampung. (2004). *Profil Kesehatan Provinsi Lampung Tahun 2003*. Buku. Dinas Kesehatan Provinsi Lampung. Bandar Lampung. 69 hlm.
- [8] Chapman, V.J. editor. (1978). *Botanical Surveys in Mangroves Communities*. Dalam The Mangroves Ecosystem: Research Methods. UNESCO, Monograph on Oceanological Methodology 8, Paris. 53-80p.
- [9] Van Steenis, C.G.G.J. (1958). *Ecology of mangroves*. Introduction to Account of The Rhizophoraceae by Ding Hou, Flora Malesiana, Ser. I, 5: 431-441.
- [10] Kantor Menteri Negara Lingkungan Hidup. 1993. *Pengelolaan Ekosistem Hutan Mangrove*. Prosiding Lokakarya Pemantapan Strategi Pengelolaan Lingkungan Wilayah Pesisir dan Lautan dalam Pembangunan Jangka Panjang Tahap Kedua. Kapal Kerinci, 11-13 September 1993, 47p.
- [11] Saputro, G.B. (2009). *Peta Mangrove Indonesia*. Buku. Pusat Survei Sumber Daya Alam Laut, Badan Koordinasi Survei dan Pemetaan Nasional (Bakosurtanal). Jakarta.
- [12] Masela, D. F. (2012). *Pengaruh struktur dan komposisi mangrove bagi kepadatan nyamuk di Desa Kopi dan Desa Minanga Kecamatan Bintauna*. Jurnal Cocos. 1(2): 1—8.
- [13] AchmadiUF.(2005).*ManajemenPenyakitBerbasisWilayah*.Jakarta:Kompas.
- [14] Aronoff, S. (1989). *Geographic Information Systems: A Management Perspective*. Ottawa: WDI Publications.
- [15] Burrough, P.A. (1986). *Principles of Geographic Information Systems for and Resources Assessment*. Clarendon Press, Oxford.
- [16] Prahasta, E. (2009). *Konsep-Konsep Dasar Sistem Informasi Geografis*. Bandung: Informatika.
- [17] Prabowo A., (2004). *Malaria, Mencegah dan Mengatasinya*. Cetakan 1. Puspa Swara. Jakarta.
- [18] HarijantoPN.(2000).*Malaria:Epidemiologi,Patogenesis,ManifestasiKlinis,danPenanganan*.EGC.Jakarta.
- [19] Departemen KesehatanRI. (2005). *PedomanPenatalaksanaan Kasus Malaria diIndonesia*.Jakarta:Depkes RI.
- [20] EbersonF.(2011).*Communicable DiseasesPart1GeneralPrinciples,Vaccine- Preventable DiseaseandMalaria*.Ethiopia:FederalDemocratic Republicof EthiopiaMinistryofHealth.
- [21] Kentrin Lingkungan Hidup Dan Kehutanan. Permenlhk tahun 2012 tentang Panduan Valuasi Ekonomi Ekosistem Hutan
- [22] Kementrian Kesehatan RI. Permenkes No 52 Tahun 2016 Tentang Stander Tarif Pelayanan Kesehatan DalamPenyelenggaraan Progam Jaminan Kesehatan
- [23] Badan Pusat Statistik Provinsi Lampung. (2015). *Lampung dalam Angka 2015*. Buku. Badan Pusat Statistik Provinsi Lampung. Bandar Lampung. Hlm 415.

ANALISIS PENGENDALIAN PERSEDIAAN DALAM MENCAPAI TINGKAT PRODUKSI *CRUDE PALM OIL* (CPO) YANG OPTIMAL DI PT. KRESNA DUTA AGROINDO LANGLING MERANGIN-JAMBI

Marcelly Widya W., Heri Wibowo, Estika Devi Erinda

Program Studi Teknik Industri Universitas Malahayati

Jl. Pramuka No.27 Kemiling Bandar Lampung 35153

Email : marcelly.widya@gmail.com, heriwibowo_ti@yahoo.co.id

Abstrak

Pengendalian produksi merupakan sistem mendayagunakan sumber daya secara efektif, untuk mencapai keseimbangan produksi dengan biaya yang minimum dan menciptakan keuntungan bagi perusahaan. PT. Kresna Duta Agroindo merupakan pabrik kelapa sawit yang beroperasi menghasilkan produk standar sehingga sebagian produk untuk persediaan. Kegiatan produksi berjalan lancar apabila perusahaan dapat mempertahankan jumlah persediaan optimal dalam memenuhi kebutuhan bahan. Penelitian ini bertujuan untuk menerapkan pengendalian produksi dalam menentukan tingkat produksi optimal Crude Palm Oil (CPO) dan biaya persediaan yang minimum. Pendekatan yang digunakan adalah dengan metode pengendalian persediaan Economic Order Quantity (EOQ). Hasil penelitian menunjukkan bahwa tingkat optimal produksi CPO sebesar 2.747.294,1 kg per bulan, interval waktu optimal yang dibutuhkan untuk memproduksi CPO selama 24 bulan, biaya total persediaan minimum setiap bulannya adalah sebesar Rp. 405.971.438,64. Metode EOQ dapat menghemat jumlah produksi CPO 4,2% yaitu sebanyak 120.641,192 kg dan biaya 11,4% sebesar Rp. 52.093.630,22.

Kata kunci: *Crude Palm Oil (CPO), Economic Order Quantity (EOQ), Pengendalian Persediaan*

1. PENDAHULUAN

Pengendalian produksi sering disebut sebagai sistem produksi yang merupakan suatu sistem untuk membuat produk (mengubah bahan baku menjadi barang) yang melibatkan fungsi manajemen untuk merencanakan dan mengendalikan proses produksi tersebut [1]. Kelangsungan kegiatan produksi akan berjalan lancar apabila perusahaan dapat mempertahankan jumlah persediaan yang optimal sehingga dapat memenuhi kebutuhan bahan dalam jumlah dan waktu yang tepat dengan total biaya seminimal mungkin, dengan kata lain dapat memenuhi kebutuhan setiap saat pada kondisi yang ekonomis ditinjau dari ongkos-ongkos yang timbul akibat adanya persediaan [2].

Masalah penentuan besarnya persediaan merupakan masalah yang penting bagi perusahaan dalam pengendalian produksi, karena pengendalian produksi mempunyai efek terhadap keuntungan perusahaan. Adanya persediaan bahan baku yang terlalu besar dibandingkan kebutuhan perusahaan dalam memproduksi, maka akan menambah beban biaya penyimpanan dan pemeliharaan dalam gudang, serta adanya kemungkinan terjadinya penyusutan kualitas yang tidak bisa dipertahankan sehingga perusahaan akan mengalami kerugian. Dan sebaliknya, jika

persediaan bahan yang terlalu kecil akan mengakibatkan kemacetan dalam produksi, dan permintaan konsumen tidak terpenuhi.

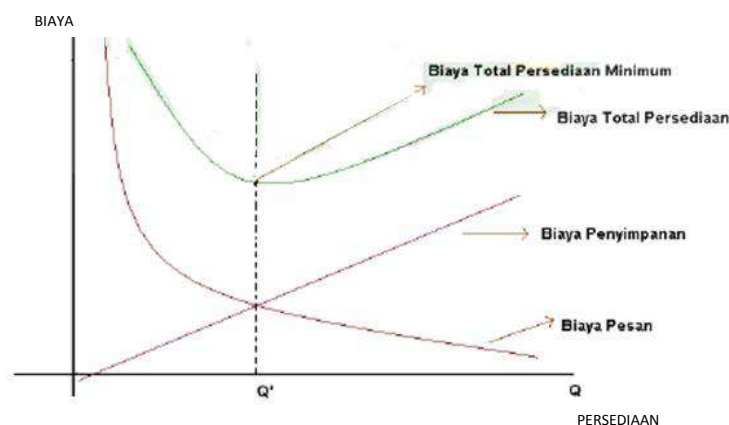
PT. Kresna Duta Agroindo merupakan salah satu perusahaan yang bergerak dalam bidang perkebunan dan pabrik kelapa sawit (PKS). Produk yang dihasilkan oleh PT. Kresna Duta Agroindo ini adalah *Crude Palm Oil* (CPO) dengan target kapasitas produksi 24,25% dari bahan baku yang masuk dan Kernel dengan target kapasitas produksi 5,6% dari bahan baku yang masuk. Kegiatan pengolahan kelapa sawit menjadi CPO di PT.Kresna Duta Agroindo mengalami masalah pada fluktuasi bahan baku yang cukup intens ketika musim panen sawit tiba, sehingga menyebabkan penumpukan bahan baku yaitu tandan buah segar (TBS). Oleh karena itu, perlu dilakukan pengendalian produksi dan pengendalian persediaan dengan perencanaan yang seefisien mungkin untuk mendapatkan biaya persediaan yang minimum.

2. LANDASAN TEORI

2.1. Persediaan

Persediaan merupakan bahan-bahan, bagian yang disediakan, dan bahan-bahan dalam proses yang terdapat dalam perusahaan untuk proses produksi, serta barang-barang jadi atau produk yang disediakan untuk memenuhi permintaan dari konsumen atau pelanggan setiap waktu yang dirawat menurut aturan tertentu dalam tempat persediaan agar selalu dalam keadaan siap pakai dan dicatat dalam bentuk buku perusahaan [3]. Pendapat lain tentang definisi persediaan, persediaan merupakan jumlah bahan-bahan *parts* yang disediakan dan bahan-bahan dalam proses yang terdapat dalam perusahaan untuk proses produksi, serta barang-barang jadi/produk yang disediakan untuk memenuhi permintaan dari komponen atau langganan setiap waktu [4].

Persediaan pada hakikatnya merupakan stok yang dibutuhkan perusahaan untuk mengatasi adanya fluktuasi permintaan. Persediaan dalam proses produksi dapat diartikan sebagai sumberdaya yang menganggur, hal ini dikarenakan sumberdaya tersebut masih menunggu dan belum digunakan pada proses berikutnya, sehingga sumberdaya tersebut harus disimpan yang akan menimbulkan biaya penyimpanan. Pengendalian persediaan digunakan untuk menentukan berapa banyak sumberdaya yang dibutuhkan dan kapan sumber daya itu dibutuhkan, sehingga dapat meminimumkan biaya persediaan.



Gambar 1. Kurva Biaya Persediaan

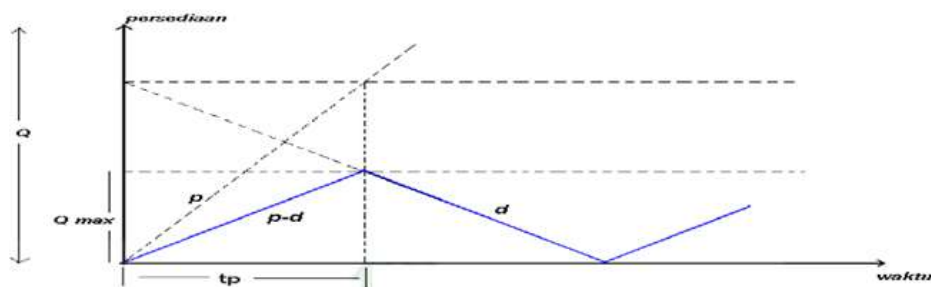
2.2. Pengendalian Persediaan *Economic Order Quantity* (EOQ)

Economic Order Quantity merupakan suatu model persediaan untuk menentukan jumlah pembelian bahan yang akan dapat mencapai biaya persediaan yang paling minimal atau dapat ditekan serendah-rendahnya, sehingga efisiensi persediaan sumberdaya dalam perusahaan dapat tercapai [5].

Perumusan model yang biasa digunakan adalah model persediaan dengan *stock*, dimana variabel-variabel yang digunakan adalah sebagai berikut :

- P = kecepatan produksi per satuan waktu
- d = jumlah penyaluran produksi per satuan waktu
- p-d = tingkat pertumbuhan persediaan
- D = permintaan pada setiap periode
- Q = jumlah pertambahan persediaan atau produksi untuk setiap kali pertambahan atau produksi
- t_p = periode waktu penambahan atau produksi
- S = jumlah persediaan yang dipesan setiap kali pesan

Model persediaan yang tepat dengan keadaan perusahaan ini dapat diilustrasikan sebagai berikut [6] :



Gambar 2 Grafik Model Persediaan

Berdasarkan gambar 2 dapat diketahui bahwa pertambahan persediaan terjadi selama t_p , maka $Q_{\max}$ tersebut akan habis terpakai, sehingga persediaan rata-ratanya menjadi [7] :

$$\frac{Q_{\max}}{2} = \frac{t_p(p-d)}{2} \quad (1)$$

untuk memenuhi persediaan Q diperlukan waktu selama t_p dengan tingkat pertambahan persediaan sebesar p maka :

$$Q = t_p \cdot p \text{ atau } t_p = \frac{Q}{p} \quad (2)$$

jika persamaan (2) disubstitusikan ke persamaan (1) persediaan rata-rata itu Q_{AVE} akan menjadi :

$$Q_{AVE} = \frac{\frac{Q}{p}(p-d)}{2} \text{ atau } Q_{AVE} = \frac{Q}{2} \left(1 - \frac{d}{p}\right) \quad (3)$$

bila biaya simpan per unit setiap periode adalah h maka biaya simpan (BS) :

$$BS = \frac{Q}{2} h \left(1 - \frac{d}{p}\right) \quad (4)$$

$$BP = \frac{D}{Q} S \quad (5)$$

Maka biaya total persediaan (BTP) adalah :

Biaya total persediaan = biaya pesan + biaya simpan

$$BTP = \frac{D}{Q} S + \frac{Q}{2} h \left(1 - \frac{d}{p}\right) \quad (6)$$

agar diperoleh biaya total persediaan minimum maka persamaan (6) harus diminimumkan untuk Q , syarat $BTP =$

$f(Q)$ minimum adalah $\frac{d(BTP)}{d(Q)}$ sehingga dari persamaan (6), diperoleh :

$$BTP = \frac{D}{Q} S + \frac{Q}{2} h \left(1 - \frac{d}{p}\right)$$

$$\frac{d(BTP)}{d(Q)} = \frac{DS}{Q^2} + \frac{h}{2} \left(1 - \frac{d}{p}\right) \text{ karena } \frac{d(BTP)}{d(Q)} = 0$$

$$\frac{DS}{Q^2} + \frac{h}{2} \left(1 - \frac{d}{p}\right) = 0$$

$$\frac{h}{2} \left(1 - \frac{d}{p}\right) = \frac{DS}{Q^2}$$

sehingga persediaan optimal untuk setiap produksi adalah :

$$Q = \sqrt{\frac{2DS}{h \left(1 - \frac{d}{p}\right)}} \quad (7)$$

dan waktu optimal yang dibutuhkan untuk satu putaran produksi adalah :

$$t_0 = \sqrt{\frac{Q_0}{d}} \quad (8)$$

substitusi persamaan (7) ke persamaan (8)

$$t_0 = \sqrt{\frac{2.S}{h.d.\left(1-\frac{d}{p}\right)}} \quad (9)$$

Bila Q optimal pada persamaan (7) disubstitusikan ke persamaan (6), maka diperoleh model matematik untuk biaya total persediaan minimum :

$$BTP = \frac{d}{S}.S + \frac{Q}{2}h\left(1-\frac{d}{p}\right)$$

$$BTP.\sqrt{\frac{2Sd}{h\left(1-\frac{d}{p}\right)}} = 2 S d$$

$$BTP = \frac{d}{\sqrt{\frac{2Sd}{h\left(1-\frac{d}{p}\right)}}}.S + \frac{\sqrt{\frac{2Sd}{h\left(1-\frac{d}{p}\right)}}}{2}.h.\left(1-\frac{d}{p}\right)$$

$$BTP.\sqrt{\frac{2Sd}{h\left(1-\frac{d}{p}\right)}} = S.d + \frac{h}{2}\left(1-\frac{d}{p}\right)$$

$$BTP = \frac{2Sd}{\sqrt{\frac{2sd}{h\left(1-\frac{d}{p}\right)}}}\sqrt{2.d.S.h\left(1-\frac{d}{p}\right)}$$

jadi biaya total persediaan minimum per satuan waktu adalah :

$$BTP = \sqrt{2.d.S.h\left(1-\frac{d}{p}\right)} \quad (10)$$

2.3. Uji Normalitas Metode Lilliefors

Metode lilliefors menggunakan data dasar yang belum diolah dalam tabel distribusi frekuensi. Data tersebut kemudian ditransformasikan dalam nilai Z untuk dapat dihitung luasan kurva normal sebagai probabilitas kumulatif normal. Probabilitas tersebut dicari bedanya dengan probabilitas kumulatif empiris. Beda terbesar dibanding dengan tabel lilliefors pada tabel nilai quantil statistik lilliefors berdistribusi normal [8].

Terdapat persyaratan untuk menggunakan metode lilliefors ini, yaitu [8]:

1. Data berskala interval atau ratio (kuantitatif).
2. Data tunggal / belum dikelompokkan pada tabel distribusi frekuensi.
3. Dapat untuk n besar maupun n kecil.
4. Ukuran sampel $n \leq 30$.

Signifikansi uji, nilai terbesar $F(z_i) - S(z_i)$ dibandingkan dengan nilai tabel Lilliefors. Jika nilai $F(z_i) - S(z_i)$ terbesar kurang dari nilai tabel Lilliefors, maka H_0 diterima ; H_1 ditolak. Jika nilai $F(z_i) - S(z_i)$ terbesar lebih besar dari nilai tabel Lilliefors, maka H_0 ditolak ; H_1 diterima. Tabel nilai Quantil Statistik Lilliefors.

Perumusan ilmu statistika juga berguna dalam pengendalian persediaan dan biasanya digunakan untuk mengetahui pola distribusi apa yang dipakai. Pola distribusi itu dapat diketahui dengan melakukan uji lilliefors. Misalkan sampel berukuran n dengan nilai data : $x_1, x_2, x_3, \dots, x_n$. Berdasarkan sampel ini akan diuji dua hipotesa, sebagai berikut :

H_0 : Data yang diperoleh berdistribusi normal.

H_1 : Data yang diperoleh tidak berdistribusi normal.

Prosedur yang harus dilakukan untuk pengujian hipotesa antara lain :

1. Nilai data $x_1, x_2, x_3, \dots, x_n$, dijadikan angka baku $Z_1, Z_2, Z_3, \dots, Z_n$ dengan menggunakan rumus :

$$Z = \frac{x_i - \bar{x}}{S}$$

dengan $\bar{x}$ = rata-rata sampel

S = simpangan baku

i = 1, 2, 3, n

untuk menghitung rata-rata sampel pengamatan digunakan rumus :

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$$

dan untuk menghitung simpangan baku dari sampel digunakan rumus :

$$S = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1}}$$

2. Hitung peluang $F(Z_1) = P(Z=Z_1)$ dengan menggunakan daftar distribusi normal standar.
3. Hitung proporsi $Z_1, Z_2, Z_3, \dots, Z_n \leq Z_1$, jika proporsi ini dinyatakan sebagai $S(Z_1)$, maka :

$$S(Z_1) = \frac{Z_1 \cdot Z_2 \cdot Z_3 \cdot \dots \cdot Z_n}{n}$$

4. Hitung selisih antara $F(Z_1)$ dengan $S(Z_1)$, yaitu :

$$|F(Z_1) - S(Z_1)|$$

5. Hitung harga maksimum antara $|F(Z_1) - S(Z_1)|$, yaitu :

$$L_{\max} = |F(Z_1) - S(Z_1)|$$

6. Kriteria pengambilan keputusan adalah :

Jika

$$L = \begin{cases} \leq L_{\alpha}(n) : \text{maka } H_0 \text{ diterima} \\ > L_{\alpha}(n) : \text{maka } H_0 \text{ ditolak} \end{cases}$$

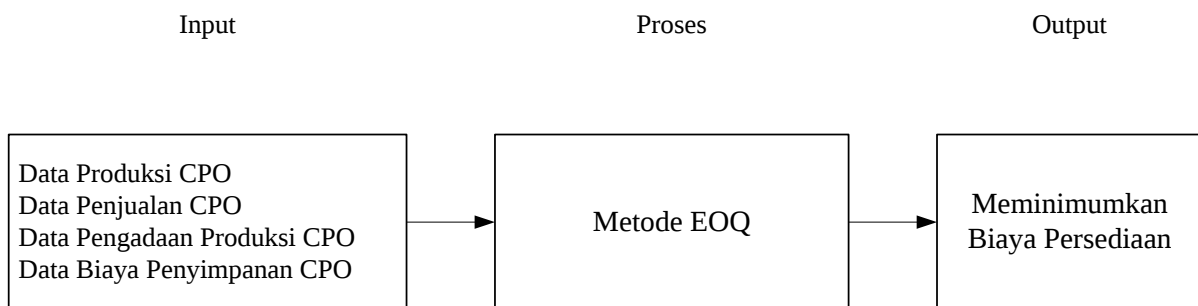
Dengan $L_{\alpha}(n)$ adalah nilai kritis uji kenormalan liliefors dengan taraf nyata α dan banyaknya sampel n .

3. METODOLOGI PENELITIAN

Berbagai penelitian tentang Pengendalian Persediaan sudah banyak dilakukan antara lain :

1. Lumempow, Veyro E.L, dkk. Meneliti tentang Aplikasi Metode EOQ pada Persediaan BBM di PT. Sarana Samudera Pacific Bitung. Penelitian tersebut menggunakan variabel persediaan produk jadi. Hasil dari penelitian tersebut adalah bahwa metode EOQ dapat menurunkan biaya persediaan sebesar 0,3 % [9].
2. Indropasto dan Erma Suryani. Penelitian tersebut menggunakan metode EOQ dan Algoritma Genetika dalam pengendalian persediaan. Hasil dari penelitian tersebut adalah bahwa metode yang diterapkan dapat menurunkan biaya persediaan [10].
3. Taufik Limansyah. Penelitian yang dilakukan adalah aplikasi metode EOQ dengan variabel penelitian faktor kadaluarsa dan unit diskon. Hasil dari penelitian tersebut adalah penerapan metode EOQ dapat meminimumkan biaya persediaan [11].

Berdasarkan penelitian yang telah dilakukan maka dalam penelitian ini menerapkan pengendalian persediaan menggunakan metode EOQ dengan variabel penelitian produk jadi dan stok persediaan. Sehingga didapatkan kerangka penelitian adalah sebagai berikut :



Gambar 3. Kerangka Pikir Penelitian

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan Data

Data-data yang diperlukan dalam penelitian ini adalah :

1. Data jumlah produksi CPO
2. Data jumlah penyaluran produksi CPO
3. Data biaya pengadaan produksi (*set-up*)
4. Data biaya penyimpanan (*carrying cost*)

Berikut ini adalah data-data yang didapat dari PT. Kresna Duta Agroindo

Tabel 1. Jumlah Produksi CPO selama 2 periode

No	Bulan	Tahun (kg)	
		Periode 1	Periode 2
1	Januari	1.753.177	2.218.256
2	Februari	2.001.295	2.222.758
3	Maret	2.315.258	2.080.752
4	April	2.515.894	2.627.914
5	Mei	2.756.753	2.962.426
6	Juni	2.892.582	3.081.413
7	Juli	3.648.882	3.761.891
8	Agustus	3.524.681	3.892.544
9	September	3.090.561	3.708.159
10	Oktober	3.199.386	3.369.942
11	November	3.007.189	2.838.407
12	Desember	2.785.445	2.575.482
Jumlah		33.491.103	35.339.344

Tabel 1 diatas merupakan data jumlah produksi CPO selama 2 periode, sedangkan untuk tabel 2 (bawah), merupakan jumlah penyaluran produksi CPO (Penjualan CPO) selama 2 periode.

Tabel 2. Jumlah Penyaluran Produksi CPO selama 2 periode

No	Bulan	Tahun (kg)	
		Periode 1	Periode 2
1	Januari	1.891.700	2.581.794
2	Februari	2.510.220	2.190.812
3	Maret	2.095.260	1.850.643
4	April	2.252.430	2.683.275
5	Mei	2.342.830	3.127.027
6	Juni	2.700.440	2.823.261
7	Juli	2.801.625	3.059.254
8	Agustus	3.602.309	3.530.498
9	September	3.175.871	2.087.209
10	Oktober	3.135.482	3.191.273
11	November	3.053.186	3.035.800
12	Desember	2.493.140	2.682.241
Jumlah		32.006.493	32.843.087

Sedangkan untuk data selanjutnya merupakan data biaya pengadaan CPO (tabel 3) dan data harga penjualan CPO (tabel 4) untuk 2 periode.

Tabel 3. Data Biaya Pengadaan CPO selama 2 periode

No	Bulan	Tahun (Rp)	
		Periode 1	Periode 2
1	Januari	1.777.262.475	6.032.429.875
2	Februari	1.476.939.425	5.171.401.850
3	Maret	2.272.885.975	3.381.160.200
4	April	3.318.899.375	7.128.285.350
5	Mei	3.854.858.450	7.123.862.825
6	Juni	4.551.584.250	7.125.825.350
7	Juli	6.994.925.820	7.354.518.025
8	Agustus	5.513.678.100	6.034.319.725
9	September	4.646.106.975	5.647.972.975
10	Oktober	6.285.377.350	3.070.849.000
11	November	6.370.798.075	2.217.024.825
12	Desember	5.678.525.550	1.778.294.850
Jumlah		52.739.375.825	62.194.802.950

Tabel 4. Data Harga Pokok Penjualan CPO selama 2 periode

No	Periode	Biaya/kg
1	Periode 1	Rp. 5.000,00
2	Periode 2	RP. 8.299,00
Jumlah		Rp. 13.299,00

4.2. Uji Kenormalan Data

Dalam pengolahan data ini, data yang sudah didapat dilakukan uji kenormalan menggunakan metode Lilliefors. Berikut ini adalah uji kenormalan untuk data penyaluran produksi CPO.

Tabel 5. Uji Kenormalan Data Penyaluran Produksi CPO periode 1

No	x_i	$x_i - \bar{X}$	$(x_i - \bar{X})^2$	z_i	F (z_i)	S (z_i)	F (z_i) - S (z_i)
1	1.891.700	-775.507,75	780.144.252.935,06	-1,56	0,0594	0,0833	0,0239
2	2.510.220	-156.987,75	46.198.9031.355,06	-1,20	0,1151	0,1666	0,0515
3	2.095.260	-571.947,75	262.684.694.520,06	-0,91	0,1814	0,2500	0,0686
4	2.252.430	-414.777,75	56.995.713.275,06	-0,41	0,3372	0,4166	0,0794
5	2.342.830	-324.377,75	7.943.755.820,06	-0,15	0,4404	0,5000	0,0594
6	2.700.440	33.232,25	15.745.795.065,06	0,22	0,5871	0,5833	0,0038
7	2.801.625	134.417,25	28.747.531.997,56	0,95	0,8289	0,9166	0,0877
8	3.602.309	935.101,25	823.286.290.876,56	1,61	0,9463	1,0000	0,0535
9	3.175.871	508.663,25	94.502.906.275,56	0,54	0,7054	0,7500	0,0446
10	3.135.482	468.274,25	78.693.854.838,06	0,49	0,6879	0,6666	0,0213
11	3.053.186	367.978,25	279.025.084.098,06	0,93	0,8238	0,8333	0,0095
12	2.493.140	-174.067,75	79.421.244.215,06	-0,50	0,3085	0,3333	0,0248
Σ	32.006.493		3.482.141.555.271,25				
$\bar{X}$	2.667.207,75						

Berdasarkan tabel diatas diperoleh $L_{\max} = 0,0877$ sedangkan L_{tabel} dengan α (0,05) dan n (12) sebesar 0.242, dikarenakan $L_{\max} \leq L_{\text{tabel}}$ maka data penyaluran produksi CPO periode 1 berdistribusi normal. Dengan menggunakan metode yang sama dengan data selanjutnya didapatkan bahwa semua data yang dikumpulkan berdistribusi normal.

4.3. Perhitungan Biaya Total Persediaan Perusahaan

Berdasarkan hasil penelitian data perusahaan dapat diketahui bahwa :

1. Laju produksi CPO per bulan

$$p = \frac{\text{jumlah produksi periode 1} + \text{jumlah produksi periode 2}}{24}$$

$$p = \frac{33.491.103 + 35.339.344}{24}$$

$$p = \frac{68.830.447}{24}$$

$$p = 2.867.935,292 \text{ kg}$$

Maka **rata-rata jumlah produksi CPO** per bulan adalah **2.867.935,292 kg**.

2. Laju penyaluran produksi CPO per bulan

$$d = \frac{\text{jumlah penyaluran periode 1} + \text{jumlah penyaluran periode 2}}{24}$$

$$d = \frac{32.006.493 + 32.843.087}{24}$$

$$d = \frac{64.894.580}{24}$$

$$d = 2.702.065,833 \text{ kg}$$

Maka **rata-rata jumlah penyaluran produksi CPO** per bulan adalah **2.702.065,833kg**.

3. Laju biaya pengadaan produksi CPO per bulan

$$S = \frac{\text{jumlah biaya pengadaan periode 1} + \text{jumlah biaya pengadaan periode 2}}{24}$$

$$S = \frac{52.739.375.825 + 62.194.802.950}{24}$$

$$S = \frac{114.934.178.775}{24}$$

$$S = \text{Rp } 4.788.924.115,625$$

Maka **rata-rata jumlah biaya pengadaan produksi CPO** per bulan adalah

Rp. 4.788.924.115,625.

4. Biaya penyimpanan CPO per kilogram

Biaya penyimpanan per kilogram CPO adalah sebesar 15% dari harga pokoknya yaitu :

$$h = 15\% \times \left(\frac{5.000 + 8.299}{2} \right)$$

$$h = \frac{15}{100} \times 6.649,5$$

$$h = 992,175$$

Maka diperoleh **biaya penyimpanan (h) CPO** perkilogram adalah sebesar **Rp. 992,175**.

5. Biaya Total Persediaan (BTP)

$$BTP = \frac{D}{Q} S + \frac{Q}{2} h \left(1 - \frac{d}{p} \right)$$

$$BTP = \frac{2.702.065,83}{2.867.935,29} (4.788.924.115,62) + \frac{2.867.935,29}{2} .992,175 \left(1 - \frac{2.702.065,83}{2.867.935,29} \right)$$

$$BTP = 287.335.446,96 + 170.729.621,90$$

$$BTP = \text{Rp. } 458.065.068,86$$

Jadi biaya total persediaan per bulan di peroleh sebesar Rp. 458.065.068,86. Maka biaya untuk pengadaan persediaan produksi CPO dalam dua periode sekaligus adalah sebesar :

$$\begin{aligned} BTP \times t &= \text{Rp. } 458.065.068,86 \times 24 \\ &= \text{Rp. } 10.993.561.652,64 \end{aligned}$$

4.4. Perhitungan Biaya Total Persediaan Menggunakan Metode Pengendalian Persediaan

Berikut ini merupakan perhitungan Biaya Total Persediaan menggunakan metode pengendalian persediaan.

1. Tingkat Optimal Produksi (Q)

Berdasarkan data yang disajikan sebelumnya, maka diperoleh nilai dari

- Rata-rata jumlah produksi CPO per bulan
 $p = 2.867.935,292 \text{ kg}$
- Rata-rata jumlah penyaluran produksi CPO per bulan
 $d = 2.702.065,833 \text{ kg}$
- Rata-rata biaya pengadaan produksi CPO per bulan
 $S = \text{Rp. } 4.788.924.115,625$
- Biaya penyimpanan (h) CPO perkilogram
 $h = \text{Rp. } 992,175$

Selanjutnya lakukan perhitungan tingkat produksi optimal CPO (Q) setiap putaran produksi dengan menggunakan rumus:

$$Q = \sqrt{\frac{2.S.d}{h\left(1 - \frac{d}{p}\right)}}$$

$$Q = \sqrt{\frac{2.(4.788.115,625).(2.702.065,833)}{(992,175)\left(1 - \frac{2.702.065,833}{2.867.935,292}\right)}}$$

$$Q = \sqrt{\frac{2.588.047.948.378.816,7213}{9,5305}}$$

$$Q = 65.935.058,42$$

Dari perhitungan diatas diperoleh tingkat produksi CPO optimal adalah sebanyak 65.935.058,42kg. Sehingga produksi CPO optimal perbulannya adalah:

$$\frac{65.935.058,42}{24} = 2.747.294,1 \text{ kg}$$

2. Interval Waktu Optimal Untuk Tiap Putaran Produksi

$$t_0 = \sqrt{\frac{2.S}{h.d\left(1 - \frac{d}{p}\right)}}$$

$$t_0 = \sqrt{\frac{2.(4.788.115,625)}{(992,175).(2.702.065,833)\left(1 - \frac{2.702.065,833}{2.867.935,292}\right)}}$$

$$t_0 = \sqrt{562,987}$$

$$t_0 = 23,72 \approx 24$$

Maka interval waktu optimal pada setiap putaran produksi adalah 24 bulan.

3. Biaya Total Persediaan Minimum Pada Produksi CPO

$$BTP = \sqrt{2.d.s.h\left(1 - \frac{d}{p}\right)}$$

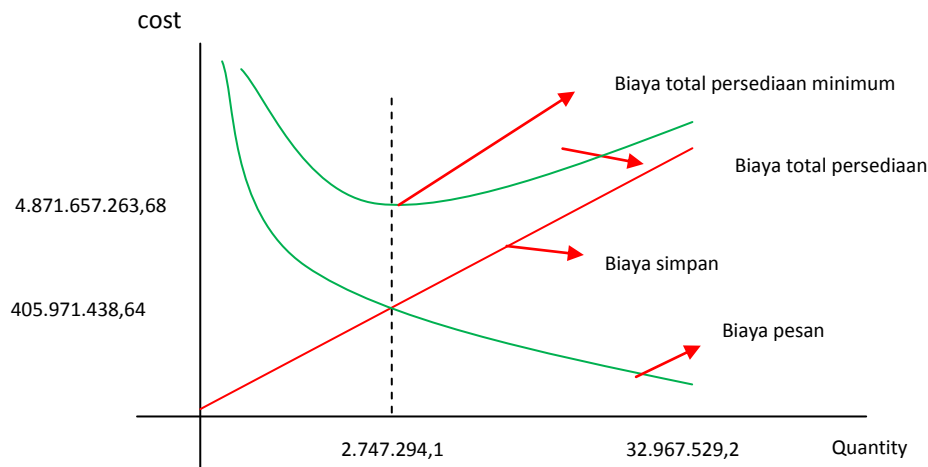
$$BTP = \sqrt{2.(2.702.065,833).(4.788.115,625).(992,175)\left(1 - \frac{2.702.065,833}{2.867.935,292}\right)}}$$

$$BTP = \sqrt{164.812.808.990.456.636,2829}$$

$$BTP = Rp.405.971.438,64$$

Karena BTP yang diperoleh dari hasil perhitungan adalah Rp.405.971.438,64per bulan, sehingga biaya pengadaan persediaan produksi dalam setiap putaran produksi optimalnya adalah :

$$BTP \times t_0 = Rp\ 405.971.438,64 \times 24 = Rp.\ 9.743.314.527,36$$



Gambar4. Grafik *Economic Order Quantity*

4.5. Pembahasan

Pengendalian produksi yang dilaksanakan pada perusahaan yang satu dengan perusahaan yang lain akan berbeda-beda tergantung pada sistem kebijaksanaan manajemen perusahaan yang digunakan. PT. Kresna Duta Agroindo merupakan perusahaan yang beroperasi untuk menghasilkan produk standar sehingga sebagian produk merupakan produk untuk persediaan, oleh karena itu sangat penting apabila perusahaan menerapkan metode pengendalian persediaan dalam merencanakan kegiatan produksinya. Berdasarkan hasil perhitungan yang dilakukan sebelumnya, dirangkum sebagai berikut :

1. Perhitungan yang dilakukan berdasarkan dengan kondisi perusahaan, diperoleh :
 - a. Laju produksi CPO setiap bulannya adalah 2.867.935,292 kg.
 - b. Interval waktu produksi CPO adalah 24 bulan (2 periode).
 - c. Biaya total persediaan CPO setiap bulannya adalah sebesar Rp. 458.065.068,86.
2. Perhitungan yang dilakukan berdasarkan dengan metode pengendalian persediaan, diperoleh :
 - a. Tingkat optimal dari produksi CPO perbulan adalah 2.747.294,1kg.
 - b. Interval waktu optimal produksi adalah 24 bulan.
 - c. Biaya total persediaan minimum CPO perbulan adalah sebesar Rp. 405.971.438,64.
3. Perbandingan perhitungan antara metode pengendalian persediaan dengan cara yang berdasarkan kondisi perusahaan, dapat dilihat pada tabel berikut:

Tabel 6 Perbandingan Perhitungan Berdasarkan Kondisi Perusahaan dan dengan Menggunakan Metode Pengendalian Persediaan

Perhitungan	Perusahaan	Metode EOQ	Selisih	%
Produksi Optimal CPO (Kg)	2.867.935,292	2.747.294,1	120.641,192	4,2
Interval Waktu (Bulan)	24	24	-	-

Biaya Total Persediaan (Rp)	458.065.068,86	405.971.438,64	52.093.630,22	11,4
-----------------------------	----------------	----------------	---------------	------

Dari tabel 6 diatas dapat dilihat bahwa adanya selisih antara perhitungan berdasarkan kondisi perusahaan dengan perhitungan menggunakan metode pengendalian persediaan. PT. Kresna Duta Agroindo belum menerapkan metode pengendalian persediaan, sehingga sering terjadi produksi tidak stabil yang akan berpengaruh terhadap pemenuhan permintaan atau kebutuhan konsumen. Selain itu juga dapat menimbulkan pemborosan biaya persediaan karna adanya persediaan berlebih. Produksi *Crude Palm Oil* (CPO) di PT. Kresna Duta Agroindo rata-rata perbulannya adalah sebesar 2.867.935,292 kgsedangkan permintaan tidak selalu tetap jumlahnya, dengan metode EOQ diperoleh jumlah produksi optimal rata-rata perbulannya adalah sebesar 2.747.294,1 kg, jumlahnya lebih sedikit dibandingkan dengan jumlah produksi yang telah dilakukan oleh PT. Kresna Duta Agroindo akan tetapi masih dapat memenuhi kebutuhan atau permintaan dari konsumen. Produksi sebelumnya perusahaan memiliki jumlah persediaan berlebih yang akan menimbulkan biaya simpan yang besar pula, dengan metode EOQ ini perusahaan dapat mengurangi jumlah produksi sebesar 4,2% yaitu sebesar 120.641,192 kg perbulannya. Apabila PT. Kresna Duta Agroindo menerapkan pengendalian persediaan dengan metode EOQ ini maka perbulannya perusahaan dapat menghemat biaya sebesar Rp. 52.093.630,22 atau 11,4% dari total biaya yang dikeluarkan. Artinya perusahaan sudah tidak melakukan pemborosan lagi dalam produksinya. Metode *Economic Order Quantity* (EOQ) memberi keuntungan bagi perusahaan dimana produksi optimal yang akan diolah telah diketahui, maka perusahaan tidak akan mengalami produksi berlebih yang akan menimbulkan biaya penyimpanan *Crude Palm Oil* (CPO) lebih banyak lagi. Kemudian apabila perusahaan mengalami kekurangan produksi maka permintaan kebutuhan konsumen akan dapat dipenuhi dengan *stock* atau persediaan *Crude Palm Oil* (CPO). Dalam jangka waktu atau periode yang sama yaitu selama 24 bulan terbukti bahwa dengan menggunakan metode pengendalian persediaan *Economic Order Quantity* lebih optimal dan biaya yang dikeluarkan lebih minimum.

Rata-rata produksi optimal *Crude Palm Oil* (CPO) dengan menggunakan metode *Economic Order Quantity* (EOQ) diperoleh sebesar 2.747.294,1 kg per bulannya, sehingga pemenuhan kebutuhan dengan *stock* dapat dilihat pada tabel berikut :

Tabel 7. Stock Tersedia Periode 1

No	Bulan	Tahun (kg)			
		2013	Optimal	Selisih	Stock
1	Januari	1.891.700	2.747.294,1	855.594,1	855.594,1
2	Februari	2.510.220	2.747.294,1	237.074,1	1.092.668,2
3	Maret	2.095.260	2.747.294,1	652.034,1	1.744.702,3
4	April	2.252.430	2.747.294,1	494.864,1	2.239.566,4
5	Mei	2.342.830	2.747.294,1	404.464,1	2.644.030,5
6	Juni	2.700.440	2.747.294,1	46.854,1	2.690.884,6
7	Juli	2.801.625	2.747.294,1	-54.330,9	2.636.553,7
8	Agustus	3.602.309	2.747.294,1	-855.014,9	1.781.538,8
9	September	3.175.871	2.747.294,1	-428.576,9	1.352.961,9
10	Oktober	3.135.482	2.747.294,1	-388.187,9	964.774,0
11	November	3.053.186	2.747.294,1	-305.891,9	658.882,1
12	Desember	2.493.140	2.747.294,1	254.154,1	913.036,2
Jumlah		32.006.493			

Tabel 8. Stock Tersedia Periode 2

No	Bulan	Tahun (kg)			
		2014	Optimal	Selisih	Stock
1	Januari	2.581.794	2.747.294,1	165.500,1	1.078.536,3
2	Februari	2.190.812	2.747.294,1	556.482,1	1.635.018,4
3	Maret	1.850.643	2.747.294,1	896.651,1	2.531.669,5
4	April	2.683.275	2.747.294,1	64.019,1	2.595.688,6
5	Mei	3.127.027	2.747.294,1	-379.732,9	2.215.955,7
6	Juni	2.823.261	2.747.294,1	-75.966,9	2.139.988,8
7	Juli	3.059.254	2.747.294,1	-311.959,9	1.828.028,9
8	Agustus	3.530.498	2.747.294,1	-783.203,9	1.044.825,0
9	September	2.087.209	2.747.294,1	660.085,1	1.704.910,1
10	Oktober	3.191.273	2.747.294,1	-443.978,9	1.260.931,2
11	November	3.035.800	2.747.294,1	-288.505,9	972.425,3
12	Desember	2.682.241	2.747.294,1	65.053,1	1.037.478,4
Jumlah		32.843.087			

Berdasarkan tabel 7 dan 8 diatas kebutuhan atau permintaan *Crude Palm Oil* (CPO) pada periode 1 dan periode 2 tidak tetap atau cenderung naik turun namun dapat terpenuhi dengan jumlah produksi optimal sebesar 2.747.294,1 kg setiap bulannya. Pada akhir periode 1 menyisakan CPO sebanyak 913.036,2 kg, *stock* ini dapat digunakan untuk pemenuhan permintaan pada periode 2, sehingga pada saat permintaan tinggi tetap dapat terpenuhi dengan adanya *stock* yang tersedia.

Di akhir periode 2 PT. Kresna Duta Agroindo memiliki *stock* CPO untuk di gunakan tahun berikutnya sebesar 3.935.867 kg, jumlah tersebut terlalu banyak sehingga dapat menimbulkan biaya simpan untuk menjaga penurunan kualitas produk yang dikeluarkan akan lebih besar lagi. Sedangkan dengan menggunakan metode pengendalian persediaan *Economic Order Quantity* perusahaan dapat menekan atau memperkecil *stock* sebesar 1.037.478,4kg, hal ini dapat menghemat biaya yang harus dikeluarkan perusahaan. Tabel diatas membuktikan bahwa perhitungan menggunakan metode pengendalian persediaan yaitu dengan *Economic Order Quantity* (EOQ), produksi optimal yang harus dikerjakan oleh perusahaan lebih sedikit namun tetap dapat memenuhi permintaan konsumen terhadap *Crude Palm Oil* (CPO) yang tidak tetap setiap bulannya sehingga dapat memberikan keuntungan bagi perusahaan.

5. SIMPULAN

Dari hasil penelitian yang dilaksanakan di PT. Kresna Duta Agroindo dan pengolahan data, pengendalian produksi *Crude Palm Oil* (CPO) menggunakan metode pengendalian persediaan *Economic Order Quantity* (EOQ) dapat disimpulkan bahwa:

1. Tingkat produksi *Crude Palm Oil* (CPO) optimal dalam pengadaanpersediaan sebesar 2.747.294,1 kg perbulan.
2. Interval waktu optimal yang dibutuhkan untuk memproduksi *Crude Palm Oil* (CPO) selama 24 bulan.
3. Biaya total persediaan *Crude Palm Oil* (CPO) minimum setiap bulannya adalah sebesar Rp. 405.971.438,64.
4. Perbandingan perhitungan berdasarkan kondisi perusahaan dengan metode pengendalian persediaan memiliki selisih jumlah produksiCPO 4,2% yaitu sebanyak 120.641,192 kg dan biaya 11,4% sebesar RP. 52.093.630,22.

KEPUSTAKAAN

- [1] Baroto, Teguh. (2002). *Perencanaan dan Pengendalian Produksi*. Jakarta. Ghalia Indonesia.
- [2] Sofyan, Diana Khairani. (2013). *Perencanaan dan Pengendalian Produksi*. Yogyakarta. Graha Ilmu.
- [3] Rangkuti, Freddy. (2004). *Manajemen Persediaan*. Jakarta. Rajawali.
- [4] Assauri, Sofjan. (2004). *Manajemen Produksi dan Operasi*. Jakarta. FEUI
- [5] Heizer, Jay dan Barry Render. (2006). *Operation Management*. Jakarta. Salemba Empat.
- [6] Ginting, Rosnani. (2007). *Sistem Produksi*. Yogyakarta. Graha Ilmu.
- [7] Ghasemi, Naser dan Behrouz A. Nadjafi. (2013). EOQ Models with Varying Holding Cost. *Journal of Industrial Mathematics*. Volume 2013, Article ID 743921.
- [8] Suryono, Hassan. (2009). *Statistika (Pedoman, Teori, dan Aplikasi)*. Jakarta. Erlangga.
- [9] Lumempouw, Veyro E.L ,et. al. (2012). Aplikasi Metode Economic Order Quantity (EOQ) pada Persediaan BBM di PT. Sarana Samudera Pacific Bitung. *Jurnal Poros Teknik Mesin Unsrat*. Vol 1, No 1.
- [10] Indropasto dan Erma Suryani. (2012). Analisis Pengendalian Persediaan Produk Dengan Metode EOQ Menggunakan Metode Algoritma Genetika untuk Mengefisienkan Biaya Persediaan. *Jurnal Teknik ITS*. Volume 1 Nomor 1.
- [11] Limansyah, Taufik. (2011). Analisis Model Persediaan Barang EOQ dengan Mempertimbangkan Faktor Kadaluarsa dan Faktor All Unit Discount. *Research Report Engineering Science*. Volume 1.

Analisis Cluster Data Longitudinal pada Pengelompokan Daerah Berdasarkan Indikator IPM di Jawa Barat

A.S Awalluddin¹, I. Taufik²
UIN SunanGunungDjati Bandung
Email : aasolih@uinsgd.ac.id

ABSTRAK

Analisis cluster dapat digunakan untuk melakukan pengelompokan objek berdasarkan kesamaan karakteristik data yang ada pada setiap objek tersebut. Umumnya analisis cluster terbatas pada struktur data cross section (data silang), dengan asumsi pada satu waktu pengamatan. Penelitian ini bertujuan untuk menentukan pendekatan baru dalam analisis cluster untuk data longitudinal dengan struktur data yang bukan hanya cross section tapi juga time series (rumpun waktu) untuk data multivariabel. Perluasan metode analisis cluster hirarki Ward dengan kombinasi analisis data cross section dan time series dijadikan sebagai dasar analisis matematik dalam penelitian ini. Implementasi metode dilakukan pada studi kasus dataIndeks Pembangunan Manusia (IPM) Jawa Barat berupa empat variabel indikator komponen IPM yaitu : Angka Harapan Hidup (AHH), Harapan Lama Sekolah (HLS), Rata-rata Lama Sekolah (RLS), dan Pengeluaran Perkapita (PP), untuk mengetahui pengelompokan kabupaten/kota, sehingga dapat diketahui daerah mana saja yang perlu mendapatkan prioritas peningkatan nilai variabel indikator IPM dalam kebijakan pembangunan Provinsi.

Kata kunci : Analisis cluster, Analisis Multivariat, Data Longitudinal, Metode Ward

1. PENDAHULUAN

Analisis *cluster* adalah teknik multivariat yang mempunyai tujuan utama untuk mengelompokkan objek penelitian berdasarkan karakteristik yang dimilikinya. Analisis *cluster* mengklasifikasi objek sehingga setiap objek yang paling dekat karakteristiknya dengan objek lain berada dalam *cluster* yang sama. *Cluster* yang baik terbentuk memiliki homogenitas internal dan heterogenitas eksternal yang tinggi. Solusi analisis *cluster* bergantung pada variabel-variabel yang digunakan sebagai dasar untuk menilai kesamaan. Metode analisis *cluster* yang ada umumnya untuk analisis data *cross section*. Kajian analisis *cluster* untuk data longitudinal masih terbatas.

Beberapa literatur telah menjelaskan analisis *cluster* secara umum dan dapat dijadikan sebagai kerangka teoretis yang dapat dijadikan dasar dalam penelitian ini diantaranya ditulis oleh [1] yang lebih mengkaji secara matematis, sedangkan [2] lebih mengkaji secara aplikatif. Landasan teoritis analisis data longitudinal dapat diperoleh dalam tulisan [3] yang membahas secara mendasar konsep data longitudinal, [4] membahas beberapa analisis dengan beragam variasi kondisi dan asumsi metode analisis data longitudinal, dan pembahasan beberapa penelitian mengenai analisis data longitudinal dikumpulkan oleh [5]. [6] menguraikan penggunaan data panel untuk analisis ekonomi.

Metode dasar analisis baik yang berkaitan dengan analisis *cluster* maupun data longitudinal, menjadi dasar pengembangan metode analisis *cluster* untuk data longitudinal yang dijelaskan dalam makalah ini. Beberapa penelitian yang berkaitan dengan analisis *cluster* dan data longitudinal diantaranya telah ditulis oleh [7] yang mengkaji analisis *cluster* untuk data panel variabel tunggal (*single*) dengan menggunakan metode regresi bertahap. [8] melakukan pengkajian pengelompokan (*clustering*) untuk data longitudinal multivariabel kontinyu dan diskrit dengan menggunakan inferensi bayesian pada model dan menggunakan simulasi MCMC (*Markov*

Chain Monte Carlo). Bagaimana dampak pengelompokan pada nilai varians dalam analisis data longitudinal dibahas oleh [9].

Makalah ini menjelaskan bagaimana analisis *cluster* untuk data longitudinal khususnya data multivariabel. Analisis *cluster* data longitudinal lebih kompleks dibandingkan dengan analisis *cluster* data *cross section*, karena bukan hanya dari dimensi objek yang diperhatikan tetapi juga dari dimensi waktu (*time series*). [10] telah mengkaji analisis *cluster* menggunakan perluasan metoda aglomerasi hirarki Ward dengan *sum of squared deviation (SSD)* untuk data longitudinal multivariabel. Makalah ini menawarkan alternatif lain pendekatan perluasan Ward dengan *sum of absolute deviation (SAD)* untuk data longitudinal multivariabel.

Penerapan metoda dilakukan untuk menentukan pengelompokan Kabupaten dan Kota di Provinsi Jawa Barat berdasarkan pada variabel-variabel indikator Indeks Pembangunan Manusia (IPM) untuk tahun 2010 – 2014. Data merupakan data longitudinal multivariabel dengan pengamatan waktu selama empat tahun dan banyaknya objek pengamatan 26 Kabupaten dan Kota dan variabel yang diukur adalah Angka Harapan Hidup (AHH), Harapan Lama Sekolah (HLS), Rata-rata Lama Sekolah (RLS), dan Pengeluaran Perkapita (PP).

Pembahasan dalam makalah ini terbagi dalam beberapa bagian. Bagian 2. Membahas struktur data longitudinal multivariabel untuk memberikan gambaran kompleksitas struktur data jika dibandingkan dengan data *cross section* maupun data longitudinal dengan variabel tunggal (*single*). Bagian 3. Membahas analisis *cluster* secara umum yang menjadi dasar pengembangan metode *cluster* data longitudinal dengan dua pendekatan yaitu perluasan Ward SSD dan alternatif lain perluasan metode Ward SAD untuk data longitudinal multivariabel. Bagian 4. Berisi penerapan metoda untuk kasus pengelompokan daerah Kabupaten dan Kota dengan data longitudinal multivariabel berdasarkan variabel indikator IPM tahun 2010-2014 dengan menganalisis perbandingan hasil pengelompokan yang diperoleh dengan kedua metode.

2. DATA LONGITUDINAL MULTIVARIABEL

Struktur data longitudinal dapat dibedakan menjadi struktur data longitudinal dengan variabel tunggal dan data longitudinal multivariabel. Struktur data longitudinal tunggal dapat dianggap sebagai struktur data dengan tabel dua dimensi, di mana baris menyatakan objek/sampel sedangkan kolom menyatakan variabel tunggal dari waktu ke waktu. Misalkan $X_i(t)$ adalah variabel pengamatan objek ke- i pada saat pengamatan ke- t . Himpunan data longitudinal terdiri dari pengamatan pada objek penelitian ke- i selama pengamatan untuk setiap $i = 1, 2, \dots, N$ dan $t = 1, 2, \dots, T$. Struktur data seperti ini sama dengan struktur data *cross section*, dengan kolom menyatakan p buah variabel, sedangkan untuk struktur longitudinal menyatakan t waktu pengamatan. Desain data longitudinal variabel tunggal dapat dilihat dalam Tabel 1. Analisis *cluster* untuk data longitudinal tunggal dapat menggunakan analisis *cluster* data *cross section*. Pengelompokan untuk data longitudinal variabel tunggal mudah dan dapat menggunakan *software* yang sudah tersedia seperti SPSS, MINITAB, dan lain sebagainya.

Tabel 1. Desain Data Longitudinal Variabel Tunggal

waktu(t)	1	2	...	T
objek(i)				
1	$X_{1(1)}$	$X_{1(2)}$	...	$X_{1(T)}$
2	$X_{2(1)}$			$X_{2(T)}$
.	.	.	.	.
.				
N	$X_{N(1)}$	.	.	$X_{N(T)}$

Struktur data longitudinal multivariabel lebih kompleks sehingga tidak dapat dibuat dalam bentuk tabel dua dimensi. Misalkan $X_{ij}(t)$ adalah variabel pengamatan untuk objek ke-i, variabel ke-j, dan saat pengamatan ke-t. Himpunan data longitudinal terdiri dari pengamatan pada objek penelitian ke-i, variabel ke-j selama pengamatan untuk setiap $i = 1, 2, \dots, N$, $j = 1, 2, \dots, p$ dan $t = 1, 2, \dots, T$. Desain struktur data longitudinal multivariabel dapat dilihat dalam Tabel 2.

Tabel 2. Desain Data Longitudinal Multivariabel

waktu (t)	1	...	t	...	T
variabel (j)	1 ... j ... p		1 ... j ... p		1 ... j ... p
objek(i)					
1	$X_{11(1)}, \dots, X_{12(1)}, \dots, X_{1p(1)}$	...	$\dots X_{1j(t)}, \dots$	...	$\dots X_{1p(T)}$
2	$X_{21(1)}, \dots, X_{22(1)}, \dots, X_{2p(1)}$	...			$\dots X_{2p(T)}$
.	.		.	.	.
.					
.					
N	$X_{N1(1)}, \dots, X_{N2(1)}, \dots, X_{Np(1)}$	...	$\dots X_{Nj(t)}, \dots$	.	$\dots X_{Np(T)}$

Untuk menentukan statistik rata-rata dan varians dari data longitudinal multivariabel, dapat ditentukan berturut turut sebagai berikut :

Rata-rata dari variabel ke-j pada waktu ke-t

$$\bar{X}_j(t) = \frac{1}{N} \sum_{i=1}^N X_{ij}(t) \quad (1)$$

Rata-rata dari variabel ke-j

$$\bar{X}_j(t) = \frac{1}{TN} \sum_{t=1}^T \sum_{i=1}^N X_{ij}(t) \quad (2)$$

Varians dari variabel ke-j pada waktu ke-t

$$Var_{x_j(t)} = \frac{1}{N-1} \sum_{i=1}^N [X_{ij}(t) - \bar{X}_j(t)]^2 \quad (3)$$

Varians dari variabel ke-j

$$Var_{x_j} = \frac{1}{T(N-1)} \sum_{i=1}^N [X_{ij}(t) - \bar{X}_j(t)]^2 \quad (4)$$

Statistik ini akan digunakan dalam analisis cluster.

3. ANALISIS CLUSTER DATA LONGITUDINAL

Dalam menentukan analisis *cluster* longitudinal, dasar pemahaman dibangun berdasarkan pada analisis *cluster* secara umum untuk data *cross section*. Dua hal utama yang perlu ditetapkan dalam analisis cluster. *Pertama*, teknik apa yang akan digunakan, apakah teknik non-hirarki atau hirarki. *Kedua*, metode pengelompokan apa yang akan dipilih, jika teknik non-hirarki yang dipilih maka metode yang dapat dipilih adalah metode partisi (*partitioning*) seperti *K-mean*, metode campuran distribusi (*mixture of distribution*), dan *density estimation*. Selain itu jika teknik hirarki yang dipilih maka metode yang dapat digunakan diantaranya : metode jarak terdekat (*single linkage*), metode jarak terjauh (*complete linkage*), metode rata-rata, metode centroid, metode jumlah kuadrat error / *sum of squares deviation* (Ward), dan lain sebagainya.

[10] menggunakan teknik aglomerasi hirarki dengan metode perluasan Ward dalam analisis *cluster* data longitudinal multivariabel yang juga digunakan sebagai salah satu metode analisis dalam makalah ini. Pengelompokan yang optimal dapat terbentuk dengan diperoleh nilai observasi yang homogen antar anggota di dalam *cluster*, tetapi berbeda secara nyata antar *cluster*. Teknik pengelompokan ini dimulai dengan n cluster dengan menganggap bahwa setiap objek/sampel adalah sebuah *cluster*. Selanjutnya pengukuran jarak antar objek dilakukan untuk mengelompokkan sesuai dengan kedekatan antar objek atau kelompok objek sebagai *cluster* baru. Proses ini terus dilakukan sampai terbentuk *cluster* yang lebih kecil. Penentuan jarak antara *cluster* baru yang terbentuk dengan sisa *cluster* ditentukan dengan metode pautan pengelompokan (*linkage method*).

Makalah ini memilih metode pautan pengelompokan Ward, metode ini tidak hanya mempertimbangkan jarak antar cluster, tetapi juga di dalam *cluster*. Berbeda dengan metode lain yang hanya mempertimbangkan jarak antar *cluster* [1]. Metode Ward juga dikenal sebagai metode jumlah kuadrat deviasi (*sum of squares deviation*). Fungsi jumlah kuadrat deviasi (SSD) untuk data longitudinal didefinisikan dalam rumus (5) berikut :

$$S_h = \sum_{t=1}^T \sum_{j=1}^p \sum_{i \in i^h} [X_{ij}(t) - \bar{X}_j^h]^2 \quad (5)$$

dengan S_h menyatakan jumlah kuadrat deviasi dari *cluster* h , dan $\bar{X}_j^h$ adalah rata-rata variabel ke-j pada waktu t untuk *cluster* h yang dapat diperoleh dengan menggunakan rumus (1) , dan i^h menyatakan seluruh objek yang termasuk *cluster* h . Perluasan metode Ward dengan menggunakan *least absolute deviation* telah diteliti oleh [11]

untuk jenis data *cross section*. Perluasan Ward dengan *sum of absolute deviation* (SAD) untuk data longitudinal yang digunakan dalam makalah ini didefinisikan dalam rumus (6) berikut :

$$M_h = \sum_{t=1}^T \sum_{j=1}^p \sum_{i \in h} |X_{ij}(t) - \bar{X}_j^h| \quad (6)$$

dengan M_h menyatakan jumlah absolut deviasi dari *cluster* h .

Pemilihan *cluster* baru dilakukan dengan teknik aglomerasi hirarki Ward dapat dilakukan dengan formula Lance-Williams [1] sebagai berikut :

$$D_{(C,AB)} = \frac{n_A + n_C}{n_A + n_B + n_C} S_{AC} + \frac{n_B + n_C}{n_A + n_B + n_C} S_{BC} - \frac{n_C}{n_A + n_B + n_C} S_{AB} \quad (7)$$

dengan n_A , n_B , n_C berturut-turut adalah banyaknya objek/sampel pada *cluster* A, B, dan C.

Menentukan nilai S_{AC} , S_{BC} , dan S_{AB} diperoleh dengan rumus (5). Formula yang sama juga diterapkan ketika menggunakan M_h dengan rumus (6). Pembentukan *cluster* data longitudinal multivariabel dapat dilakukan dengan langkah berikut :

1. Tentukan data $X_{ij}(t)$; $i=1,2,...,N$, $j=1,2,...,p$, $t=1,2,..,T$
2. Mulai dengan N *cluster*
3. Hitung persamaan (5) dan (6)
4. Pilih min $\{S_h\}$ dan min $\{M_h\}$ sebagai *cluster* $N-1$
5. Hitung $D_{(C, AB)}$ pada persamaan (7)
6. Ulangi langkah 3 s.d 5 untuk *cluster* $N-2$, $N-3$, dan seterusnya
7. Sampai terbentuk 1 *cluster* untuk setiap metode

4. ANALISIS CLUSTER DATA IPM JAWA BARAT

Human Development Index (HDI) atau sering dikenal Indeks Pembangunan Manusia (IPM) merupakan indikator yang dapat menjelaskan perkembangan pembangunan manusia secara representatif. Tiga dimensi dasar yang menjadi tolak ukur besar kecilnya nilai IPM di suatu wilayah adalah harapan hidup yang tinggi, pendidikan yang cukup dan standar hidup yang layak. Dalam perkembangannya indikator dan perhitungan IPM mengalami perubahan sejak tahun 2010, perbedaan yang cukup signifikan ini menjadikan perhatian pemerintah daerah untuk memperhatikan variabel indikator dalam rangka meningkatkan nilai IPM. Tidak hanya itu perhatian pengelompokan Kabupaten dan Kota berdasarkan capaian indikator IPM masing-masing menjadi perhatian untuk menentukan prioritas pembangunan daerah.

Penerapan metode analisis *cluster* data longitudinal yang telah dijelaskan dalam pembahasan sebelumnya dilakukan pada data IPM Kabupaten dan Kota di Provinsi Jawa Barat dengan metode perhitungan baru. Data diambil dari 26 Kabupaten dan Kota di Jawa Barat (tidak termasuk Kab. Pangandaran) selama periode 2010-2014 bersumber dari BPS Jawa Barat [12]. Empat variabel indikator IPM digunakan yaitu : Angka Harapan Hidup (AHH), Harapan Lama Sekolah (HLS), Rata-rata Lama Sekolah (RLS), dan Pengeluaran Perkapita (PP). Perhitungan metode baru digunakan dalam menentukan nilai dari setiap variabel. Secara desriptif statistik data yang digunakan dapat dilihat dalam Tabel 3 dan rata-rata IPM Kab/Kota dapat dilihat dalam Tabel 4.

Tabel 3. Statistik Data Variabel Indikator IPM

Variabel	Mean	Var	Max	Min
AHH	71,29	1,67	74,18	67,54
HLS	11,69	0,95	13,71	9,62
RLS	7,72	1,52	10,78	4,93
PP	9344,96	2178,87	15048,47	6149,57

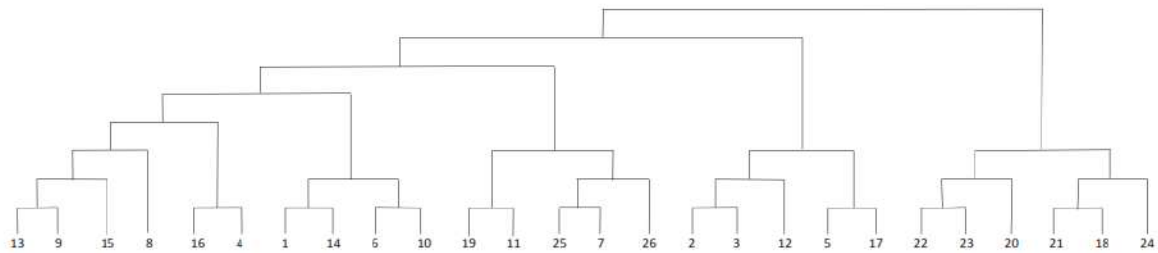
Tabel 4. Rata-rata IPM Kab/Kota Jawa Barat

No.	Kabupaten	Rata-rata	No.	Kabupaten/Kota	Rata-rata
1	Bogor	65,78	14	Purwakarta	66,23
2	Sukabumi	62,36	15	Karawang	65,89
3	Cianjur	60,40	16	Bekasi	69,24
4	Bandung	68,17	17	Bandung Barat	63,02
5	Garut	61,14	18	Kota Bogor	72,24
6	Tasikmalaya	61,63	19	Kota Sukabumi	69,67
7	Ciamis	66,25	20	Kota Bandung	78,29
8	Kuningan	65,57	21	Kota Cirebon	71,88
9	Cirebon	64,57	22	Kota Bekasi	77,89
10	Majalengka	63,18	23	Kota Depok	77,55
11	Sumedang	67,36	24	Kota Cimahi	75,02
12	Indramayu	62,19	25	Kota Tasikmalaya	67,86
13	Subang	64,78	26	Kota Banjar	67,57

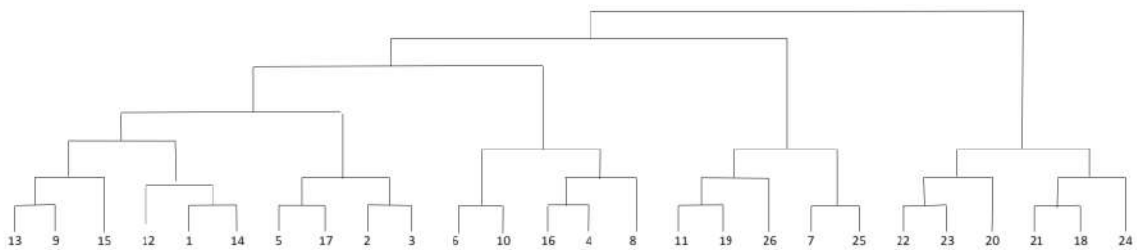
Penentuan *cluster* dengan metode yang telah dijelaskan dilakukan dengan dua pendekatan yaitu penerapan analisis *cluster* data longitudinal dengan perluasan Ward SSD rumus (7) dan (9) dan perluasan Ward SAD rumus (8) dan (9). Hasil pengelompokan ditunjukkan berturut-turut dalam dendrogram **Gambar 1** dan **Gambar 2**.

Secara umum kedua metode menghasilkan lima pengelompokan (*cluster*) besar dengan dua *cluster* memiliki objek yang sama. *Cluster* pertama yaitu : Kota Bekasi, Kota Depok, Kota Bandung, Kota Cirebon, Kota Bogor, dan Kota Cimahi. *Cluster* kedua yaitu : Kota Sukabumi, Kab. Sumedang, Kota Tasikmalaya, Kota Cimahi, dan Kota Banjar. Kedua *cluster* secara berturut-turut menunjukkan bahwa dari keempat indikator IPM kedua cluster di atas merupakan daerah yang paling baik dan baik jika dilihat dari keempat variabel indikator IPM. Tiga

cluster lainnya dari kedua metode memberikan hasil yang sedikit berbeda. Khusus untuk *cluster* dengan objek daerah Kab. Sukabumi, Kab. Cianjur, Kab. Garut, dan Kab. Bandung Barat, hasil pengelompokan metode perluasan Ward SAD, ditambah Kab. Indramayu yang diperoleh metode perluasan Ward SSD menunjukkan *cluster* yang paling rendah, sehingga kelima daerah tersebut dapat dikategorikan sebagai daerah prioritas dalam pembangunan di Provinsi Jawa Barat untuk meningkatkan nilai keempat variabel indikator IPM.



Gambar 1. Dendrogram Perluasan Ward SSD



Gambar 2. Dendrogram Perluasan Ward SAD

5. SIMPULAN

Analisis *cluster* dengan data longitudinal dengan memperhatikan dimensi waktu dan objek berbeda pendekatannya dengan analisis umum yang digunakan karena mengabaikan pengaruh perubahan waktu terhadap perubahan nilai variabel. Alternatif analisis *cluster* untuk data longitudinal telah diuraikan dalam makalah ini dengan dua pendekatan alternatif yaitu jumlah kuadrat deviasi (SSD) dan jumlah absolut deviasi (SAD) untuk data longitudinal dengan perluasan Ward sebagai *linkage method*.

Pengelompokan studi kasus pada data empat variabel indikator IPM untuk 26 Kabupaten dan Kota di Jawa Barat dari 5 *cluster* daerah yang terbentuk, memberikan hasil yang sama untuk tiga *cluster*. Daerah dengan *cluster* sangat baik berdasarkan empat variabel indikator IPM yaitu Kota Bekasi, Kota Depok, Kota Bandung, Kota Cirebon, Kota Bogor, dan Kota Cimahi. Sedangkan *cluster* rendah yang perlu mendapat perhatian pemerintah baik Provinsi maupun Pusat untuk menjadikan prioritas dalam meningkatkan variabel indikator IPM yaitu : Kab. Sukabumi, Kab. Cianjur, Kab. Garut, Kab. Bandung Barat, dan Kab. Indramayu

6. UCAPAN TERIMA KASIH

Penelitian ini dibiayai oleh dana BOPTN UIN Sunan Gunung Djati. Kami mengucapkan terima kasih atas dukungan yang diberikan.

KEPUSTAKAAN

- [1] Rencher, A.C., 2002. *Methods of Multivariate Analysis*, 2nd Edition, New York; John Wiley & Sons.
- [2] Hair, J.E., dkk., 1998. *Multivariate Data Analysis fifth Ed.*, Prentice Hall International.
- [3] Toon W. Taris, 2000. *A Primer in Longitudinal Data Analysis.*, Sage Publications Ltd.
- [4] Frees, Edward. W, 2004. *Longitudinal and Longitudinal Data : Analysis and Applications in Social Science.*, Cambridge University Press.
- [5] Hsio, Cheng, dkk., 2004. *Analysis of Longitudinals and Limited Dependent Variable Models.*, Cambridge University Press.
- [6] Baltagi, Badi H., 2005. *Econometric Analysis of Panel Data third edition*, John Wiley and Sons.
- [7] Mouchart, Michel, Jeroen V.K., 2005. *Clustered panel data Clustered Panel Data Models: An Efficient Approach for Nowcasting from Poor Data.* International Journal of Forecasting 5:577-594.
- [8] Komarek, A., Kamarkova, L., 2013. *Clustering for multivariate continuous and discrete longitudinal data.* The Annals of Applied Statistics, 7(1), 177–200.
- [9] C. Skinner., M.T. Vieire., 2007. *Variance Estimation in the Analysis of Clustered Longitudinal Survey Data.*, Survey Methodology Vol.33, No.1, pp.3-12.
- [10] Zheng, B., Li, S., 2014. *Multivariable Panel Data Cluster Analysis and Its Application.*, Computer Modelling & New Technologies 18., 553-557.
- [11] Strauss, T., J Von Maltitz, M., 2014. *Statistical Classification of Languages : Generalising Method for Use with Manhattan Distance.*, Technical Report., University of The Free State.
- [12] Badan Pusat Statistika Provinsi Jawa Barat. *Indeks Pembangunan Manusia Metode Baru Provinsi Jawa Barat dan Kabupaten/Kota Tahun 2010-2014.* (Online). (<http://jabar.bps.go.id/linkTabelStatis/view/id/95>)

Indek Pembangunan Manusia dan Faktor Yang Mempengaruhinya di Daerah Perkotaan Provinsi Lampung

Ahmad Rifa'i¹ dan Hartono²

¹Jurusan Ilmu Adm. Bisnis FISIP Unila, Jln Sumatri Brojonegoro No. 01 Bandar Lampung,
HP. 0812 79 481 545, Email: rifaiunila@gmail.com

²Jurusan Ilmu Adm. Bisnis FISIP Unila, Jln Sumatri Brojonegoro No. 01 Bandar Lampung

ABSTRAK

Tujuan penelitian ini adalah menghitung pengaruh Angka Harapan Hidup, Rata-Rata Lama Sekolah, Banyaknya Tenaga Kesehatan, Persentase Penduduk Miskin, Pendapatan Per Kapita, dan Pertumbuhan Ekonomi terhadap Indek Pembangunan Manusia (IPM) Kota Bandar Lampung dan Kota Metro di Provinsi Lampung Tahun 2002 – 2015. Data yang kami gunakan adalah data sekunder. Teknik analisis data menggunakan regresi linear berganda. Kami juga menggunakan uji asumsi klasik, yaitu uji gejala multikolineritas, heteroskedastisitas, dan autokolerasi untuk memperoleh pemeriksa hasil estimasi yang bersifat BLUE (Best Linear Unbiased Estimator). Hasil penelitian menunjukkan pemeriksa hasil estimasi terbebas dari gejala asumsi klasik sehingga model yang diperoleh memenuhi sifat Best Linear Unbiased Estimator. Secara parsial Angka Harapan Hidup berpengaruh signifikan terhadap IPM dengan tingkat signifikansi 99%. Sedangkan Rata-Rata Lama Sekolah, Banyaknya Tenaga Kesehatan, Persentase Penduduk Miskin, Pendapatan Per Kapita, dan Pertumbuhan Ekonomi tidak berpengaruh signifikan terhadap IPM, karena tingkat signifikansi yang diperoleh dibawah 95%. Secara simultan seluruh faktor berpengaruh signifikan terhadap IPM dengan tingkat signifikansi 99% dan tingkat keeratan hubungan antara variabel dependent dengan variabel independent-nya (koefisien determinasi – R^2) mencapai 67,05%.

Keywords: Human Developmen Index, Life Expectacy at Birth, Mean Years Schooling, Health Personal, Poverty, GDP per Capita, dan Growth.

I. PENDAHULUAN

Pembangunan manusia adalah suatu proses untuk memperbanyak pilihan-pilihan yang dimiliki oleh manusia. Diantara banyak pilihan tersebut, pilihan yang terpenting adalah untuk berumur panjang dan sehat, untuk berilmu pengetahuan, dan untuk mempunyai akses terhadap sumber daya yang dibutuhkan agar dapat hidup secara layak [1]. Dalam [2] dinyatakan bahwa salah satu alat ukur yang dianggap dapat merefleksikan status pembangunan manusia adalah *Human Development Index* (HDI) atau Indek Pembangunan Manusia (IPM). Indeks Pembangunan Manusia (IPM) mengukur capaian pembangunan manusia berbasis sejumlah komponen dasar kualitas hidup.

IPM merupakan suatu indeks komposit yang mencakup tiga bidang pembangunan manusia yang di anggap sangat mendasar yaitu usia hidup (*longevity*) dan sehat, pengetahuan (*knowledge*), dan standar hidup layak (*decent living*). Untuk mengukur dimensi kesehatan, digunakan angka harapan hidup waktu lahir. Selanjutnya untuk mengukur dimensi pengetahuan digunakan gabungan indikator angka melek huruf dan rata-rata lama sekolah. Adapun untuk mengukur dimensi hidup layak digunakan indikator kemampuan daya beli masyarakat terhadap sejumlah kebutuhan pokok yang dilihat dari rata-rata besarnya pengeluaran per kapita sebagai pendekatan pendapatan yang mewakili capaian pembangunan untuk hidup layak.

[1] dinyatakan bahwa komponen Komponen Indeks Pembangunan Manusia (IPM) terdiri dari empat unsur, yaitu angka harapan hidup, angka melek huruf, rata-rata lama sekolah, dan pengeluaran riil per kapita yang

disesuaikan. Beberapa penelitian tentang faktor yang mempengaruhi IPM diantaranya [3], [4], [5], [6], [7] menunjukkan hasil yang bervariasi. Misalnya untuk faktor kesehatan dalam [6] dan [7] berpengaruh signifikan terhadap IPM, sedangkan dalam [5] dan [1] tidak berpengaruh signifikan. Berdasarkan penelitian yang dilakukan oleh beberapa peneliti tersebut yang dilakukan diantaranya di Jawa Timur, Kota Semarang, Provinsi Bali, dan Provinsi Papua terdapat beberapa faktor yang mempengaruhi IPM, yaitu angka partisipasi sekolah (APS), kesehatan, PDRB per kapita, pendidikan, pengeluaran pemerintah fungsi kesehatan dan/atau fungsi pendidikan, pengeluaran rumah tangga untuk makanan dan/atau pendidikan, rasio jumlah penduduk terhadap jumlah bidan/jumlah dokter/jumlah perawat, rasio kemiskinan terhadap jumlah penduduk, rasio murid SMA terhadap guru, rumah tangga dengan akses air bersih, dan tingkat partisipasi angkatan kerja (TPAK). Berdasarkan fakta tersebut penelitian ini ingin mengetahui faktor apa saja yang mempengaruhi capaian IPM di Kota Bandar Lampung dan Kota Metro Lampung Tahun 2002 – 2015.

II. METODE PENELITIAN

Jenis penelitian ini adalah *explanatory research* yaitu penelitian yang menggunakan pengujian hipotesis. Penelitian eksplanatori atau penelitian penjelasan menganalisis hubungan antar variabel penelitian dan menguji hipotesis yang telah dirumuskan. Objek penelitian ini adalah Indeks Pembangunan Manusia (IPM), Angka Harapan Hidup (AHH), Rata-Rata Lama Sekolah (MYS), Banyaknya Tenaga Kesehatan (*Health Personal*/HP), Persentase Penduduk Miskin (POV), Pendapatan Perkapita (PC), dan Pertumbuhan Ekonomi (GRW). Data yang digunakan adalah data sekunder yang diambil dari publikasi Badan Pusat Statistik (BPS) Kota Bandar Lampung, Kota Metro, dan Badan Pusat Statistik (BPS) Provinsi Lampung. Teknik analisis data menggunakan regresi linear berganda. Model yang digunakan dalam penelitian ini adalah:

$$IPM_{i,t} = \alpha_{0i,t} + \alpha_1 AHH_{i,t} + \alpha_2 GRW_{i,t} + \alpha_3 MYS_{i,t} + \alpha_4 PC_{i,t} + \alpha_5 HP_{i,t} + \alpha_6 POV_{i,t} + \varepsilon_t \quad (1)$$

dimana:

IPM = *Indek Pembangunan Manusia*, Ukuran capaian pembangunan manusia yang berbasis komponen dasar kualitas hidup

AHH = *Angka Harapan Hidup*, Rata-rata perkiraan banyak tahun yang dapat ditempuh oleh seseorang selama hidup

GRW = *Growth*, Pertumbuhan ekonomi tahunan.

MYS = *Rata-Rata Lama Sekolah*, Jumlah tahun yang digunakan oleh penduduk usia 15 tahun keatas dalam menjalani pendidikan formal

PC = *Per Capita*, Pendapatan per kapita masyarakat atas dasar harga konstan 2000

HP = *Banyaknya Tenaga Kesehatan*, Jumlah tenaga kesehatan yang ada dan dimiliki oleh daerah observasi

POV = *Persentase Penduduk Miskin*, Persentase penduduk yang berada dibawah Garis Kemiskinan (GK)

ε = Error term; i = Kab/Kota; t = tahun ke- t ; α_0 = konstanta; $\alpha_1, \alpha_2, \alpha_3, \alpha_4, \alpha_5, \alpha_6$ = koefisien regresi.

Dalam penelitian ini juga dilakukan pengujian ketepatan asumsi model melalui uji Asumsi Klasik yaitu uji gejala Multikolinearitas, gejala Heteroskedastisitas, dan gejala Autokorelasi. Tujuannya adalah untuk memperoleh pemeriksa hasil estimasi yang bersifat *BLUE (Best Linear Unbiased Estimator)*. Pengujian gejala Multikolinearitas dilakukan menggunakan matrik korelasi dimana jika masing-masing variabel bebas berkorelasi lebih dari 80%, maka termasuk yang memiliki hubungan yang tinggi dan ada indikasi adanya multikolinearitas. Pengujian gejala Heteroskedastisitas dilakukan menggunakan metode *Breusch-Pagan-Godfrey Heteroscedasticity Test*. Langkah pengujiannya adalah jika diketahui sebuah persamaan [8]:

$$Y_i = \beta_1 + \beta_2 X_{2i} + \beta_3 X_{3i} + u_i \quad (2)$$

Maka langkah tes *Breusch-Pagan-Godfrey Heteroscedasticity* adalah:

1. Regresikan persamaan awal (persamaan 2), dapatkan nilai residunya, $\hat{u}_i$.
2. Lakukan *auxiliary regression* dengan menggunakan nilai $\hat{u}_i$ sebagai variabel *dependentnya*, yaitu:

$$\hat{u}_i^2 = \alpha_1 + \alpha_2 X_{2i} + \alpha_3 X_{3i} + \alpha_4 X_{2i}^2 + \alpha_5 X_{3i}^2 + \alpha_6 X_{2i} X_{3i} + v_i \quad (3)$$

3. Ambil nilai R^2 -nya untuk dilakukan pengujian hipotesis.

Hipotesisnya adalah, H_0 : tidak ada heteroskedastisitas

H_a : ada heteroskedastisitas

dengan persamaan:

$$n \cdot R^2 \sim \chi^2_{df} \quad (4)$$

Jika nilai χ^2 persamamaan (5.7) $> \chi^2_{tabel}$, maka tolak H_0

Jika nilai χ^2 persamamaan (5.7) $< \chi^2_{tabel}$, maka terima H_0

Pengujian gejala Autokorelasi dilakukan dengan metode *Breusch-Godfrey (BG) test* [8]. Metode ini mengasumsikan bahwa selain adanya *first order autoregressive*, dalam sebuah estimasi bisa terjadi *second, third, ..., n autoregressive*. Adapun langkah dalam metode ini adalah jika diketahui sebuah persamaan model Logit sebagai berikut:

$$Y_i = \alpha_1 + \alpha_2 X_{it} + u_t \quad (5)$$

Jika diasumsikan *error term* u_t mengikuti *autoregressive order* ke- p , $AR(p)$ sehingga persamaannya menjadi:

$$u_t = p_1 u_{t-1} + p_2 u_{t-2} + \dots + p_p u_{t-p} + e_t \quad (6)$$

Hipotesis nol-nya adalah $H_0 = p_1 = p_2 = \dots = p_p = 0$; tidak ada autokorelasi pada setiap order.

Jika nilai: $(n-p) \chi^2 > \chi^2 (df=p)$, maka tolak H_0 yang berarti ada autokorelasi

Jika nilai: $(n-p) \chi^2 < \chi^2 (df=p)$, maka terima H_0 .

Pengujian parameter regresi dilakukan melalui uji pengaruh parsial (Uji-t), uji pengaruh secara simultan (Uji-F), dan pengujian koefisien determinasi (R^2). Untuk menguji tingkat signifikansi masing-masing variabel secara parsial dilakukan dengan uji-t. Nilai t -statistik dapat di cari dengan menggunakan rumus:

$$t_i = \frac{\bar{\beta}_i}{S_i} \quad (7)$$

dimana hipotesisnya adalah: $H_0 : \beta_i = 0$, dimana $i = 1, 2, 3 \dots n$

$$H_a : \beta_i \neq 0$$

Dengan menggunakan tingkat derajat kepercayaan minimal 95% ($\alpha=5\%$) dan nilai *degree of freedom* adalah $df=n-k-1$, dimana n =jumlah sampel, k =jumlah variabel bebas, maka kriteria uji yang digunakan adalah:

- (a). Jika nilai t -statistik $>$ t -tabel, maka H_0 ditolak, H_a diterima (signifikan)
- (b). Jika nilai t -statistik $<$ t -tabel, maka H_0 diterima, H_a ditolak (tidak signifikan)

Selanjutnya untuk menguji tingkat signifikansi seluruh variabel secara bersama-sama (simultan) dilakukan dengan menggunakan uji-F. Nilai F -statistiknya dihitung dengan rumus:

$$F = \frac{R^2 / (k - 1)}{(1 - R^2) / (N - k)} \quad (8)$$

dimana hipotesisnya adalah: $H_0 : \beta_1 = \beta_2 = \beta_3 = \dots = \beta_i = 0$

$$H_a : \beta_i \neq \beta_j \neq 0$$

Dengan menggunakan tingkat derajat kepercayaan minimal 95% ($\alpha=5\%$) dan nilai *degree of freedom* adalah $df_1=n-k$ dan $df_2=k-1$, maka kriteria uji yang digunakan adalah:

- (a). Jika nilai F -statistik $>$ F -tabel, maka H_0 ditolak, H_a diterima (signifikan)
- (b). Jika nilai F -statistik $<$ F -tabel, maka H_0 diterima, H_a ditolak (tidak signifikan).

Sedangkan untuk menguji kekuatan hubungan variabel *dependent* dengan variabel *independent* digunakan uji koefisien determinansi (R^2). Koefisien determinansi (R^2) adalah angka yang menunjukkan besarnya proporsi atau persentase variasi variabel *dependent* yang dijelaskan oleh variabel *independent* dalam sebuah model regresi, atau besarnya kemampuan varian/ penyebaran variabel bebas yang menerangkan variabel tidak bebas. Nilai R^2 adalah $0 < R^2 < 1$, yang berarti semakin mendekati angka 1 berarti model tersebut semakin baik karena semakin dekat hubungan antara variabel bebas dengan variabel tidak bebas. Dengan kata lain semakin mendekati angka 1 variabel tak bebas hampir seluruhnya dipengaruhi oleh variabel bebas, dan sebaliknya.

III. HASIL DAN PEMBAHASAN

3.1 Analisis Statistik Hasil Penelitian

Bagian ini menyajikan hasil estimasi persamaan (1). Variabel dependen dalam penelitian ini adalah Indeks Pembangunan Manusia (IPM), sedangkan variabel independennya adalah Angka Harapan Hidup (AHH), Pertumbuhan Ekonomi (GRW), Banyaknya Tenaga Kesehatan (*Health Personal/HP*), Rata-Rata Lama Sekolah (MYS), Pendapatan Perkapita (PC), dan Persentase Penduduk Miskin (POV). Hasil estimasi terhadap persamaan (1) seperti terlihat pada Tabel 1.

Tabel 1. Hasil Estimasi Terhadap Persamaan (1)

Dependent Variable: IPM
Method: Least Squares
Date: 10/21/16 Time: 08:43
Sample: 1 28
Included observations: 28

Variable	Coefficient	Std. Error	t-Statistic	Prob. Ket.
C	-3.371255	11.84504	-0.284613	0.7787 ^{ns}
AHH	0.887240	0.162010	5.476462	0.0000 ^{***}
GRW	0.039669	0.195882	0.202517	0.8415 ^{ns}
HP	-0.000179	0.000535	-0.334719	0.7412 ^{ns}
MYS	1.611902	0.835970	1.928180	0.0675 [*]
PC	-3.91E-08	4.51E-08	-0.865854	0.3964 ^{ns}
POV	-0.072288	0.066526	-1.086607	0.2895 ^{ns}
R-squared	0.743720	Mean dependent var		74.36857
Adjusted R-squared	0.670498	S.D. dependent var		1.741784
S.E. of regression	0.999824	Akaike info criterion		3.049843
Sum squared resid	20.99261	Schwarz criterion		3.382894
Log likelihood	-35.69780	F-statistic		10.15696
Durbin-Watson stat	0.875135	Prob(F-statistic)		0.000026

Sumber: Hasil estimasi dengan Program EVIEWS 3.0

* Significance at $\alpha = 10\%$

** Significance at $\alpha = 5\%$

*** Significance at $\alpha = 1\%$

^{ns} Not Significance

Jika dituliskan dalam bentuk persamaan regresi maka hasil estimasi terhadap persamaan (1) dapat dituliskan sebagai berikut:

$$\begin{aligned}
 \text{IPM} = & -3.371255 + 0.887240 \text{ AHH} + 0.039669 \text{ GRW} - 0.000179 \text{ HP} + 1.611902 \text{ MYS} \\
 & (-0.284613) \quad (5.476462) \quad (0.202517) \quad (-0.334719) \quad (1.928180) \\
 & - 3.91 \times 10^{-08} \text{ PC} - 0.072288 \text{ POV} \\
 & (-0.865854) \quad (-1.086607)
 \end{aligned} \tag{9}$$

3.1.1 Pengujian Ketepatan Asumsi Model

Pengujian ketepatan asumsi model meliputi pengujian terhadap kemungkinan munculnya gejala Multikolinearitas, Heteroskedastisitas, dan Autokorelasi hasil estimasi pada persamaan (9). Pengujian ini dimaksudkan untuk memperoleh pemeriksa (hasil estimasi) yang memiliki sifat *BLUE* (*Best Linear Unbiased Estimator*).

1. Uji Gejala Multikolinearitas

Hasil pengujian adanya gejala Multikolinearitas menggunakan matrik korelasi dapat dilihat pada tabel 2. Dari matrik tersebut terlihat bahwa antar *independent variable*-nya (tidak termasuk variabel *dependent*) tidak ada

yang berkorelasi kuat dimana korelasi antar *independent variable*-nya di bawah 80%. Dengan kata lain hasil estimasi persamaan (9) terbebas dari gejala Multikolinearitas.

Tabel 2. Hasil Uji Gejala Multikolinearitas Persamaan (9)

Variabel	IPM	AHH	GRW	HP	MYS	PC	POV
IPM	1.000000	0.805969	-0.070350	0.232357	0.329781	0.007355	-0.133198
AHH	0.805969	1.000000	-0.290119	0.150804	0.128654	-0.145123	0.036295
GRW	-0.070350	-0.290119	1.000000	0.172234	0.167600	0.149604	-0.554981
HP	0.232357	0.150804	0.172234	1.000000	0.701195	0.635943	0.033844
MYS	0.329781	0.128654	0.167600	0.701195	1.000000	0.788956	0.018269
PC	0.007355	-0.145123	0.149604	0.635943	0.788956	1.000000	-0.011842
POV	-0.133198	0.036295	-0.554981	0.033844	0.018269	-0.011842	1.000000

Sumber: Hasil estimasi dengan Program EVIEWS 3.0

2. Uji Gejala Heteroskedastisitas.

Hasil pengujian adanya gejala Heteroskedastisitas menggunakan metode *Breusch-Pagan-Godfrey (BPG) Test* dapat dilihat pada tabel 3. Hasil pengujian menunjukkan nilai *chi-square* hitung = $\chi^2_{\text{hit}} = n.R^2$ adalah = 21,94475, sehingga dari hasil pengujian hipotesis didapatkan:

$\chi^2_{\text{hit}} = n.R^2 < \chi^2_{\text{tabel}} (5\%; df = 21) = 21,94475 < 32,6705$ yang berarti bahwa H_0 yang menyatakan tidak ada Heteroskedastisitas dapat diterima. Dengan demikian dapat disimpulkan hasil estimasi persamaan (9) terbebas dari gejala Heteroskedastisitas.

3. Uji Gejala Autokorelasi.

Hasil pengujian adanya gejala Autokorelasi dengan menggunakan metode *Breusch-Godfrey (BG) Serial Correlation LM Test* dapat dilihat pada tabel 4. Hasil pengujian menunjukkan nilai *chi-square* hitung $\chi^2_{\text{hit}} = (n-p)R^2$ adalah = 10,28302. Sehingga dari hasil pengujian hipotesis didapatkan:

$\chi^2_{\text{hit}} = (n-p)R^2 < \chi^2_{\text{tabel}} (5\%; df = 6) = 10,28302 < 32,6705$ yang berarti bahwa H_0 yang menyatakan tidak ada Autokorelasi dapat diterima. Dengan demikian dapat disimpulkan hasil estimasi persamaan (9) terbebas dari gejala Autokorelasi.

3.1.2 Pengujian Parameter Regresi (Uji Hipotesis)

1. Uji Pengaruh Parsial (Uji-t).

Uji-t dilakukan untuk melihat pengaruh *independent variable* secara parsial terhadap *dependent variable* dimana masing-masing *independent variable* yaitu AHH, GRW, HP, MYS, PC, dan POV akan diuji pengaruhnya/perannya terhadap IPM pada persamaan (9). Cara pengujiannya yaitu dengan membandingkan nilai *t* hasil estimasi (*t*-hitung) dengan nilai *t*-tabel pada *degree of freedom* (*df*) dan tingkat kepercayaan (α) tertentu. Dalam penelitian ini nilai *t*-tabel untuk *df* = 21 (*df* = *n*-*k*-1 = 28-6-1) pada tingkat kepercayaan 95% (α = 5%) adalah $t_{0,05 (21)} = 2.080$. Nilai *t*-hitung hasil estimasi seperti dalam tabel 1 adalah untuk variabel AHH = 5.476462, GRW = 0.202517, HP = -0.334719, MYS = 1.928180, PC = -0.865854, dan POV = -1.086607. Berdasarkan nilai *t*-hitung hasil estimasi pada masing-masing variabel maka dapat disimpulkan bahwa pada tingkat kepercayaan 95% (α = 5%), maka:

- 1) AHH (*t*-hitung > *t*-tabel); maka H_0 di tolak, yang berarti variabel angka harapan hidup (AHH) berpengaruh secara signifikan terhadap IPM.
- 2) GRW (*t*-hitung < *t*-tabel); maka H_0 di terima, yang berarti variabel pertumbuhan ekonomi (GRW) berpengaruh tidak signifikan terhadap IPM.
- 3) HP (*t*-hitung < *t*-tabel); maka H_0 di terima, yang berarti variabel banyaknya tenaga kesehatan (HP) berpengaruh tidak signifikan terhadap IPM.
- 4) MYS (*t*-hitung < *t*-tabel); maka H_0 di terima, yang berarti variabel rata-rata lama sekolah (MYS) berpengaruh tidak signifikan terhadap IPM.
- 5) PC (*t*-hitung < *t*-tabel); maka H_0 di terima, yang berarti variabel pendapatan perkapita (PC) berpengaruh tidak signifikan terhadap IPM.
- 6) POV (*t*-hitung < *t*-tabel); maka H_0 di terima, yang berarti variabel persentase penduduk miskin (POV) berpengaruh tidak signifikan terhadap IPM.

1. Uji Simultan (Uji-F).

Uji-F dilakukan untuk melihat pengaruh *independent variable* terhadap *dependent variable* secara simultan pada persamaan (9). Cara pengujiannya yaitu dengan membandingkan nilai *F* hasil estimasi (*F*-hitung) dengan nilai *F*-tabel pada *degree of freedom* (*df*) dan tingkat kepercayaan (α) tertentu. Dalam penelitian ini nilai *F*-tabel untuk *df*₁ = 5 (*df*₁ = *k*-1 = 6-1) dan *df*₂ = 21 (*df*₂ = *n*-*k* = 28-7) pada tingkat kepercayaan 95% (α = 5%) adalah $F_{0,05 (5 ; 21)} = 2,71$. Sedangkan nilai *F*-hitung hasil estimasi adalah 10,15696. Dengan demikian *F*-hitung > *F*-tabel sehingga H_0 di tolak yang berarti bahwa *dependent variable* yang terdiri dari AHH, GRW, HP, MYS, PC, dan POV dan secara simultan berpengaruh terhadap IPM.

2. Uji Koefisien Determinasi (R^2).

Hasil estimasi menunjukkan bahwa nilai koefisien determinasi (R^2 = *adjusted R-squared*) adalah 0,670498. Hal ini berarti variansi dari *independent variable* mampu menjelaskan 67,0498% terhadap variansi *dependent variable*. Dengan kata lain hubungan/peranan AHH, GRW, HP, MYS, PC, dan POV terhadap pembentukan indek pembangunan manusia (IPM) adalah 67,0498%. Sedangkan sebanyak 32,9502% faktor yang berhubungan

dengan pembentukan indeks pembangunan manusia (IPM) di Kota Bandar Lampung dan Kota Metro periode 2002 – 2015 disebabkan oleh variabel lain diluar variabel yang diteliti dalam penelitian ini.

Tabel 3. Hasil Uji Gejala Heteroskedastisitas Persamaan (9)

White Heteroskedasticity Test:				
F-statistic	4.530113	Probability		0.003654
Obs*R-squared	21.94475	Probability		0.038144
Test Equation:				
Dependent Variable: RESID^2				
Method: Least Squares				
Date: 10/22/16 Time: 18:43				
Sample: 1 28				
Included observations: 28				
Variable	Coefficient	Std. Error	t-Statistic	Prob.
C	-1054.043	440.4309	-2.393209	0.0302
AHH	36.17171	10.61442	3.407790	0.0039
AHH^2	-0.257081	0.075472	-3.406330	0.0039
GRW	0.589388	1.137661	0.518069	0.6120
GRW^2	-0.022030	0.084121	-0.261882	0.7970
HP	0.002337	0.001002	2.333262	0.0340
HP^2	-4.99E-07	3.14E-07	-1.588250	0.1331
MYS	-43.74856	25.04092	-1.747083	0.1011
MYS^2	2.125649	1.239015	1.715595	0.1068
PC	6.91E-08	1.49E-07	0.463517	0.6496
PC^2	-4.60E-15	4.37E-15	-1.052260	0.3093
POV	0.363112	0.394902	0.919498	0.3724
POV^2	-0.008042	0.013428	-0.598897	0.5582
R-squared	0.783741	Mean dependent var		0.749736
Adjusted R-squared	0.610734	S.D. dependent var		0.956874
S.E. of regression	0.597005	Akaike info criterion		2.110633
Sum squared resid	5.346217	Schwarz criterion		2.729157
Log likelihood	-16.54886	F-statistic		4.530113
Durbin-Watson stat	1.361073	Prob(F-statistic)		0.003654

Sumber: Hasil estimasi dengan Program EViews 3.0

Tabel 4. Hasil Uji Gejala Autokorelasi Persamaan (9)

Breusch-Godfrey Serial Correlation LM Test:				
F-statistic	5.513848	Probability		0.012933
Obs*R-squared	10.28302	Probability		0.005849
Test Equation:				
Dependent Variable: RESID				
Method: Least Squares				
Date: 10/22/16 Time: 18:25				
Variable	Coefficient	Std. Error	t-Statistic	Prob.
C	0.728796	9.975200	0.073061	0.9425
AHH	-0.045210	0.138525	-0.326371	0.7477
GRW	-0.069340	0.179646	-0.385981	0.7038
HP	0.000119	0.000458	0.260584	0.7972
MYS	0.344157	0.717560	0.479621	0.6370

PC	-6.95E-09	3.81E-08	-0.182583	0.8571
POV	-0.043623	0.059096	-0.738175	0.4694
RESID(-1)	0.692572	0.227151	3.048950	0.0066
RESID(-2)	-0.098589	0.244548	-0.403150	0.6913
R-squared	0.367251	Mean dependent var		1.74E-15
Adjusted R-squared	0.100830	S.D. dependent var		0.881762
S.E. of regression	0.836127	Akaike info criterion		2.735019
Sum squared resid	13.28306	Schwarz criterion		3.163228
Log likelihood	-29.29027	F-statistic		1.378462
Durbin-Watson stat	1.965244	Prob(F-statistic)		0.267528

Sumber: Hasil estimasi dengan Program EVIEWS 3.0

3.2 Pembahasan Faktor Penentu Indek Pembangunan Manusia (IPM) Di Kota Bandar Lampung dan Kota Metro

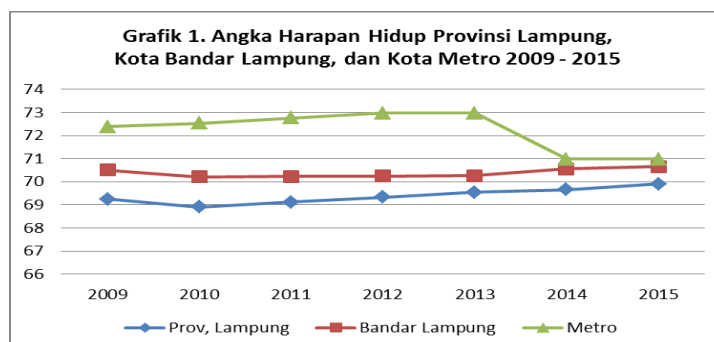
Pada bagian ini diuraikan pengaruh masing-masing variabel bebas terhadap IPM di Kota Bandar Lampung dan Kota Metro periode 2002 – 2015. Tujuannya adalah untuk mengetahui besarnya pengaruh masing-masing *independent variable*, faktor-faktor yang menyebabkan adanya signifikansi dan/atau ketidaksignifikanan pengaruh tersebut, dan dukungan teori yang mendukung adanya pengaruh tersebut.

1. Angka Harapan Hidup (AHH) Penduduk Kota Bandar Lampung dan Kota Metro

Berdasarkan hasil estimasi pada tabel 1 didapatkan bahwa angka harapan hidup (AHH) berpengaruh positif dan signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Artinya setiap ada peningkatan satu tahun angka harapan hidup bagi penduduk maka akan meningkatkan indek pembangunan manusia sebesar 0,887240 satuan indek. Angka harapan hidup (AHH) pada waktu lahir merupakan rata-rata perkiraan banyak tahun yang dapat ditempuh oleh seseorang selama hidup. Semakin panjang harapan penduduk untuk hidup (harapan panjang umur) maka nilai IPM juga akan semakin tinggi. Angka harapan hidup (AHH) saat lahir dapat diartikan sebagai jumlah tahun yang diharapkan dapat dicapai oleh bayi yang baru lahir untuk hidup, dengan asumsi bahwa pola angka kematian menurut umur pada saat kelahiran sama sepanjang usia bayi. Angka harapan hidup (AHH) juga merupakan umur panjang dan hidup sehat.

Angka harapan hidup (AHH) saat lahir juga merepresentasikan dimensi umur panjang dan hidup sehat terus meningkat dari tahun ke tahun. Di Provinsi Lampung (grafik 1), angka harapan hidup (AHH) untuk Kota Bandar Lampung dan Kota Metro secara rata-rata lebih tinggi dibandingkan dengan angka harapan hidup (AHH) Provinsi Lampung. Namun demikian selama periode pengamatan, secara rata-rata pertumbuhan angka harapan hidup (AHH) Provinsi Lampung lebih tinggi, yaitu sebesar 0,29% per tahun. Sedangkan rata-rata pertumbuhan angka harapan hidup (AHH) untuk Kota Bandar Lampung 0,036% per tahun dan Kota Metro -0,32% per tahun. Makna yang terkandung dalam data angka harapan hidup (AHH) adalah misalnya pada tahun 2014 angka harapan hidup di kota Bandar Lampng adalah 70,55 tahun. Maka angka tersebut menunjukkan bahwa rata-rata tahun hidup yang akan dijalani oleh anak-anak (bayi) yang lahir pada tahun 2014 diperkirakan akan hidup sampai umur 70,55 tahun. Dalam BPS (2016) dinyatakan capaian pembangunan manusia di suatu wilayah pada waktu tertentu dapat dikelompokkan ke dalam empat kelompok. Pengelompokan ini bertujuan untuk mengorganisasikan wilayah-wilayah menjadi kelompok-kelompok yang sama dalam hal pembangunan

manusia. Pengelompokan tersebut IPM adalah (a) Kelompok “sangat tinggi” ($IPM \geq 80$); (b) Kelompok “tinggi” ($70 \leq IPM < 80$); (c) Kelompok “sedang” ($60 \leq IPM < 70$); dan (d) Kelompok “rendah” ($IPM < 60$).



Gambar 1. Angka Harapan Hidup Provinsi Lampung, Kota Bandar Lampung, dan Kota Metro 2009-2015

Berdasarkan pengelompokan tersebut, nilai IPM untuk Provinsi Lampung masuk dalam kelompok “sedang”. Nilai indeks pembangunan manusia (IPM) untuk Kota Bandar Lampung dan Kota Metro masuk dalam kelompok “tinggi”. Dalam BPS (2016) dinyatakan sebagian besar Kabupaten/Kota yang ada di Provinsi Lampung (14 Kab/Kota) memiliki IPM dengan status “sedang”. Satu kabupaten pada tahun 2015, yaitu Kabupaten Mesuji berstatus “rendah”. Sedangkan dua kota di Lampung yaitu Kota Metro dan Kota Bandar Lampung keduanya berstatus “tinggi”. Selanjutnya dalam BPS (2016) dinyatakan pada periode 2014-2015, tercatat tiga kabupaten dengan kemajuan pemba-ngunan manusia paling cepat, yaitu Kabupaten Lampung Selatan (2,31%), Kabupaten Mesuji (1.83%), dan Kabupaten Pesawaran (1,62%). Kemajuan pembangunan manusia di Kabupaten Mesuji dan Kabupaten Pesawaran didorong oleh dimensi pendidikan, sementara di Kabupaten Lampung Selatan lebih dikarenakan perbaikan standar hidup layak. Sementara itu, kemajuan pembangunan manusia di Kota Metro (0,16%), Kabupaten Tulang Bawang (0,36%), dan Kabupaten Lampung Utara (0,48%) tercatat paling lambat di Provinsi Lampung.

Hasil penelitian ini sejalan dengan [4] yang menyatakan bahwa angka harapan hidup masyarakat terus mengalami perubahan. Berkaitan dengan dimensi kesehatan (hidup sehat) yang merupakan bagian dari angka harapan hidup (AHH), maka hasil penelitain ini juga sejalan dengan hasil [6]. [6] menunjukkan bahwa pengeluaran pemerintah fungsi kesehatan di Provinsi Lamung berpengaruh positif dan signifikan terhadap IPM, dimana setiap peningkatan pengeluaran pemerintah fungsi kesehatan sebesar Rp 1 Milyar akan meningkatkan IPM sebesar 0,00569. Selain itu hasil penelitian ini juga sejalan dengan [1] menyatakan rumah tangga dengan akses air bersih memiliki hubungan positif terhadap IPM, dimana secara umum IPM Provinsi Jawa Timur dari tahun 2004 hingga 2011 mengalami peningkatan secara terus menerus.

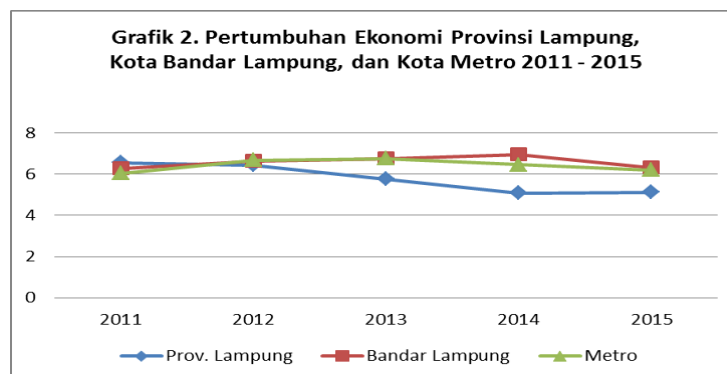
Indek pembangunan manusia dihitung berdasarkan rata-rata geometrik indeks kesehatan, indeks pengetahuan, dan indeks pengeluaran. IPM dibentuk oleh tiga dimensi dasar, yaitu umur panjang dan hidup sehat (*a long and healthy life*), pengetahuan (*knowledge*), dan standard hidup layak (*decent standard of living*). Umur panjang dan hidup sehat digambarkan oleh angka harapan hidup saat lahir (AHH). Karena angka harapan hidup (AHH) saat lahir juga merepresentasikan dimensi umur panjang dan hidup sehat yang terus mening-kat dari tahun ke tahun,

maka hasil penelitian ini berbeda dengan hasil penelitian [5] yang menunjukkan pengeluaran rumah tangga untuk kesehatan tidak berpengaruh signifikan terhadap pembentukan IPM. Pemerintah daerah Kab/Kota di Provinsi Bali perlu meningkatkan kualitas dan fasilitas pendidikan dan kesehatan serta layanan pendidikan gratis dan pengobatan gratis bagi masyarakat miskin.

2. Pertumbuhan Ekonomi (GRW) Kota Bandar Lampung dan Kota Metro

Berdasarkan hasil estimasi seperti yang ditunjukkan pada tabel 1 didapatkan bahwa pertumbuhan ekonomi (GRW) berpengaruh positif dan tidak signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Artinya setiap ada kenaikan satu persen akan meningkatkan indeks pembangunan manusia sebesar 0,039669 satuan indeks, namun demikian peningkatan ini pengaruhnya tidak signifikan. Meskipun hasil estimasi tidak signifikan, namun hal yang menarik dari penelitian ini adalah tanda koefisien regresi dari GRW adalah positif (+). Tanda positif ini berarti hasil estimasi ini sesuai dengan harapan *apriori*, yaitu bahwa pertumbuhan ekonomi yang semakin tinggi akan meningkatkan faktor pembentukan indeks pembangunan manusia. Gujarati (2003:219) menyatakan model regresi/estimasi yang baik tidak semata-mata melihat nilai *adjusted R²* yang tinggi, tetapi juga harus mempertimbangkan koefisien regresinya apakah nyata secara statistik (*statistically significant*) dan tanda koefisien regresinya apakah sesuai dengan harapan *apriori*.

Berdasarkan grafik 2, pertumbuhan ekonomi di Kota Bandar Lampung dan Kota Metro berada di atas Provinsi Lampung. Bahkan rata-rata pertumbuhannya kurun waktu 2011 – 2015 jauh di atas Provinsi Lampung yaitu (-5,80%) untuk Provinsi Lampung, sedangkan Kota Bandar Lampung (0,28%) dan Kota Metro (0,88%). Suatu perekonomian dikatakan mengalami pertumbuhan ekonomi jika jumlah produksi barang dan jasanya mengalami peningkatan. Tujuan utama menghitung pertumbuhan ekonomi adalah untuk melihat apakah kondisi perekonomian tersebut semakin membaik atau semakin memburuk. Adanya pertumbuhan ekonomi yang semakin membaik juga akan menjamin semakin membuka kesempatan kerja karena adanya peningkatan output, meningkatkan pendapatan masyarakat, dan pada akhirnya akan membantu memperbaiki/ meningkatkan nilai IPM daerah tersebut.



Gambar 2. Pertumbuhan Ekonomi Hidup Provinsi Lampung, Kota Bandar Lampung, dan Kota Metro 2011-2015

Untuk menghitung nilai pertumbuhan ekonomi digunakan nilai PDRB berdasarkan harga konstan (*PDRB real*) tahun tertentu. Penggunaan *PDRB real* bertujuan untuk menghilangkan pengaruh perubahan (kenaikan) harga (inflasi) sehingga perubahan nilai PDRB sekaligus menunjukkan perubahan jumlah kuantitas barang dan jasa

selama periode pengamatan. Penggunaan variabel PDRB, yang merupakan indikator untuk menghitung pertumbuhan ekonomi (GRW), karena peningkatan PDRB dipandang akan membantu menjadi faktor peningkatan IPM. Prathama & Mandala (2001:17) menyatakan terdapat beberapa alasan mengapa PDRB dapat dijadikan sebagai indikator kemajuan sebuah perekonomian yaitu, *pertama*, besarnya output daerah (PDRB) merupakan gambaran awal tentang seberapa efisien sumber daya yang ada dalam perekonomian (tenaga kerja, barang modal, uang dan *entrepreneurship*) digunakan untuk memproduksi barang dan jasa. Secara umum makin besar pendapatan daerah maka semakin baik efisiensi alokasi sumber daya ekonominya. *Kedua*, besarnya output daerah (PDRB) merupakan gambaran awal tentang produktifitas dan kemakmuran suatu daerah. *Ketiga*, besarnya output daerah (PDRB) merupakan gambaran awal tentang masalah-masalah struktural (mendasar) yang dihadapi suatu perekonomian. Jika sebagian besar output daerah dinikmati oleh sebagian kecil penduduk, maka perekonomian tersebut menghadapi masalah dengan distribusi pendapatan. Jika sebagian besar output daerah berasal dari sektor pertanian (ekstraktif), maka perekonomian menghadapi masalah ketimpangan produksi.

Hasil estimasi terhadap pertumbuhan ekonomi (GRW) yang tidak signifikan tersebut menginformasikan beberapa kemungkinan yang terjadi pada pertumbuhan ekonomi (GRW) di Kota Bandar Lampung dan Kota Metro. *Pertama*, pertumbuhan ekonomi yang telah dicapai oleh Kota Bandar Lampung dan Kota Metro tidak berpengaruh terhadap upaya menaikkan pendapatan seluruh lapisan masyarakat secara merata serta pertumbuhan ekonomi yang telah dicapai ternyata tidak bisa mengurangi gap pendapatan antara orang kaya dan orang miskin. *Kedua*, pertumbuhan ekonomi yang tinggi (*growth oriented*) kemungkinan justru hanya memicu munculnya kesenjangan pendapatan dan *in-equality*. Proses pembangunan ekonomi yang *growth oriented* pada keadaan-keadaan tertentu yaitu pada saat keterbelakangan ekonomi dan kemajuan pembangunan ekonomi yang tinggi, ternyata sama-sama tidak mampu merubah kondisi masyarakat ke keadaan yang lebih baik. Hal ini berarti seperti juga pendapat Adelman & Morris (1973:6) dalam Bashri (2003) tidak adanya *trickle down effect* yang bersifat otomatis yang mengalirkan hasil pembangunan kepada semua golongan/lapisan masyarakat. *Ketiga*, pertumbuhan ekonomi yang tinggi yang selama ini dicapai oleh Kota Bandar Lampung dan Kota Metro ternyata tidak mampu mengurangi faktor pembentuk IPM. Kenaikan pertumbuhan ekonomi tersebut hanya bisa dinikmati oleh sebagian kecil masyarakat. Pertumbuhan ekonomi yang tinggi hanya bisa dinikmati oleh sebagian kecil orang kaya, sementara bagian terbesar masyarakat tidak mampu menikmatinya. Keadaan ini seperti dinyatakan oleh Todaro, (2000:206) sesuai dengan teori “*trade off between growth and equity*” yang menyatakan bahwa pertumbuhan ekonomi yang tinggi akan menimbulkan ketimpangan yang semakin besar dalam pembagian pendapatan atau makin tidak merata, dan sebaliknya upaya pemerataan dapat terwujud dalam pertumbuhan ekonomi yang rendah.

Namun demikian dalam jangka panjang prestasi pertumbuhan ekonomi (GRW) yang tinggi yang telah dicapai oleh Kota Bandar Lampung dan Kota Metro lambat laun akan menimbulkan pemerataan pendapatan. Hal ini sesuai dengan pandangan Kuznets dalam Wie (1983:4) yang menjelaskan mengenai hubungan jangka panjang antara pertumbuhan ekonomi dan pembagian pendapatan. Teori/pandangan ini menyatakan bahwa proses pembangunan ekonomi pada tahap awal umumnya disertai oleh kemerosotan yang cukup besar dalam pembagian pendapatan, yang baru berbalik menuju suatu pemerataan yang lebih besar dalam pembagian pendapatan pada tahap pembangunan lebih lanjut. Hipotesis Kuznets ini berupa kurva *U* terbalik dimana ketika

pembangunan baru dimulai, distribusi pendapatan akan makin tidak merata (dan tidak bisa dinikmati oleh seluruh masyarakat), namun setelah mencapai suatu tingkat pembangunan tertentu, distribusi pendapatan semakin merata (kemakmuran). Dengan semakin makmurnya masyarakat, maka dapat meningkatkan faktor pembentuk IPM.

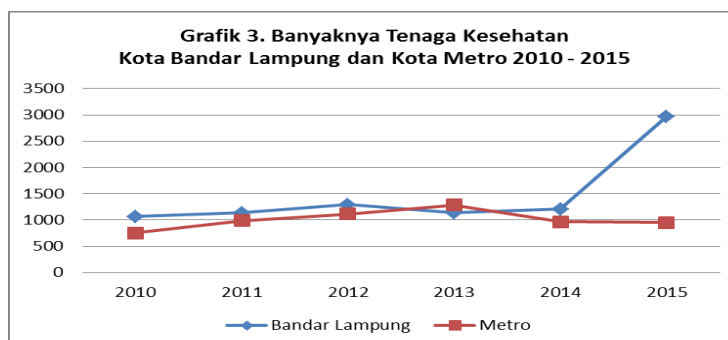
3. Banyaknya Tenaga Kesehatan (HP) Kota Bandar Lampung dan Kota Metro

Berdasarkan hasil estimasi seperti yang ditunjukkan pada tabel 1 didapatkan bahwa banyaknya tenaga kesehatan (HP) berpengaruh negatif dan tidak signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Tanda koefisien hasil estimasi ini bertentangan dengan harapan *a priori*, karena koefisien hasil estimasi HP adalah negatif. Artinya setiap ada kenaikan tenaga kesehatan (HP) sebesar satu orang, maka akan menurunkan faktor penyebab pembentukan indeks pembangunan manusia (IPM) sebesar 0,000179 satuan indeks. Sehubungan sesuai dengan harapan *a priori* hasil estimasi terhadap HP ini memiliki tanda positif (+) yaitu setiap ada kenaikan tenaga kesehatan (HP) maka akan menaikkan faktor pembentukan IPM. Hasil koefisien estimasi yang negatif ini mengindikasikan kemungkinan, *pertama*, telah terjadi kekurangoptimalan managerial dalam pengelolaan SDM tenaga kesehatan di Kota Bandar Lampung dan Kota Metro. SDM tenaga kesehatan yang ada belum secara maksimal dimanfaatkan untuk meningkatkan kesejahteraan masyarakat.

Kedua, kemungkinan produktifitas dan kinerja SDM kesehatan yang ada di Kota Bandar Lampung dan Kota Metro masih rendah. Sehingga keberadaan (kerja-kerja) mereka belum mampu dirasakan oleh masyarakat dalam rangka meningkatkan kesehatan masyarakat. Karena IPM dihitung berdasarkan rata-rata geometrik indeks kesehatan, indeks pengetahuan, dan indeks pengeluaran. *Ketiga*, kemungkinan jumlah tenaga kesehatan yang ada di Kota Bandar Lampung dan Kota Metro masih sangat kurang. Sehingga persentasenya tidak berimbang dengan jumlah penduduk yang harus dilayani, meskipun berdasarkan grafik 3 jumlahnya terus meningkat dengan rata-rata pertumbuhan Kota Bandar Lampung (26,86%) dan Kota Metro (5,56%). Hasil penelitian yang sama dengan penelitian ini adalah hasil penelitian Erwin, dkk (2014:140-153) yang menunjukkan pengeluaran rumah tangga untuk kesehatan tidak berpengaruh signifikan terhadap kesejahteraan masyarakat (IPM) Kab/Kota di Provinsi Bali. Pemda perlu meningkatkan kualitas, fasilitas kesehatan, layanan pendidikan gratis, dan pengobatan gratis bagi masyarakat miskin.

Hasil penelitian ini berbeda dengan hasil penelitian Ayunanda & Ismaini (2013:237-242) menyatakan secara umum, IPM Provinsi Jawa Timur dari tahun 2004 hingga 2011 mengalami peningkatan secara terus menerus dan rumah tangga dengan akses air bersih dan jumlah sarana kesehatan memiliki hubungan positif terhadap IPM. Untuk meningkatkan IPM dapat dilakukan dengan cara meningkatkan jumlah sarana kesehatan dan persentase rumah tangga dengan akses air bersih. Hasil penelitian yang sama ditunjukkan oleh penelitian Wyati & Teguh (2011:28-39) yang menunjukkan faktor yang mempengaruhi mutu sumberdaya manusia di Kota Semarang adalah kesehatan. Berdasarkan derajat kesehatan maupun pelayanan kesehatan dapat diketahui bahwa angka harapan hidup masyarakat terus mengalami perubahan. Penelitian Wyati dan Teguh menyarankan perlu

meningkatkan peringkat angka harapan hidup dan memperhatikan mutu pelayanan kesehatan baik dari segi jumlah tenaga dokter maupun jumlah puskesmas.



Gambar 3. Banyaknya Tenaga Kesehatan
Kota Bandar Lampung dan Kota Metro 2010-2015

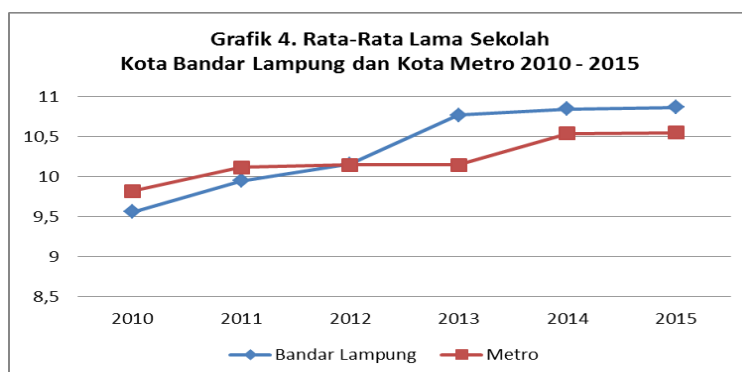
Selain itu hasil penelitian ini juga berbeda dengan hasil penelitian Prayudha (2013:243-257) menunjukkan pengeluaran pemerintah fungsi kesehatan juga berpengaruh positif dan signifikan terhadap IPM, dimana setiap peningkatan pengeluaran pemerintah fungsi kesehatan sebesar Rp 1 Milyar akan meningkatkan IPM sebesar 0,00569. Selain itu juga perlu pengontrolan keefektifitasan alokasi kesehatan yang dikeluarkan dalam rangka meningkatkan IPM. Hasil penelitian Kacaribu (2013) juga menunjukkan hasil yang sama, yaitu rasio jumlah penduduk terhadap jumlah dokter, rasio jumlah penduduk terhadap jumlah bidan, dan rasio jumlah penduduk terhadap jumlah perawat mempengaruhi IPM di Kab/Kota di Provinsi Papua.

4. Rata-Rata Lama Sekolah (MYS) Kota Bandar Lampung dan Kota Metro

Hasil estimasi seperti yang ditunjukkan pada tabel 1 menunjukkan bahwa bahwa rata-rata lama sekolah (MYS) berpengaruh positif, tetapi tidak signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Artinya setiap ada kenaikan MYS sebesar satu tahun, maka akan meningkatkan faktor penyebab pembentukan IPM sebesar 1.611902 satuan indek. Meskipun hasil estimasi tidak signifikan, namun hal yang menarik dari penelitian ini adalah tanda koefisien regresi dari MYS adalah positif (+). Tanda positif ini berarti hasil estimasi ini sesuai dengan harapan *a priori*, yaitu bahwa rata-rata lama sekolah (MYS) yang semakin tinggi akan meningkatkan faktor penyebab pembentukan IPM. Gujarati (2003:219) menyatakan model regresi/estimasi yang baik tidak semata-mata melihat nilai *adjusted R²* yang tinggi, tetapi juga harus mempertimbangkan koefisien regresinya apakah nyata secara statistik (*statistically significant*) dan tanda koefisien regresinya apakah sesuai dengan harapan *a priori*.

Rata-rata lama sekolah (MYS) menggambarkan jumlah tahun yang digunakan oleh penduduk usia 15 tahun keatas dalam menjalani pendidikan formal. Semakin lama penduduk menjalani pendidikan formalnya, maka akan meningkatkan faktor pembentukan indek pembangunan manusia (IPM). Rata-rata lama sekolah di Kota Bandar Lampung dan Kota Metro menunjukkan *trend* yang terus meningkat dengan pertumbuhan 2,18% untuk Kota Bandar Lampung dan 1,21% untuk Kota Metro. Rata-rata lama sekolah untuk Kota Bandar Lampung (10,36 tahun) dan Kota Metro (10,22 tahun). Ini berarti tingkat pendidikan rata-rata penduduk usia 15 tahun

keatas di Kota Bandar Lampung dan Kota Metro sudah tamat SMP dan yang Sederajat, dan bahkan sudah masuk kelas 1 (kelas 10) SMU. Hal ini juga bermakna Program Nasional Wajib Belajar (Wajar) Sembilan Tahun telah sukses dijalankan di Kota Bandar Lampung dan Kota Metro.



Gambar 4. Rata-rata Lama Sekolah
Kota Bandar Lampung dan Kota Metro 2010-2015

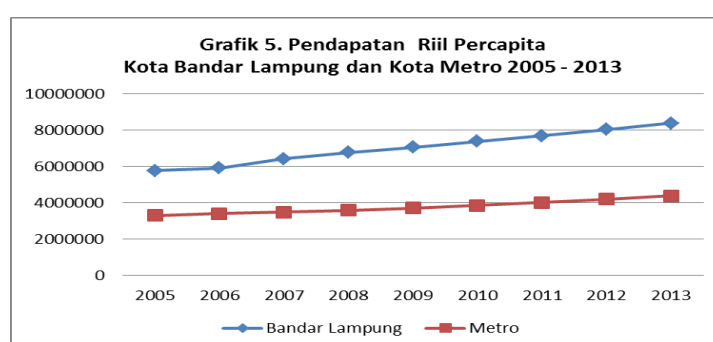
Hasil penelitian ini berbeda dengan hasil penelitian [4] menunjukkan faktor yang mempengaruhi mutu sumberdaya manusia di Kota Semarang adalah pendidikan. Pada bidang pendidikan dapat diketahui bahwa angka melek huruf dan rata-rata lama sekolah dari tahun ke tahun cenderung sama atau tidak mengalami perubahan. Penelitian Wyati dan Teguh menyarankan perlu meningkatkan angka melek huruf, rata-rata lama sekolah, dan angka partisipasi murni untuk usia 13-15 tahun dan 16-18 tahun. Meskipun berbeda, hasil penelitian ini masih memiliki tanda koefisien estimasi yang positif (+), yaitu setiap ada kenaikan MYS sebesar satu tahun, maka akan meningkatkan faktor penyebab pembentukan IPM sebesar 1.611902 satuan indek di Kota Bandar Lampung dan Kota Metro.

5. Pendapatan Perkapita (PC) Kota Bandar Lampung dan Kota Metro

Berdasarkan hasil estimasi seperti yang ditunjukkan pada tabel 1 didapatkan bahwa pendapatan perkapita penduduk (PC) berpengaruh negatif, tetapi tidak signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Tanda koefisien hasil estimasi ini bertentangan dengan harapan *a priori*, karena koefisien hasil estimasi PC adalah negatif. Artinya setiap ada kenaikan PC sebesar Rp 1.000.000,- maka akan menurunkan faktor penyebab pembentukan IPM sebesar $3,91 \times 10^{-8}$ satuan indek. Seharusnya sesuai dengan harapan *a priori* hasil estimasi terhadap PC ini memiliki tanda positif (+) yaitu setiap ada kenaikan pendapatan perkapita maka akan meningkatkan faktor pembentukan indek pemba-ngunan manusia. Hasil koefisien estimasi yang negatif ini mengindikasikan kemungkinan telah terjadi ketimpangan pendapatan di Kota Bandar Lampung dan Kota Metro. Dimana kenaikan pendapatan perkapita tersebut hanya dapat dinikmati oleh sebagian kecil penduduk di kedua kota tersebut. Sementara sebagian besar penduduk masih tetap tidak memiliki kemampuan untuk mengakses peningkatan pendapatan perkapita tersebut. Munculnya ketimpangan pendapatan ini akan menghambat upaya pembentukan IPM. Hasil penelitian Rifa'i (2010: 317-327) menunjukkan diindikasikan terjadi ketimpangan distribusi pendapatan dan kegagalan kebijakan pemerintah dalam mendistribusikan hasil-hasil pemba-ngunan secara merata ke seluruh penduduk di Kota Bandar Lampung

dan Kota Metro. Indikasi ini terlihat dari adanya pengaruh yang positif dan signifikan dari pendapatan perkapita terhadap kemiskinan absolut. Artinya semakin besar nilai pendapatan perkapita maka kemiskinan absolut (ketimpangan pendapatan) juga semakin meningkat.

Dalam BPS (2016) dinyatakan bahwa komponen IPM adalah pengeluaran riil per kapita yang disesuaikan (PDRB/Capita atas dasar harga konstan). Nilai pendapatan per kapita dapat diperoleh dengan cara membagi besarnya PDRB dengan jumlah penduduk pada tahun yang bersangkutan. Artinya semakin tinggi nilai PDRB, maka akan semakin tinggi nilai pendapatan per kapita. Berdasarkan grafik 5 nilai pendapatan per kapita atas dasar harga konstan 2000 Kota Bandar Lampung dan Kota Metro terus meningkat, yaitu Kota Bandar Lampung (4,8%) dan Kota Metro (3,6%). Sedangkan nilai rata-rata pendapatan per kapita untuk kurun waktu 2005 – 2013 untuk Kota Bandar Lampung (Rp 7.057.434) dan Kota Metro (Rp 3.775.888).



Gambar 5. Pendapatan Riil Per Kapita
Kota Bandar Lampung dan Kota Metro 2005-2013

Sama halnya dengan pendapatan perkapita, besarnya nilai pertumbuhan ekonomi diperoleh dengan cara melihat perubahan (peningkatan atau penurunan) dari tahun ke tahun nilai PDRB berdasarkan harga konstan (PDRB *real*). Dengan demikian hasil penelitian ini sejalan dengan hasil penelitian [9] yang menyatakan bahwa pertumbuhan ekonomi yang tinggi (*growth oriented*) justru hanya memicu munculnya kesenjangan pendapatan dan *in-equality*. Hal senada juga dihasilkan dalam penelitian [10] yang menyatakan bahwa pertumbuhan ekonomi tidak berpengaruh terhadap upaya menaikkan pendapatan penduduk miskin serta pertumbuhan ekonomi tidak bisa mengurangi gap pendapatan antara orang kaya dan orang miskin.

Hasil penelitian ini juga sejalan dengan penelitian [11] yang menunjukkan di Indonesia tahun 1985-1996 yang mengindikasikan telah terjadi ketimpangan dalam pemerataan hasil-hasil pembangunan di Indonesia pada kurun waktu tersebut. Selain itu hasil penelitian ini sejalan dengan penelitian [12] yang menunjukkan bahwa tingginya pertumbuhan pendapatan per kapita tidak akan terlalu berdampak apabila tidak disertai dengan perbaikan dalam hal distribusi pendapatan. Hal senada juga dihasilkan dalam penelitian yang dilakukan [13] mengungkapkan tentang peran pembangunan ekonomi di negara-negara berkembang dimana negara-negara tersebut bukan saja menghadapi kemerosotan dalam ketimpangan relatif akibat pertumbuhan ekonomi. Proses pembangunan ekonomi yang *growth oriented* pada keadaan-keadaan tertentu yaitu pada saat keterbelakangan ekonomi dan kemajuan pembangunan ekonomi yang tinggi, ternyata sama-sama menimbulkan keadaan yang lebih buruk bagi

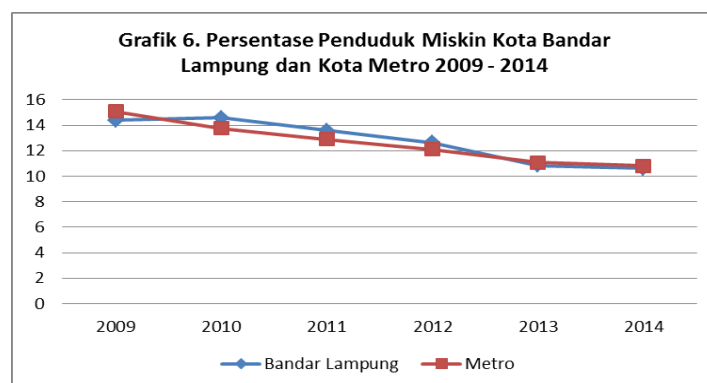
sekitar 60% penduduk yang tergolong miskin dan tidak adanya *trickle down effect* yang bersifat otomatis yang mengalirkan hasil pembangunan kepada masyarakat luas.

Koefisien hasil estimasi PC yang berlawanan dengan harapan *a priori* ini, seharusnya koefisien estimasi bertanda negatif (+), juga mengindikasikan beberapa hal, *pertama*, produktifitas penduduk pada kedua kota tersebut rendah. Dimana meskipun mereka mengalami kenaikan pendapatan perkapita tetapi mereka masih tetap berada dalam kondisi belum sejahtera. *Kedua*, upah tenaga kerja di kedua kota tersebut rendah. Dimana meskipun mereka mengalami kenaikan pendapatan perkapita, yang di-*proxy* dengan naiknya upah tenaga kerja, namun kenaikan ini tetap juga tidak mampu menciptakan kesejahteraan. Dimana kemungkinan pendapatan yang mereka terima belum bisa digunakan untuk memenuhi/ mencukupi kebutuhan pendidikan dan kesehatan. *Ketiga*, telah terjadi kegagalan pemerataan hasil-hasil pembangunan sehingga menimbulkan ketimpangan pendapatan. Pemerintah daerah tentunya telah berupaya melakukan pembangunan dan pertumbuhan ekonomi, namun kebijakan pemerintah tidak berhasil menyebarkan hasil-hasil pembangunan dan pertumbuhan tersebut secara merata pada seluruh penduduk. Sebagian kecil penduduk ada yang menikmati hasil-hasil pembangunan dan pertumbuhan dan sebagian besar penduduk tidak dapat menikmati hasil pembangunan tersebut, misalnya untuk memenuhi kebutuhan pendidikan dan kesehatan sebagai unsur utama pembentuk IPM.

6. Persentase Penduduk Miskin (POV) Kota Bandar Lampung dan Kota Metro

Hasil estimasi seperti yang ditunjukkan pada tabel 1 menunjukkan bahwa bahwa persentase penduduk miskin (POV) berpengaruh negatif, tetapi tidak signifikan terhadap faktor-faktor yang mempengaruhi pembentukan indek pembangunan manusia (IPM) di Kota Bandar Lampung dan Kota Metro. Artinya setiap ada penurunan POV sebesar satu persen, maka akan meningkatkan faktor penyebab pembentukan IPM sebesar 0.072288 satuan indek. Meskipun hasil estimasi tidak signifikan, namun hal yang menarik dari penelitian ini adalah tanda koefisien regresi dari POV adalah negatif (-). Tanda negatif ini berarti hasil estimasi ini sesuai dengan harapan *a priori*, yaitu bahwa persentase penduduk miskin (POV) yang semakin rendah akan meningkatkan faktor penyebab pembentukan IPM. [8] menyatakan model regresi/ estimasi yang baik tidak semata-mata melihat nilai *adjusted R²* yang tinggi, tetapi juga harus mempertimbangkan koefisien regresinya apakah nyata secara statistik (*statistically significant*) dan tanda koefisien regresinya apakah sesuai dengan harapan *a priori*. Berdasarkan grafik 6, persentase penduduk miskin (POV) di Kota Bandar Lampung dan Kota Metro menunjukkan trend yang terus menurun. Selama kurun waktu 2010 – 2015 persentase penduduk miskin terus menurun dengan pertumbuhan penurunannya untuk Kota Bandar Lampung (-5,78%) dan untuk Kota Metro (-6,38%). Rata-rata persentase penduduk miskin yang ada di Kota Bandar Lampung adalah 12,78% dan di Kota Metro adalah 12,62%.

Hasil penelitian ini berbeda dengan hasil penelitian Kacaribu (2013) menunjukkan bahwa rasio kemiskinan terhadap jumlah penduduk mempengaruhi IPM di Kab/Kota di Provinsi Papua. Hasil penelitian yang sama juga ditunjukkan oleh penelitian [14] dimana rasio penduduk miskin berpengaruh negatif dan signifikan terhadap IPM di Indonesia. Saran yang berkaitan dengan hasil penelitian ini yaitu pemerintah Indonesia diharapkan dapat mengurangi jumlah penduduk miskin sehingga akan meningkatkan indeks pembangunan manusia di Indonesia.



Gambar 6. Persentase Penduduk Miskin
Kota Bandar Lampung dan Kota Metro 2009-2014

IV. KESIMPULAN

Kesimpulan penelitian Analisis Indeks Pembangunan Manusia (IPM) dan Faktor-Faktor Yang Mempengaruhi Di Kota Bandar Lampung dan Kota Metro adalah sebagai berikut:

1. Secara parsial variabel angka harapan hidup (AHH) berpengaruh signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro. Sedangkan variabel pertumbuhan ekonomi (GRW), banyaknya tenaga kesehatan (HP), rata-rata lama sekolah (MYS), pendapatan perkapita (PC), dan persentase penduduk miskin (POV) berpengaruh tidak signifikan.
2. Secara simultan variabel angka harapan hidup (AHH), pertumbuhan ekonomi (GRW), banyaknya tenaga kesehatan (HP), rata-rata lama sekolah (MYS), pendapatan perkapita (PC), dan persentase penduduk miskin (POV) berpengaruh signifikan terhadap faktor-faktor yang mempengaruhi pembentukan IPM di Kota Bandar Lampung dan Kota Metro.
3. Koefisien Determinasi (R^2) hasil estimasi menunjukkan bahwa nilai koefisien determinasi adalah 0,670498. Hal ini berarti variansi dari *independent variable* mampu menjelaskan 67,0498% terhadap variansi *dependent variable*.

V. KEPUSTAKAAN

- [1] BPS. 2016. Berita Resmi Statistik No. 15/06/18/TAHUN I, 16 Juni 2016. BPS Provinsi Lampung
- [2] Rifa'i, A. & Triatmojo, F. 2011. Studi Karakteristik Rumah Tangga Miskin Di Kota Metro. *Laporan Hibah Bappeda Kota Metro Oktober 2011*.
- [3] Ayunanda, M. & Ismaini, Z. 2013. *Analisis Statistika Faktor Yang Mempengaruhi Indeks Pembangunan Manusia Di Kabupaten/Kota Provinsi Jawa Timur Dengan Menggunakan Regresi Panel* . Jurnal Sains Dan Seni Pomits Vol. 2, No.2, (2013) 2337-3520 (2301-928X Print)
- [4] Wyati, S. & Teguh, A. 2011. *Analisis Faktor-Faktor Yang Mempengaruhi Indeks Pembangunan Manusia (IPM) di Kota Semarang*. Jurnal Dinamika Sosbud Volume 13 Nomor 1, Juni 2011: 28 – 39. ISSN 1410-9859.
- [5] Erwin, N. Nyoman D.S., & Djayastra, I.K. 2014. *Analisis Faktor-Faktor Yang Mempengaruhi Kesejahteraan Masyarakat Kabupaten/Kota Di Provinsi Bali*. E-Jurnal Ekonomi Dan Bisnis Universitas Udayana Volume. 03 No. 03 Tahun 2014.
<http://id.portalgaruda.org/index.php?ref=browse&mod=viewarticle&article=151110>
<http://ojs.unud.ac.id/index.php/EEB/article/viewFile/7619/6044>
- [6] Prayudha, A. 2013. *Determinants Of Human Development In Lampung Province*. Jurnal Ekonomi Pembangunan (JEP)-Vol. 2, No.3, September 2013
- [7] Kacaribu, R.D. 2013. *Analisis Indeks Pembangunan Manusia Dan Faktor-Faktor Yang Memengaruhi Di Provinsi Papua*. <http://repository.ipb.ac.id/handle/123456789/63886>
- [8] Gujarati, D.N. 2003. *Basic Econometrics*. Fourth Edition. McGrawHill Singapore.
- [9] Tambunan, T.T.H.. 2001. *Transformasi Ekonomi Di Indonesia*. Jakarta : Salemba Empat.
- [10] Foster, E. J. & Szekely, M. 2002. *Is Economic Growth Good for the Poor? Tracking Low Incomes Using General Means*. Report on Symposium on Poverty Measurement, Mexico.
- [11] Booth, A. 2000. *Poverty and Inequality in the Soeharto Era: An Assesment*. Bulletin of Indonesian Economics Studies, Vol.36, No.1.
- [12] Iradian, G. 2005. *Inequality, Poverty, and Growth: Cross Country Evidence*. IMF Working Paper. Middle East and Central Asia Departement.
- [13] Bashri, Y. 2003. *Mau Ke Mana Pembangunan Ekonomi Indonesia, Prisma Pemikiran Prof. Dr. Dorodjatun Kuntjoro-Jakti*. Jakarta: Prenada.
- [14] Syarifah. 2012. *Faktor—faktor yang Mempengaruhi Indeks Pembangunan Manusia di Indonesia Tahun 2005-2009*. Universitas Negeri Semarang. <http://lib.unnes.ac.id/13471/>

Parameter Estimation of Bernoulli Distribution using Maximum Likelihood and Bayesian Methods

Nurmaita Hamsyah¹⁾, Khoirin Nisa¹⁾, & Warsono¹⁾

¹⁾Department of Mathematics, Faculty of Mathematics and Science, University of Lampung
Jl. Prof. Dr. Sumantri Brojonegoro No. 1 Bandar Lampung
Phone Number +62 721 701609 Fax +62 721 702767
E-mail: itamath98@gmail.com

ABSTRACT

The term parameter estimation refers to the process of using sample data to estimate the parameters of the selected distribution. There are several methods that can be used to estimate distribution parameter(s). In this paper, the maximum likelihood and Bayesian methods are used for estimating parameter of Bernoulli distribution, i.e. θ , which is defined as the probability of success event for two possible outcomes. The maximum likelihood and Bayesian estimators of Bernoulli parameter are derived, for the Bayesian estimator the Beta prior is used. The analytical calculation shows that maximum likelihood estimator is unbiased while Bayesian estimator is asymptotically unbiased. However, empirical analysis by Monte Carlo simulation shows that the mean square errors (MSE) of the Bayesian estimator are smaller than maximum likelihood estimator for large sample sizes.

Keywords: Bernoulli distribution, beta distribution, conjugate prior, parameter estimation.

1. PENDAHULUAN

Parameter estimation is a way to predict the characteristics of a population based on the sample taken. In general, parameter estimation is classified into two types, namely point estimation and interval estimation. The point estimation of a parameter is a value obtained from the sample and is used as a parameter estimator whose value is unknown.

Several point estimation methods are used to calculate the estimator, such as moment method, maximum likelihood method, and Bayesian method. The moment method predicts the parameters by equating the values of sample moments to the population moment and solving the resulting equation system [1]. The maximum likelihood (ML) method uses differential calculus to determine the maximum of the likelihood function to obtain the parameters estimates. The Bayesian method differs from the traditional methods by introducing a frequency function for the parameter being estimated namely prior distribution. The Bayesian method combines the prior distribution and sample distribution. The prior distribution is the initial distribution that provides information about the parameters. The sample distribution combined with the prior distribution provides a new distribution i.e. the posterior distribution that expresses a degree of confidence regarding the location of the parameters after the sample is observed [2].

Researches on parameter estimation using various methods of various distributions have been done, for example: Bayesian estimation of exponential distribution [3], [4], ML and Bayesian estimations of Poisson distribution [5], Bayesian estimation of Poisson-Exponential distribution [6], and Bayesian estimation of Rayleigh distribution [7].

The difference between the ML and the Bayesian methods is that the ML method considers that the parameter is

an unknown quantity of fixed value and the inference is based only on the information in the sample; while the Bayesian method considers the parameter as a variable that describes the initial knowledge of the parameters before the observation is performed and expressed in a distribution called the prior distribution. After the observation is performed, the information in the prior distribution is combined with the sample data information through Bayesian theorem, and the result is expressed in a distribution form called the posterior distribution, which further becomes the basis for inference in the Bayesian method [8].

The Bayesian method has advantages over other methods, one of which is the Bayesian method can be used for drawing conclusions in complicated or extreme cases that cannot be handled by other methods, such as in complex hierarchical models. In addition, if the prior information does not indicate complete and clear information about the distribution of the prior, appropriate assumptions may be given to its distribution characteristics. Thus, if the prior distribution can be determined, then a posterior distribution can be obtained which may require mathematical computation [8].

This paper examines the parameter estimation of Bernoulli distribution using ML and Bayesian methods. A review of Bernoulli distribution and Beta distribution is presented in Section 2. The research methodology is described in Section 3. Section 4 provides the results and discussion. Finally, the conclusion is given in Section 5.

2. THEORETICAL FRAMEWORK

2.1 Bernoulli Distribution

Bernoulli distribution was introduced by Swiss mathematician, Jacob Bernoulli (1654-1705). It is the probability distribution resulting from two outcomes or events in a given experiment, i.e. success ($X = 1$) and fail ($X = 0$), with the probability of the success is θ and the probability of failure is $1 - \theta$.

Definition

A random variable X is called a Bernoulli random variable (or X is Bernoulli distributed) if and only if its probability distribution is given by

$$f(x; \theta) = \theta^x (1 - \theta)^{1-x}, \text{ for } x = 0, 1.$$

Proposition 1

Bernoulli distribution $f(x; \theta)$ has mean and variance as follows:

$$\mu = \theta \text{ and } \sigma^2 = \theta(1 - \theta).$$

Proof :

The mean of Bernoulli random variable X is

$$\begin{aligned}
 \mu &= E(X) \\
 &= \sum_{x=0}^1 xf(x; \theta) \\
 &= \sum_{x=0}^{x=1} x\theta^x(1-\theta)^{1-x} \\
 &= 0 \cdot \theta(1-\theta)^{1-0} + 1 \cdot (1-\theta)^{1-1} = \theta.
 \end{aligned}$$

The variance, i.e. $\sigma^2 = E(X - \mu)^2 = E(X^2) - [E(X)]^2$, of Bernoulli distribution is obtained as follows:

$$\begin{aligned}
 E(X^2) &= \sum_{x=0}^1 x^2 f(x; \theta) \\
 &= \sum_{x=0}^1 x^2 f(x; \theta) \\
 &= \sum_{x=0}^1 x^2 \theta(1-\theta)^{1-x} \\
 &= 0^2 \cdot \theta(1-\theta)^{1-0} + 1^2 \cdot (1-\theta)^{1-1} = \theta.
 \end{aligned}$$

Then,

$$\sigma^2 = E(X - \mu)^2 = \theta - \theta^2 = \theta(1 - \theta).$$

2.2. Beta Distribution

Definition

A random variable X is called a betarandom variable with parameters a and b if the density function of X is given by

$$f(x) = \begin{cases} \frac{1}{B(a, b)} x^{a-1} (1-x)^{b-1}, & 0 < x < 1 \\ 0, & \text{lainnya} \end{cases}$$

where $B(a, b)$ is betafunction defined as

$$B(a, b) = \int_0^1 x^{a-1} (1-x)^{b-1} dx ; a > 0, b > 0. \quad (1)$$

Proposition 2

The beta function and gamma function is connected by

$$B(a, b) = \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}. \quad (2)$$

Proof :

$$\Gamma(a)\Gamma(b) = \int_{x=0}^{\infty} x^{a-1} e^{-x} dx \cdot \int_{y=0}^{\infty} y^{b-1} e^{-y} dy$$

$$= \int_{y=0}^{\infty} \int_{x=0}^{\infty} x^{a-1} y^{b-1} e^{-x-y} dx dy.$$

Let $x = f(z, t) = zt$ and $y = g(z, t) = z(1-t)$,

$$\begin{aligned} \Gamma(a)\Gamma(b) &= \int_{z=0}^{\infty} \int_{t=0}^1 (zt)^{a-1} [z(1-t)]^{b-1} e^{-z} |J(z, t)| dt dz \\ &= \int_{z=0}^{\infty} \int_{t=0}^1 (zt)^{a-1} [z(1-t)]^{b-1} e^{-z} z dt dz \\ &= \int_{z=0}^{\infty} \int_{t=0}^1 z^{a-1+b-1+1} e^{-z} t^{a-1} (1-t)^{b-1} dt dz \\ &= \int_{z=0}^{\infty} z^{a+b-1} e^{-z} dz \cdot \int_{t=0}^1 t^{a-1} (1-t)^{b-1} dt \\ &= \Gamma(a+b)B(a, b). \end{aligned}$$

Then,

$$B(a, b) = \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}.$$

Proposisi 3

The mean and variance of beta distribution with parameters a and b are

$$\mu = \frac{a}{a+b} \quad \text{and} \quad \sigma^2 = \frac{ab}{(a+b+1)(a+b)^2}.$$

Proof :

The proposition can be proved by using the moment of beta distribution as follows:

$$\begin{aligned} E(X^n) &= \frac{1}{B(a, b)} \int_0^1 x^n x^{a-1} (1-x)^{b-1} dx \\ &= \frac{1}{B(a, b)} \int_0^1 x^{(a+n)-1} (1-x)^{b-1} dx. \end{aligned}$$

From equations (1) and (2) we obtain

$$\begin{aligned} E(X^n) &= \frac{B(a+n, b)}{B(a, b)} \\ &= \frac{\frac{\Gamma(a+n)\Gamma(b)}{\Gamma(a+b+n)}}{\frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}} \end{aligned}$$

$$\begin{aligned}
 &= \frac{\Gamma(a+n)\Gamma(b)}{\Gamma(a+b+n)} \times \frac{\Gamma(a+b)}{\Gamma(a)\Gamma(b)} \\
 &= \frac{\Gamma(a+n)\Gamma(a+b)}{\Gamma(a+b+n)\Gamma(a)}.
 \end{aligned} \tag{3}$$

Thus the mean and variance of beta distribution will be obtained by substituting $n = 1$ and $n = 2$ to equation (3), then

$$\begin{aligned}
 \text{Mean}(X) = E(X^1) &= \frac{\Gamma(a+1)\Gamma(a+b)}{\Gamma(a+b+1)\Gamma(a)} \\
 &= \frac{a\Gamma(a)\Gamma(a+b)}{(a+b)\Gamma(a+b)\Gamma(a)} \\
 &= \frac{a}{a+b}
 \end{aligned}$$

and $\text{Var}(X) = \sigma^2 = E(X^2) - [E(X)]^2$.

Since

$$\begin{aligned}
 E(X^2) &= \frac{\Gamma(a+2)\Gamma(a+b)}{\Gamma(a+b+2)\Gamma(a)} \\
 &= \frac{(a+1)\Gamma(a+1)\Gamma(a+b)}{(a+b+1)\Gamma(a+b+1)\Gamma(a)} \\
 &= \frac{(a+1)a\Gamma(a)\Gamma(a+b)}{(a+b+1)(a+b)\Gamma(a+b)\Gamma(a)} \\
 &= \frac{(a+1)a}{(a+b+1)(a+b)},
 \end{aligned}$$

then

$$\begin{aligned}
 \text{Var}(X) &= \frac{(a+1)a}{(a+b+1)(a+b)} - \left(\frac{a}{a+b}\right)^2 \\
 &= \frac{(a+1)a}{(a+b+1)(a+b)} - \frac{a^2}{(a+b)^2} \\
 &= \frac{(a+b)(a^2+a) - a^2(a+b+1)}{(a+b)^2(a+b+1)} \\
 &= \frac{a^3 + a^2b + a^2 + ab - a^3 - a^2b - a^2}{(a+b)^2(a+b+1)} \\
 &= \frac{ab}{(a+b)^2(a+b+1)}.
 \end{aligned}$$

3. RESEARCH METHOD

The research method for estimating the parameter of Bernoulli distribution in this paper can be described as follows. For ML estimation, the parameter estimation is done by differentiating partially the log of the likelihood

function and equation it by zero,

$$\frac{\partial \ln L(\theta)}{\partial \theta} = 0$$

to obtain ML estimator ($\hat{\theta}_{ML}$). The second derivation assessment is performed to show that the resulted $\hat{\theta}$ truly maximize the likelihood function. For the Bayesian method, the parameter estimation is done through the following steps:

1. Form the likelihood function of Bernoulli distribution as follows:

$$L(x_1, x_2, \dots, x_n | \theta) = \prod_{i=1}^n f((x_i) | \theta).$$

2. Calculate the joint probability distribution, which is obtained by multiplying the likelihood function and the prior distribution,

$$H(x_1, x_2, \dots, x_n; \theta) = L(x_1, x_2, \dots, x_n | \theta) \cdot \pi(\theta).$$

3. Calculate the marginal probability distribution function,

$$p(x_1, x_2, \dots, x_n) = \int H(x_1, x_2, \dots, x_n; \theta) d\theta.$$

4. Calculate the posterior distribution by dividing the joint probability distribution function by the marginal function,

$$\pi(\theta | x_1, x_2, \dots, x_n) = \frac{H(x_1, x_2, \dots, x_n; \theta)}{p(x_1, x_2, \dots, x_n)}.$$

The Bayesian parameter estimate of θ is then produced as the mean of the posterior distribution.

After the parameter estimate of θ is obtained by MLE and Bayesian methods, the evaluation of the estimators is performed by assessing their bias, variance, and mean square error.

4. RESULT AND DISCUSSION

4.1. The ML Estimator of the Bernoulli Distribution Parameter (θ)

Let $X_1, X_2, \dots, X_n$ be Bernoulli distributed random sample with $X_i \sim \text{Bernoulli}(\theta)$, where $\theta \in \Omega = (0,1)$. The probability function of X_i is

$$f(x_i; \theta) = \theta^{x_i} (1 - \theta)^{1-x_i} \text{ with } x_i \in \{0,1\}.$$

The likelihood function of Bernoulli distribution is given by

$$\begin{aligned} L(\theta) &= f(x_1, x_2, \dots, x_n; \theta) \\ &= \prod_{i=1}^n f(x_i; \theta) \\ &= \prod_{i=1}^n \theta^{x_i} (1 - \theta)^{1-x_i} \end{aligned}$$

$$\begin{aligned}
 &= \theta^{x_1}(1-\theta)^{1-x_1} \cdot \theta^{x_2}(1-\theta)^{1-x_2} \dots \theta^{x_n}(1-\theta)^{1-x_n} \\
 &= \theta^{\sum_{i=1}^n x_i} (1-\theta)^{n-\sum_{i=1}^n x_i}.
 \end{aligned} \tag{4}$$

The natural logarithm of the likelihood function is then

$$\begin{aligned}
 \ln L(\theta) &= \ln[\theta^{\sum_{i=1}^n x_i} (1-\theta)^{n-\sum_{i=1}^n x_i}] \\
 &= \ln \theta^{\sum_{i=1}^n x_i} + \ln(1-\theta)^{n-\sum_{i=1}^n x_i} \\
 &= \sum_{i=1}^n x_i \ln \theta + (n - \sum_{i=1}^n x_i) \ln(1-\theta).
 \end{aligned} \tag{5}$$

The ML estimate value of θ is obtained by differentiating equation (5) with respect to θ and equating the differential result to zero, i.e.

$$\begin{aligned}
 \frac{\partial}{\partial \theta} \ln L(\theta) &= \frac{\partial}{\partial \theta} \left[\sum_{i=1}^n x_i \ln \theta + \left(n - \sum_{i=1}^n x_i \right) \ln(1-\theta) \right] = 0 \\
 \frac{\sum_{i=1}^n x_i}{\theta} - \frac{n - \sum_{i=1}^n x_i}{1-\theta} &= 0 \\
 (1-\theta) \sum_{i=1}^n x_i - \theta \left(n - \sum_{i=1}^n x_i \right) &= 0 \\
 \sum_{i=1}^n x_i - \theta \sum_{i=1}^n x_i - n\theta + \theta \sum_{i=1}^n x_i &= 0 \\
 \sum_{i=1}^n x_i &= n\theta,
 \end{aligned}$$

then we obtain

$$\hat{\theta} = \frac{1}{n} \sum_{i=1}^n x_i.$$

To show that $\hat{\theta}$ is the value that maximizes the likelihood function $L(\theta)$, it must be confirmed that the second derivative of the likelihood function for $\theta = \hat{\theta}$ is negative:

$$\begin{aligned}
 \frac{\partial^2}{\partial \theta^2} \ln L(\theta) &= \frac{\partial^2}{\partial \theta^2} \left[\sum_{i=1}^n x_i \ln \theta + \left(n - \sum_{i=1}^n x_i \right) \ln(1-\theta) \right] \\
 &= \frac{-\sum_{i=1}^n x_i}{\theta^2} - \frac{n - \sum_{i=1}^n x_i}{(1-\theta)^2} \\
 &= \frac{-(1-\theta)^2 \sum_{i=1}^n x_i - \theta^2 (n - \sum_{i=1}^n x_i)}{\theta^2 (1-\theta)^2} \\
 &= \frac{-\sum_{i=1}^n x_i + 2\theta \sum_{i=1}^n x_i - \theta^2 \sum_{i=1}^n x_i - n\theta^2 + \theta^2 \sum_{i=1}^n x_i}{\theta^2 (1-\theta)^2}
 \end{aligned}$$

$$= \frac{-n\theta^2 + 2\theta \sum_{i=1}^n x_i - \sum_{i=1}^n x_i}{\theta^2(1-\theta)^2} < 0.$$

Since $\hat{\theta}$ maximizes the likelihood function, we conclude that the ML estimator of θ is given by

$$\hat{\theta}_{ML} = \frac{1}{n} \sum_{i=1}^n x_i.$$

4.2. The Bayesian Estimator of the Bernoulli Distribution Parameter(θ)

To estimate θ using Bayesian method, it is necessary to choose the initial information of a parameter called the prior distribution, denoted by $\pi(\theta)$, to be applied to the basis of the method namely the conditional probability. In this paper, the prior selection for Bernoulli distribution refers to the formation of its likelihood function. From equation (4) we have

$$\pi(\theta) \propto \theta^{\sum_{i=1}^n x_i} (1-\theta)^{1-\sum_{i=1}^n x_i}.$$

A distribution having probability function in the same form as the above expression is the beta distribution with density function

$$f(\theta; a, b) = \frac{1}{B(a, b)} \theta^{a-1} (1-\theta)^{b-1}, 0 < \theta < 1$$

where $a - 1 = \sum_{i=1}^n x_i$, $b - 1 = n - \sum_{i=1}^n x_i$, and $\frac{1}{B(a, b)}$ are factors required for the density function to be satisfied.

The prior distribution is combined with the sample distribution to produce a new distribution called posterior distribution and denoted by $\pi(\theta|x_1, x_2, \dots, x_n)$. Posterior distribution is obtained by dividing the joint density distribution by the marginal distribution.

Joint probability density function of $(x_1, x_2, \dots, x_n)$ is given by:

$$\begin{aligned} H(x_1, x_2, \dots, x_n; \theta) &= L(x_1, x_2, \dots, x_n|\theta) \cdot \pi(\theta) \\ &= \theta^{\sum_{i=1}^n x_i} (1-\theta)^{n-\sum_{i=1}^n x_i} \cdot \frac{1}{B(a, b)} \theta^{a-1} (1-\theta)^{b-1} \\ &= \frac{1}{B(a, b)} \theta^{a+\sum_{i=1}^n x_i-1} (1-\theta)^{b+n-\sum_{i=1}^n x_i-1} \end{aligned} \quad (6)$$

and the marginal function of $(x_1, x_2, \dots, x_n)$ is obtained as follows:

$$p(x_1, x_2, \dots, x_n) = \int_0^1 H(x_1, x_2, \dots, x_n; \theta) d\theta.$$

Using equation (6) we have

$$p(x_1, x_2, \dots, x_n) = \int_0^1 \frac{1}{B(a, b)} \theta^{a+\sum_{i=1}^n x_i-1} (1-\theta)^{b+n-\sum_{i=1}^n x_i-1} d\theta$$

$$\begin{aligned}
 &= \frac{1}{B(a, b)} \int_0^1 \theta^{a+\sum_{i=1}^n x_i-1} (1-\theta)^{b+n-\sum_{i=1}^n x_i-1} d\theta \\
 &= \frac{1}{B(a, b)} B(a + \sum_{i=1}^n x_i, b + n - \sum_{i=1}^n x_i).
 \end{aligned} \tag{7}$$

Then from equation (6) and (7) the posterior distribution can be written as follows:

$$\begin{aligned}
 \pi(\theta|x_1, x_2, \dots, x_n) &= \frac{H(x_1, x_2, \dots, x_n; \theta)}{p(x_1, x_2, \dots, x_n)} \\
 &= \frac{\frac{1}{B(a, b)} \theta^{a+\sum_{i=1}^n x_i-1} (1-\theta)^{b+n-\sum_{i=1}^n x_i-1}}{\frac{1}{B(a, b)} B(a + \sum_{i=1}^n x_i, b + n - \sum_{i=1}^n x_i)} \\
 &= \frac{\theta^{a+\sum_{i=1}^n x_i-1} (1-\theta)^{b+n-\sum_{i=1}^n x_i-1}}{B(a + \sum_{i=1}^n x_i, b + n - \sum_{i=1}^n x_i)}.
 \end{aligned} \tag{8}$$

The posterior distribution expressed in equation (8) is obviously following beta distribution also with parameter $(a + \sum_{i=1}^n x_i)$ and $(b + n - \sum_{i=1}^n x_i)$, or

$$\hat{\theta} \sim \text{Beta}(a + \sum_{i=1}^n x_i, b + n - \sum_{i=1}^n x_i).$$

Since the prior and posterior distribution of Bernoulli follows the same distribution, i.e. the Beta distribution, beta distribution is called as the conjugate prior of the Bernoulli distribution. The posterior mean is used as the parameter estimate θ in Bayesian method. Using Proposition 2, the Bayesian estimator of parameter θ is obtained as follows:

$$\begin{aligned}
 \hat{\theta}_B &= \frac{a + \sum_{i=1}^n x_i}{a + \sum_{i=1}^n x_i + b + n - \sum_{i=1}^n x_i} \\
 &= \frac{a + \sum_{i=1}^n x_i}{a + b + n}.
 \end{aligned}$$

4.3. Evaluation of the Estimators Properties

The parameter estimation of the Bernoulli distribution is obtained by the MLE and Bayesian methods yields different estimates. The best estimator has to meet the following properties:

1. Unbiased

An estimator is called to be unbiased if its expected values is equal to the estimated parameter, i.e. $\hat{\theta}$ is an unbiased estimator of θ if $E(\hat{\theta}) = \theta$. The bias of an estimator is then given by:

$$\text{Bias}(\hat{\theta}) = E(\hat{\theta}) - \theta. \tag{9}$$

Let $X_1, X_2, \dots, X_n$ are Bernoulli(θ) random sample observations. Since $\hat{\theta}_{ML} = \frac{1}{n} \sum_{i=1}^n x_i$ is the ML estimator of θ , its expected value is as follows:

$$\begin{aligned}
 E(\hat{\theta}_{ML}) &= E\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \\
 &= \frac{1}{n} E\left(\sum_{i=1}^n x_i\right) \\
 &= \frac{1}{n} \sum_{i=1}^n E(x_i) \\
 &= \frac{1}{n} \cdot n\theta = \theta.
 \end{aligned} \tag{10}$$

Since $E(\hat{\theta}_{ML}) = \theta$, $\hat{\theta}_{MLE}$ is an unbiased estimator of θ .

Now consider the Bayesian estimator of θ i.e. $\hat{\theta}_B = \frac{a + \sum_{i=1}^n x_i}{a + b + n}$. The expected value of Bayesian estimator is given by

$$\begin{aligned}
 E(\hat{\theta}_B) &= E\left(\frac{a + \sum_{i=1}^n x_i}{a + b + n}\right) \\
 &= \frac{1}{a + b + n} E\left(a + \sum_{i=1}^n x_i\right) \\
 &= \frac{1}{a + b + n} \left[E(a) + E\left(\sum_{i=1}^n x_i\right) \right] \\
 &= \frac{1}{a + b + n} \left[E(a) + \sum_{i=1}^n E(x_i) \right] \\
 &= \frac{1}{a + b + n} (a + n\theta).
 \end{aligned} \tag{11}$$

Since $E(\hat{\theta}_B) \neq \theta$, $\hat{\theta}_B$ is a biased estimator of θ . The bias value of $\hat{\theta}_B$ is:

$$\begin{aligned}
 Bias(\hat{\theta}_B) &= E(\hat{\theta}_B) - \theta \\
 &= \frac{a + n\theta}{a + b + n} - \theta.
 \end{aligned} \tag{12}$$

Although $\hat{\theta}_B$ is a biased estimator of θ , it can be shown that $\hat{\theta}_B$ is asymptotically unbiased. The proof is given as follows:

$$\begin{aligned}
 \lim_{n \rightarrow \infty} E(\hat{\theta}_B) &= \lim_{n \rightarrow \infty} \frac{a + n\theta}{a + b + n} \\
 &= \lim_{n \rightarrow \infty} \frac{\frac{a}{n} + \frac{n\theta}{n}}{\frac{a}{n} + \frac{b}{n} + \frac{n}{n}} \\
 &= \lim_{n \rightarrow \infty} \frac{\frac{a}{n} + \theta}{\frac{a}{n} + \frac{b}{n} + 1}
 \end{aligned}$$

$$= \frac{\theta}{1} = \theta. \quad (13)$$

Since $\lim_{n \rightarrow \infty} E(\hat{\theta}_B) = \theta$, $\hat{\theta}_B$ is an asymptotically unbiased estimator of θ .

2. Efficiency

The efficiency of an estimator is observed from its variance. The best parameter estimator is the one that has the smallest variance. This is because the variance of an estimator is a measure of the spread of the estimator around its mean.

The variance of ML estimator $\hat{\theta}_{ML}$ is:

$$\begin{aligned} Var(\hat{\theta}_{ML}) &= Var\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \\ &= \frac{1}{n^2} Var\left(\sum_{i=1}^n x_i\right) \\ &= \frac{1}{n^2} \sum_{i=1}^n Var(x_i) \\ &= \frac{1}{n^2} n\theta(1 - \theta) \\ &= \frac{1}{n} \theta(1 - \theta). \end{aligned} \quad (14)$$

While the variance of the Bayesian estimator $\hat{\theta}_B$ is given by:

$$\begin{aligned} Var(\hat{\theta}_B) &= Var\left(\frac{a + \sum_{i=1}^n x_i}{a + b + n}\right) \\ &= \frac{1}{(a + b + n)^2} Var\left(a + \sum_{i=1}^n x_i\right) \\ &= \frac{1}{(a + b + n)^2} \left[Var(a) + \sum_{i=1}^n Var(x_i) \right]. \end{aligned}$$

Since $Var(a) = 0$ and $Var(x_i) = \theta(1 - \theta)$, we obtain

$$Var(\hat{\theta}_B) = \frac{1}{(a+b+n)^2} n\theta(1 - \theta). \quad (15)$$

From equation (10), it is shown that the ML estimator is unbiased, whereas from equations (11) and (12) it is shown that Bayesian estimator is biased. As a result, the efficiency of the two methods cannot be compared because the efficiency of estimators applies to unbiased estimators.

3. Consistency

The consistency of the estimators is evaluated from their mean square error (MSE). The MSE can be expressed as

$$MSE(\hat{\theta}) = E(\hat{\theta} - \theta)^2 = Var(\hat{\theta}) + (bias\hat{\theta})^2. \quad (16)$$

If the sample size grows infinitely, a consistent estimator will give a perfect point estimate to θ . Mathematically, θ is a consistent estimator if and only if

$$E(\hat{\theta} - \theta)^2 \rightarrow 0 \text{ when } n \rightarrow \infty,$$

which means that the bias and the variance approaches to 0 if $n \rightarrow \infty$.

Substituting equation (10) and (14) to equation (16), the MSE of ML estimator $\hat{\theta}_{MLE}$ is then

$$\begin{aligned} E(\hat{\theta}_{MLE} - \theta)^2 &= Var(\hat{\theta}_{MLE}) + (bias\hat{\theta}_{MLE})^2 \\ E(\hat{\theta}_{MLE} - \theta)^2 &= Var(\hat{\theta}_{MLE}) = \frac{1}{n}\theta(1 - \theta). \end{aligned}$$

For $n \rightarrow \infty$, we have

$$\lim_{n \rightarrow \infty} E(\hat{\theta}_{MLE} - \theta)^2 = \lim_{n \rightarrow \infty} \frac{1}{n}\theta(1 - \theta) = 0. \quad (17)$$

In the same manner, by substituting equation (12) and (15) the MSE of Bayesian estimator $\hat{\theta}_B$ is:

$$\begin{aligned} E(\hat{\theta}_B - \theta)^2 &= Var(\hat{\theta}_B) + (bias\hat{\theta}_B)^2 \\ E(\hat{\theta}_B - \theta)^2 &= \left[\frac{1}{(a+b+n)^2} n\theta(1 - \theta) \right] + \left(\frac{a+n\theta}{a+b+n} - \theta \right)^2. \end{aligned}$$

For $n \rightarrow \infty$, we have

$$\lim_{n \rightarrow \infty} (\hat{\theta}_B - \theta)^2 = \lim_{n \rightarrow \infty} \left[\left\{ \frac{1}{(a+b+n)^2} n\theta(1 - \theta) \right\} + \left(\frac{a+n\theta}{a+b+n} - \theta \right)^2 \right] = 0. \quad (18)$$

From equation (17) and (18), we can conclude that ML and Bayesian estimators are consistent estimators of θ .

4.4. Empirical Comparison of the Properties of ML and Bayesian Estimators

To compare the ML and Bayesian estimators of θ , a Monte Carlo simulation using R program was conducted. The simulation was performed by generating Bernoulli distributed data with $\theta = 0.1, 0.3$, and 0.5 and eight different sample sizes, i.e. $n = 20, 50, 100, 300, 500, 1000, 5000$, and 10000 . The simulation was repeated 1000 times for each combination of θ and n . The generated data were used to estimate parameter θ using the two methods. Furthermore, the bias and MSE of both estimators were calculated using the formulas in equations (9) and (16) and the results are presented in Table 1.

Table 1. The bias and MSE of ML and Bayesian estimators of θ

θ	N	Bias		MSE	
		ML	Bayesian ($\alpha = 1, \beta = 1$)	ML	Bayesian ($\alpha = 1, \beta = 1$)
0,1	20	0,001200	0,223084	0,031478	0,558149
	50	0,002180	0,036602	0,015264	0,041740
	100	0,000270	0,009328	0,007058	0,008843
	300	0,000413	0,001075	0,001901	0,002906
	500	0,000210	0,000364	0,001183	0,000551
	1000	0,000195	0,000091	0,000503	0,000128
	5000	0,000003	0,000003	0,000143	0,000004
	10000	0,000114	0,000001	0,000184	0,000002
0,3	20	0,003100	0,536609	0,001652	0,493287
	50	0,001300	0,090000	0,003113	0,142692
	100	0,000630	0,021341	0,001583	0,067615
	300	0,000003	0,002205	0,000327	0,002307
	500	0,000312	0,000845	0,000111	0,000899
	1000	0,000545	0,000207	0,000444	0,000217
	5000	0,000384	0,000008	0,000403	0,000018
	10000	0,000023	0,000002	0,000013	0,000004
0,5	20	0,001450	0,714234	0,023000	1,453419
	50	0,002260	0,097036	0,011566	0,147045
	100	0,003860	0,024341	0,008602	0,024833
	300	0,000143	0,002856	0,001792	0,003374
	500	0,000146	0,001004	0,001139	0,003040
	1000	0,000432	0,000252	0,000929	0,000253
	5000	0,000132	0,000009	0,000232	0,000056
	10000	0,000066	0,000002	0,000016	0,000003

Table 1 shows the bias and MSE values of ML and Bayesian estimates for a successful probability of $\theta = 0.1, 0.3$ and 0.5 . From the table it can be seen ML estimator produces smaller biases than Bayesian estimates for finite sample (i.e. $n < 1000$). However, when the sample size equal or larger than 1000 (i.e. 5000 and 10000, the biases of the Bayesian estimator are smaller than the ML estimator. Even though the bias values of ML estimates changes inconsistently throughout the sample sizes, analytically it has been proved that ML estimator is an unbiased estimator. This appears to be different from the bias values for Bayesian estimator. This is because for all the considered success probabilities of the bias values become smaller when the sample size increases, although analytically it is found that Bayesian estimator is a biased estimator. As a result the efficiency of the two estimators cannot be compared. Therefore, to compare the best estimators we use MSE of both estimators. This is because MSE considers both the bias and variance values.

The MSE values of ML and Bayesian estimators that have been shown in Table 1 have similarities, i.e. the MSE value decreases as the sample size increases and it closes to 0. Thus, both estimators are consistent estimators. This also corresponds to the results obtained analytically. Based on the simulation results in this study, it can be seen that for the larger sample sizes Bayesian estimator is better than ML estimator. This is because the MSE value of Bayesian estimator is smaller than the ML estimator. As shown in Table 1, when $\theta = 0.1$, the MSE value

of the Bayesian estimator is smaller than the ML estimator for $n = 500, 1000$, and 10000 ; and when $\theta = 0.3$ and 0.5 , the MSE values of the Bayesian estimator are smaller than the ML estimator for $n = 1000, 5000$, and 10000 .

5. CONCLUSION

In this paper, we derived the ML and Bayesian estimator (using beta prior) of Bernoulli distribution parameter. Analytically we show that the ML estimator is an unbiased estimator and Bayesian estimator is a biased estimator for parameter θ . However, Bayesian estimator is asymptotically unbiased. Based on the simulation result, both ML and Bayesian estimator are consistent estimators of θ because the two estimators satisfy the property of consistency, i.e. $E(\hat{\theta} - \theta)^2 \rightarrow 0$ when $n \rightarrow \infty$. The simulation result also shows that the Bayesian estimator using beta prior is better than the MLE method for large sample sizes ($n \geq 1000$).

REFERENCES

- [1]. Bain, L.J. and Engelhardt, M. (1992). *Introduction to Probability and Mathematical Statistics*. Duxbury Press, California.
- [2]. Walpole, R.E dan Myers, R.H. (1995). *Ilmu Peluang dan Statistika untuk Insinyur dan Ilmuwan*. ITB, Bandung.
- [3]. Al-Kutubi H. S., Ibrahim N.A. (2009). Bayes Estimator for Exponential Distribution with Extension of Jeffery Prior Information. *Malaysian Journal of Mathematical Sciences*. 3(2):297-313.
- [4]. Nurlaila, D., Kusnandar D., & Sulistianingsih, E. (2013). Perbandingan Metode *Maximum Likelihood Estimation* (MLE) dan Metode Bayes dalam Pendugaan Parameter Distribusi Eksponensial. *Buletin Ilmiah Mat. Stat. dan Terapannya*. 2(1):51-56.
- [5]. Fikhri, M., Yanuar, F., & Yudiantri A. (2014). Pendugaan Parameter dari Distribusi Poisson dengan Menggunakan Metode *Maximum Likelihood Estimation* (MLE) dan Metode Bayes. *Jurnal Matematika UNAND*. 3(4):152-159.
- [6]. Singh S. K., Singh, U., & Kumar, M. (2014). Estimation for the Parameter of Poisson-Exponential Distribution under Bayesian Paradigm. *Journal of Data Science*. 12:157-173.
- [7]. Gupta, I. (2017). Bayesian and E-Bayesian Method of Estimation of Parameter of Rayleigh Distribution-A Bayesian Approach under Linex Loss Function. *International Journal of Statistics and Systems*. 12(4):791-796.
- [8]. Box, G.E.P & Tiao, G.C. (1973). *Bayesian Inference in Statistical Analysis*. Addison-Wesley Publishing Company, Philippines.

Proses Pengamanan Data Menggunakan Kombinasi Metode Kriptografi Data *Encryption Standard dan Steganografi End Of File*

Dedi Darwis^{*1) **1)}, Wamiliana²⁾, Akmal Junaidi³⁾

^{*1)} Program Studi Sistem Informasi, Universitas Teknokrat Indonesia

^{**1)} Mahasiswa S3 Doktor Ilmu MIPA, Universitas Lampung

E-mail : darwisdedi@teknokrat.ac.id

²⁾ Jurusan Matematika, FMIPA, Universitas Lampung

³⁾ Jurusan Ilmu Komputer, FMIPA, Universitas Lampung

E-mail : wamiliana.1963@unila.ac.id

ABSTRAK

Metode Data Encryption Standard (DES) merupakan salah satu metode Kriptografi yang dapat digunakan sebagai alternatif pengamanan data seperti keamanan password, keamanan jaringan ataupun pengiriman pesan rahasia dalam pertukaran informasi, namun hasil kriptografi hanya dalam bentuk kode penyandian kata sehingga menimbulkan kecurigaan orang lain untuk memecahkan kode penyandian tersebut. Pada penelitian ini dikembangkan kombinasi kriptografi yang lebih aman sehingga pesan rahasia dalam bentuk hasil enkripsi yang akan dikirim disisipkan terlebih dahulu ke media digital seperti image dan media lainnya yang disebut teknik steganografi, sehingga pihak lain tidak akan curiga bahwa dalam media digital tersebut terdapat pesan rahasia. Metode yang digunakan dalam penelitian ini untuk steganografi adalah End Of File (EOF) karena metode ini mampu mempertahankan kualitas citra yang baik karena antara cover image dan stego image memiliki kemiripan yang signifikan sehingga tidak menimbulkan kecurigaan. Berdasarkan hasil pengujian yang dilakukan pesan rahasia berhasil dienkripsi dengan menggunakan panjang kunci 56 bit kemudian disisipkan pada citra digital menggunakan EOF dan memiliki hasil pengujian bahwa pesan yang disisipkan dapat diambil kembali untuk dapat dilakukan proses dekripsi.

Kata kunci: EOF, DES, Kriptografi, Steganografi

1. PENDAHULUAN

Data dan Informasi merupakan hal yang sangat penting untuk dijaga keamanannya atau kerahasiannya agar pihak-pihak yang tidak berkompeten terhadap data tersebut misalkan Laporan Keuangan, keamanan password, dan keamanan jaringan ataupun data-data penting lainnya. Salah satu cara untuk mengamankan data adalah menggunakan teknik Kriptografi yaitu ilmu yang mempelajari teknik – teknik matematika yang berhubungan dengan aspek keamanan informasi seperti kerahasiaan, integritas data, serta otentikasi[1]. Pada implementasinya banyak sekali metode kriptografi yang dapat digunakan, dari mulai kriptografi klasik sampai kriptografi modern. Salah satu metode yang cukup populer adalah *Data Encryption Standard* (DES) karena dapat mengatasi masalah-masalah yang terjadi seperti pencurian data, penyalahgunaan data, dan kerahasiaan data[2]. Namun saat ini penggunaan teknik kriptografi masih belum cukup dalam melakukan proses pengamanan data karena masih menimbulkan kecurigaan karena teks yang disandikan masih terlihat walaupun dalam bentuk-bentuk simbol-simbol sehingga membuat orang lain ingin memecahkan sandi tersebut, maka dari itu diperlukan teknik lain yang dapat dikombinasikan dengan teknik kriptografi.

Salah satu teknik yang dapat dikombinasikan adalah steganografi yang merupakan suatu cabang ilmu yang mempelajari tentang bagaimana menyembunyikan suatu informasi rahasia di dalam suatu informasi lainnya[3]. Sama halnya seperti kriptografi, steganografi juga memiliki banyak metode yang dapat digunakan, salah satunya adalah metode *End Of File (EOF)* di mana metode ini berfokus pada ukuran suatu citra digital untuk menyisipkan pesan pada *file* yang terakhir. Pada metode *End Of File*, data yang telah dienkripsi akan disisipkan pada nilai akhir file gambar, sehingga hanya menambah ukuran file dan terdapat penambahan garis-garis pada bagian bawah file gambar tersebut sehingga tidak banyak merubah kualitas citra aslinya[4].

Pada penelitian ini akan dilakukan kombinasi dengan menggabungkan teknik kriptografi menggunakan metode *Data Encryption Standard* dan teknik steganografi *End Of File* dengan cara informasi rahasia dalam bentuk *text* dienkripsi terlebih dahulu, lalu hasil enkripsi akan dimasukkan ke dalam sebuah citra digital menggunakan format *JPEG*. Diharapkan dengan penggabungan dua metode ini dapat meningkatkan keamanan data pada pesan-pesan rahasia sehingga pihak yang tidak berkompeten tidak akan dapat mengetahui kerahasiaan informasi yang disimpan.

2. LANDASAN TEORI

2.1 Tinjauan Pustaka

1. Pada [5] digunakan Algoritma *Data Encryption Standard (DES)* dan dibangun menggunakan bahasa pemrograman Java untuk menerapkan metode DES dengan kunci simetri sepanjang 64 bit. Hasil dari penelitian ini adalah dengan adanya aplikasi kriptografi yang dikembangkan berdasarkan algoritma DES, maka data-data penting dapat diamankan (dienkripsi) ketika hendak dikirim melalui media internet [5].
2. Pada [2] masalah diterapkan Algoritma DES pada keamanan data pesan text dilakukan dengan cara permutasi pada blok plainteks permutasi awal kemudian di-*enciphering* sebanyak 16 kali putaran serta dalam merancang keamanan data pesan teks dilakukan perancangan input, proses, *output* dan salah satunya dapat melakukan program *visual basic 6.0* dan dapat melakukan *initial permutation* kemudian dipermutasikan dengan matriks permutasi balikan (*invers initial permutation*) menjadi blok *Ciphertext*.
3. Pada [6] digunakan. Metode *End Of File* untuk menyembunyikan pesan rahasia ke dalam citra digital dengan format ekstensi berupa *Jpg* dan selanjutnya pesan akan disisipkan pada file terakhir pada sebuah citra, dengan metode ini ukuran file jumlahnya tidak terbatas. Hasil dari penelitian ini menunjukkan bahwa pesan teks yang disisipkan ke dalam file citra dapat diambil kembali dari citra tersebut, kemudian berdasarkan pengujian yang dilakukan Pengujian *imperectibility* memberikan hasil steganografi pada gambar, dengan metode kuesioner yang menghasilkan 70% mahasiswa dibidang komputer tidak mengetahui tentang Steganografi dan 100% menyatakan gambar hasil steganografi tidak dapat terlihat oleh indra mata manusia secara kasat mata. Pada proses pengujian tahap fidelity tidak nampak nilai MSE yang hanya menghasilkan nilai "0" dan PSNR menghasilkan nilai " ∞ " (tak hingga) dikarenakan metode yang digunakan menyisipkan pesan di akhir file tanpa merubah nilai intensitas warna pikselnya dan Pembuatan aplikasi steganografi dapat diterapkan dan dijalankan dengan *mobile smartphone android*.
4. Pada [4] metode *Cryptosystem* digunakan untuk mengenkripsi data atau pesan rahasia yang berupa teks angka dengan jumlah maksimum yang dimasukkan adalah 24 digit angka kemudian hasil enkripsi (*Ciphertext*) akan disembunyikan ke dalam suatu file gambar yang berformat bitmap dengan ukuran minimum 25x25. Selanjutnya, dilakukan proses ekstraksi dan dekripsi *Ciphertext*, sehingga diperoleh kembali *Plaintext* yang berupa data teks angka. Hasil pengujian menunjukkan Algoritma Rabin Public Key tidak aman untuk serangan chosen-*Ciphertext* attack karena untuk kombinasi *Plaintext* dan kunci yang merupakan angka kelipatan "11" akan menghasilkan *Ciphertext* yang

merupakan angka kelipatan “11” juga. Sehingga, seorang kriptanalis dapat mengetahui bentuk *Plaintext* yang sebenarnya. sedangkan pada metode *End Of File*, data yang telah dienkripsi akan disisipkan pada nilai akhir file gambar, sehingga akan menambah ukuran file dan terdapat penambahan garis-garis pada bagian bawah file gambar tersebut.

2.2 Konsep Kriptografi

Kriptografi yaitu berasal dari bahasa Yunani, *Crypto* dan *graphia*. *Crypto* berarti *secret* (rahasia) dan *graphia* berarti *writing* (tulisan). Menurut terminologinya, kriptografi adalah ilmu dan seni untuk menjaga keamanan pesan ketika pesan dikirim dari suatu tempat ke tempat yang lain[7].

Algoritma dalam kriptografi di bagi menjadi dua, yaitu :[7]

1. Algoritma simetris adalah algoritma yang menggunakan kunci yang sama untuk proses enkripsi dan proses dekripsi. Adapun contoh algoritma kunci simetris adalah DES (*Data Encryption Standard*), Blowfish, Twofish, MARS, IDEA, 3DES (DES diaplikasikan 3 kali), AES (*Advanced Encryption Standard*) yang bernama asli Rijndael, *Vigenere*, dan lain - lain.
2. Algoritma asimetris adalah algoritma yang menggunakan kunci yang berbeda untuk proses enkripsi dan dekripsi. Adapun contoh algoritma yang menggunakan kunci asimetris adalah RSA (*Riverst Shamir Adleman*) dan ECC (*Elliptic Curve Cryptography*).

2.3 Algoritma Data Encryption Standard

Standar Enkripsi Data (*Data Encryption Standard* – DES) merupakan algoritma enkripsi yang paling banyak dipakai di dunia, yang diadopsi oleh NIST (*National Institute of Standards and Technology*) sebagai standar pengolahan informasi Federal AS. Secara umum standar enkripsi data terbagi menjadi tiga kelompok yaitu pemrosesan kunci, enkripsi data 64 bit dan dekripsi data 64 bit yang mana satu kelompok saling berinteraksi satu sama lain[7].

Pada akhir 1960 IBM memulai riset proyek *Lucifer* yang dipimpin oleh Horst Feistel untuk kriptografi komputer. Proyek ini berakhir tahun 1971 dan *Lucifer* pertama kali dikenal sebagai blok kode pada pengoperasian blok 64 bit dan menggunakan ukuran kunci 128 bit. Setelah IBM mengembangkan sistem enkripsi yang dikomersilkan maka *Lucifer* disebut dengan DES (*Data Encryption Standard*). Proyek ini dipimpin oleh Walter Tuchman. Hasil dari riset ini merupakan versi *Lucifer* yang bersifat menentang pemecahan algoritma kriptografi.

DES merupakan sistem kriptografi simetri dan tergolong jenis blok kode. DES beroperasi pada ukuran blok 64 bit. DES mengenkripsikan 64 bit teks asli menjadi 64 bit teks kode dengan menggunakan 56 bit kunci internal (*internal key*) atau upa-kunci (*subkey*). Kunci internal dibangkitkan dari kunci eksternal (*external key*) yang panjangnya 64 bit[7].

Skema global dari algoritma DES adalah sebagai berikut :

1. Blok teks-asli dipermutasi dengan matrik permutasi awal (*initial permutatiaon* atau IP). Bisa ditulis $x_0 = IP(x) = LOR_0$, dimana L_0 terdiri dari 32 bit pertama dari x_0 dan 32 bit terakhir dari R_0 .
2. Hasil permutasi awal kemudian di-*enciphering* sebanyak 16 kali (16 putaran). Setiap putaran menggunakan kunci internal yang berbeda dengan perhitungan $L_i R_i$ $1 \leq i \leq 16$, dengan mengikuti aturan $L_i = R_i - IR_i = L_i - I \pm f(R_i - 1, K_i)$
3. Hasil enciphering kemudian dipermutasikan dengan matriks permutasi balik (*invers initialpermutation* atau IP-1) menjadi blok teks kode. IP-1 ke bitsring $R16L16$, memperoleh teks-kode y , kemudian $y = IP-1(R16L16)$.

2.4 Konsep Steganografi

Steganografi merupakan suatu cabang ilmu yang mempelajari tentang bagaimana menyembunyikan suatu informasi rahasia di dalam suatu informasi lainnya[3].

Steganography membutuhkan dua aspek yaitu media penyimpan dan informasi rahasia yang akan disembunyikan. Metode *steganography* sangat berguna jika digunakan pada *steganography* komputer karena banyak format *file digital* yang dapat dijadikan media untuk menyembunyikan pesan. *Steganography digital* menggunakan media *digital* sebagai wadah penampung, misalnya teks, citra, suara, dan *video*. Data rahasia yang disembunyikan juga dapat berupa teks, citra, suara, atau *video*.

Steganography memanfaatkan kekurangan-kekurangan sistem indera manusia seperti mata (*Human Visual System*) dan telinga (*Human Auditory System*), sehingga tidak diketahui kehadirannya oleh indera manusia (indera penglihatan atau indera pendengaran) dan mampu menghadapi proses-proses pengolahan sinyal *digital* dengan tidak merusak kualitas data yang telah disisipi sampai pada tahap tertentu.

Terdapat tiga aspek yang perlu diperhatikan dalam menyembunyikan pesan: kapasitas, keamanan, dan ketahanan. Kapasitas merujuk kepada besarnya informasi yang dapat disembunyikan oleh media, keamanan merujuk kepada ketidakmampuan pihak lain untuk mendeteksi keberadaan informasi yang disembunyikan, dan ketahanan merujuk kepada sejauh mana medium *steganography* dapat bertahan sebelum pihak lain menghancurkan informasi yang disembunyikan.

Perbedaan *steganography* dengan *cryptography* terletak pada bagaimana proses penyembunyian data dan hasil akhir dari proses tersebut. *Cryptography* melakukan proses pengacakan data aslinya sehingga menghasilkan data terenkripsi yang benar – benar acak dan berbeda dengan aslinya, sedangkan *steganography* menyembunyikan dalam data lain yang akan ditumpanginya tanpa mengubah data yang ditumpanginya tersebut sehingga data yang ditumpanginya sebelum dan setelah proses penyembunyian hamper sama.

2.5 Citra Digital

Citra adalah gambar pada bidang dua dimensi. Dalam tinjauan matematis, citra merupakan fungsi kontinu dari intensitas cahaya pada bidang dua dimensi. Di dalam komputer, citra digital disimpan sebagai suatu *file* dengan format tertentu. Contoh format citra digital adalah *.bmp*, *.jpg*, *.png*, *.gif* dan sebagainya. Ukuran citra digital dinyatakan dalam *pixel* (*picture element*)[8].

2.6 Algoritma End Of File

Metode ini pesan yang disisipkan jumlahnya tak terbatas. Akan tetapi efek sampingnya adalah ukuran file menjadi lebih besar dari ukuran semula. Ukuran file yang terlalu besar dari yang seharusnya, tentu akan menimbulkan kecurigaan bagi yang mengetahuinya. Oleh karena itu dianjurkan agar ukuran pesan dan ukuran citra yang digunakan proporsional [9].

Proses penyisipan pesan dengan metode EOF dapat dituliskan dalam algoritma sebagai berikut:

1. Inputkan pesan yang akan disisipkan.
2. Ubah pesan menjadi kode desimal.
3. Inputkan citra grayscale yang akan disisipi pesan.
4. Dapatkan nilai derajat keabuan masing-masing piksel.
5. Tambahkan kode desimal pesan sebagai nilai derajat keabuan diakhir citra.
6. Petakan menjadi citra baru.

Sedangkan ekstraksi pesan yang sudah disisipkan dengan metode EOF dapat dilakukan dengan algoritma berikut:

1. Inputkan image yang sudah mengandung pesan.
2. Dapatkan nilai derajat keabuan citra tersebut.
3. Ubah nilai tersebut menjadi karakter pesan.

2.7 PSNR (*Peak Signal to Noise Ration*) dan MSE (*Mean Square Error*)

PSNR merupakan parameter yang digunakan mengukur kualitas citra yang dihasilkan. Metode PSNR adalah ukuran perbandingan antara nilai piksel *cover image* dengan nilai piksel pada citra *stego* yang dihasilkan. Sebelum menentukan PSNR terlebih dahulu ditentukan nilai rata-rata kuadrat *absolute error* antara *cover image* dengan citra *stego* yaitu nilai MSE (*Mean Square Error*). Berikut ini rumus MSE untuk *cover image* berwarna[10]:

$$MSE_{AVG} = \frac{MSE_R + MSE_G + MSE_B}{X.Y} \quad (1)$$

Keterangan:

MSE_{AVG} = Nilai rata-rata MSE *cover image*.

MSE_R = Nilai MSE warna merah.

MSE_G = Nilai MSE warna hijau.
 MSE_B = Nilai MSE warna Biru.
 $X.Y$ = Dimensi gambar.

Berikut ini adalah rumus atau formula ang digunakan untuk menghitung PSNR [10]:

$$PSNR = 10_{\log 10} \left(\frac{255^2}{MSE} \right) \quad (2)$$

Keterangan:

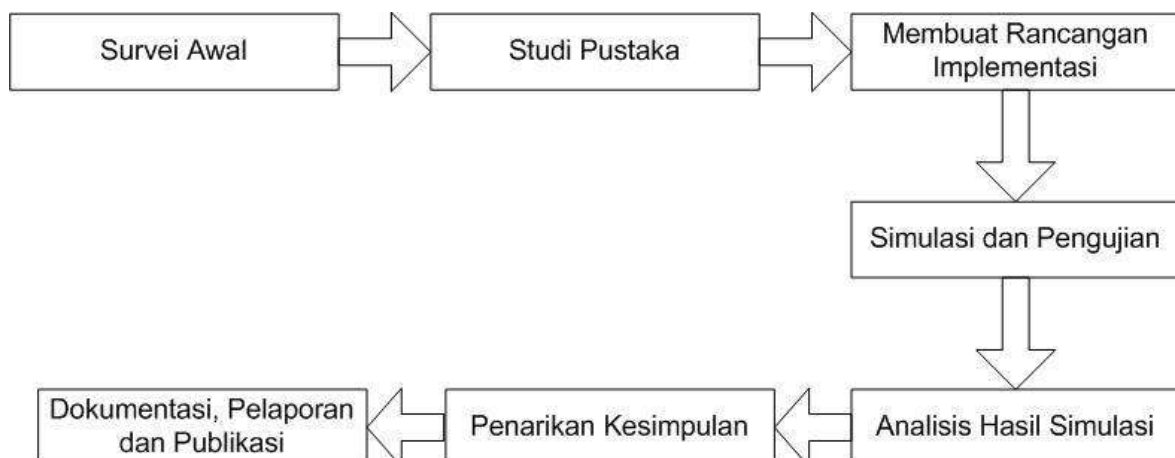
PSNR = Nilai PNSR citra digital.
MSE = Nilai *Mean Square Error* dari citra.

Citra *stego* dapat dikatakan berkualitas baik jika nilai PNSR dari citra *stego* tersebut bernilai tinggi. Terdapat sedikit perbedaan antara *cover image* dan citra *stego* setelah penambahan pesan rahasia. Tingkatan kualitas nilai PNSR berbanding terbalik dengan nilai MSE, semakin tinggi nilai PNSR semakin rendah nilai MSE. Semakin tinggi kualitas yang dihasilkan dari citra *stego* maka semakin rendah nilai dari MSE.

3. METODOLOGI PENELITIAN

3.1 Tahapan Penelitian

Penelitian ini terdiri dari beberapa tahap yaitu studi pustaka dan literature, Analisis Data, Metode dan Pemodelan Desain, Implementasi, Pengujian, kesimpulan dan publikasi yang ditunjukkan pada gambar 1.



Gambar 1. Tahapan Penelitian

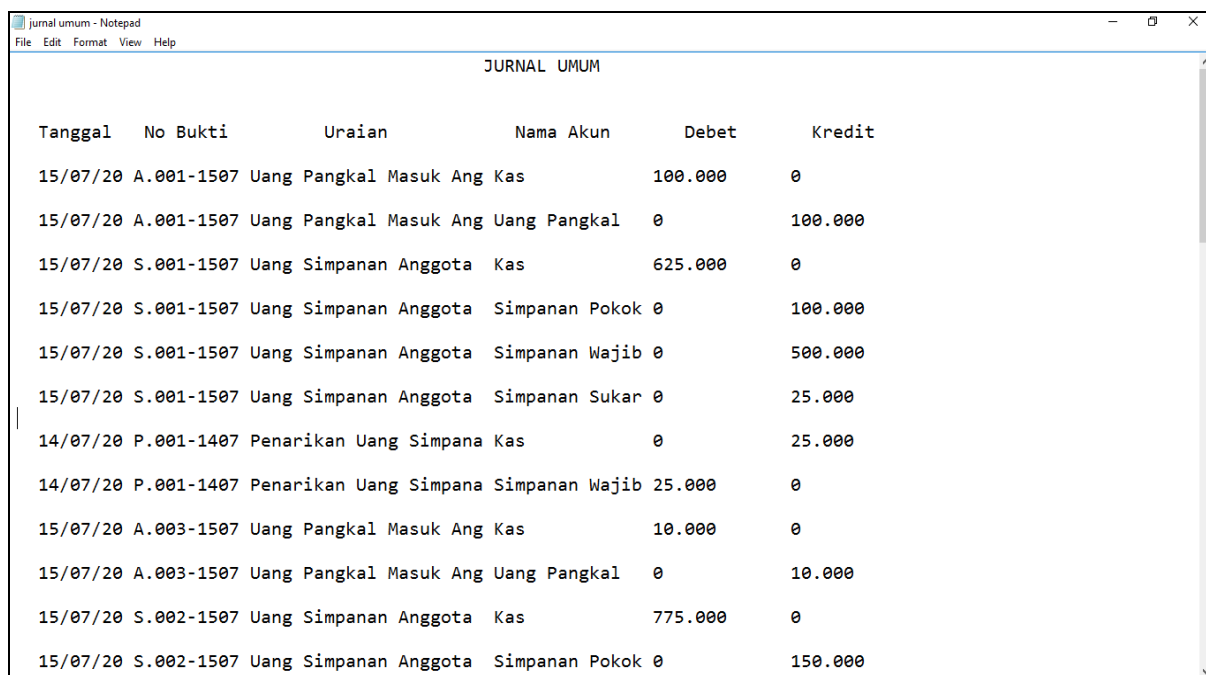
4. HASIL DAN PEMBAHASAN

4.1 Pengujian Enkripsi Kriptografi dan Steganografi

Pada proses pengujian kriptografi akan dilakukan proses pengujian pada aplikasi dari mulai file dalam bentuk *Plaintext* sampai dengan *Ciphertext* menggunakan metode *Data Encryption Standard* sampai proses *Encoding*/enkripsi steganografi menggunakan metode *End Of File*. Adapun proses pengujian enkripsi kriptografi dan steganografi adalah sebagai berikut :

4.1.1 Data Plaintext

Data sampel yang akan digunakan pada pengujian ini untuk proses pengamanan adalah jurnal laporan keuangan, karena data ini bersifat penting dan rahasia dalam bentuk format file txt. Adapun contoh data *Plaintext* yang digunakan dapat dilihat pada gambar 2.



JURNAL UMUM					
Tanggal	No Bukti	Uraian	Nama Akun	Debet	Kredit
15/07/20	A.001-1507	Uang Pangkal Masuk Ang	Kas	100.000	0
15/07/20	A.001-1507	Uang Pangkal Masuk Ang	Uang Pangkal	0	100.000
15/07/20	S.001-1507	Uang Simpanan Anggota	Kas	625.000	0
15/07/20	S.001-1507	Uang Simpanan Anggota	Simpanan Pokok	0	100.000
15/07/20	S.001-1507	Uang Simpanan Anggota	Simpanan Wajib	0	500.000
15/07/20	S.001-1507	Uang Simpanan Anggota	Simpanan Sukar	0	25.000
14/07/20	P.001-1407	Penarikan Uang Simpana	Kas	0	25.000
14/07/20	P.001-1407	Penarikan Uang Simpana	Simpanan Wajib	25.000	0
15/07/20	A.003-1507	Uang Pangkal Masuk Ang	Kas	10.000	0
15/07/20	A.003-1507	Uang Pangkal Masuk Ang	Uang Pangkal	0	10.000
15/07/20	S.002-1507	Uang Simpanan Anggota	Kas	775.000	0
15/07/20	S.002-1507	Uang Simpanan Anggota	Simpanan Pokok	0	150.000

Gambar 2. Data Plaintext

Setelah data *Plaintext* dipersiapkan dalam format txt, selanjutnya meng-upload data tersebut ke aplikasi yang sudah dibuat.

4.1.2 Form Enkripsi

Form ini berfungsi untuk melakukan proses enkripsi data, di mana data *Plaintext* yang sudah dipersiapkan akan di-upload ke dalam aplikasi. Form enkripsi dapat dilihat pada gambar 3.

Form Enkripsi DES

Masukan Kunci

 Ambil File Data

 Enkripsi File

 Simpan Hasil Enkripsi

Enkripsi Dalam Milidetik

Plainteks Karakter

 Proses Selanjutnya

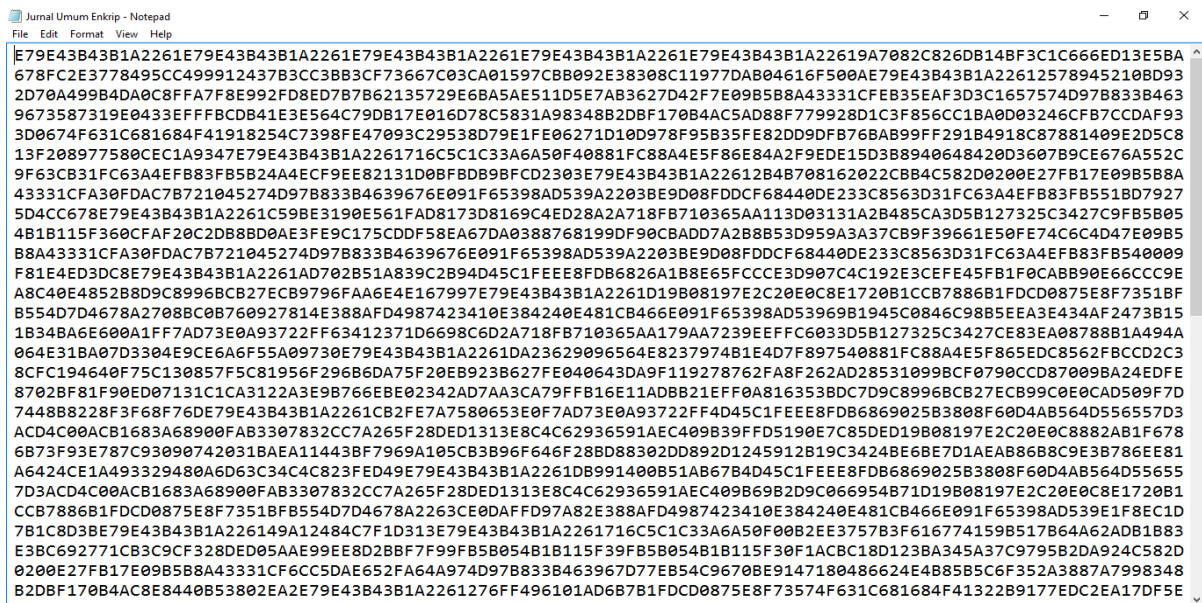
 Batal

Gambar 3. Form Enkripsi

Cara penggunaan form enkripsi adalah pertama memasukkan kunci enkripsi dalam hal ini pengguna bebas memasukkan kunci dalam bentuk apapun, kunci yang dimasukkan akan mempengaruhi hasil enkripsi. Selanjutnya pengguna dapat mengambil data *Plaintext* dengan cara memilih Tombol Ambil File Data lalu memilih Tombol Enkripsi File dan terakhir memilih Tombol Simpan Hasil Enkripsi. Pada gambar 3 terlihat dengan panjang karakter 3.196 dapat melakukan enkripsi hanya dalam 61 milidetik. Selanjutnya akan dilakukan proses *Encoding* steganografi dengan cara memilih Tombol Proses Selanjutnya.

4.1.3 Hasil *Ciphertext*

Ciphertext merupakan data hasil dari enkripsi, metode *Data Encryption Standard* mengubah data *Plaintext* menjadi *Ciphertext* dalam format hexadecimal dengan panjang kunci 54 bit. Adapun hasil file *Ciphertext* dapat dilihat pada gambar 4.



Gambar 4. Hasil Cipertext

Dengan hasil *Ciphertext* dari aplikasi yang sudah dibuat, maka dipastikan bahwa data sudah cukup aman karena orang lain yang tidak berkompeten dalam melihat data tidak akan bisa membaca informasi, karena hasil

Ciphertext menjadi sandi-sandi yang tidak dapat dipahami. Namun akan lebih aman lagi jika data hasil *Ciphertext* ini diamankan ke dalam citra digital menggunakan teknik steganografi metode *End Of File*.

4.1.4 Form *Encoding*/Enkripsi Steganografi

Form ini berfungsi untuk memasukkan gambar asli (cover image) yang akan disisipkan file hasil *Ciphertext* lalu diproses menjadi gambar yang sudah disisipkan pesan (stego image). Adapun form *Encoding* steganografi dapat dilihat pada gambar 5.



Gambar 5. Form *Encoding*/Enkripsi Steganografi

Cara melakukan proses *Encoding* adalah memilih gambar sebagai cover image dengan browse file yang telah disediakan aplikasi, lalu mengambil file *Ciphertext* dan memasukkan password sebagai kunci proses steganografi, lalu selanjutnya memilih Tombol Simpan Gambar dan Proses Enkripsi.

4.2. Pengujian *Decoding*/Dekripsi Steganografi dan Dekripsi Kriptografi

Pada proses pengujian ini akan dilakukan terlebih dahulu proses *Decoding*/enkripsi steganografi menggunakan metode *End Of File*, selanjutnya dekripsi kriptografi menggunakan metode *Data Encryption Standard*. Adapun proses pengujian adalah sebagai berikut :

4.2.1 Form *Decoding*/Dekripsi Steganografi

Form ini berfungsi untuk melakukan proses *Decoding*/dekripsi yaitu memisahkan stego image dan file isi pesan. Adapun form *Decoding* steganografi dapat dilihat pada gambar 6.

Gambar 6. Form Decoding/Dekripsi Steganografi

Cara melakukan proses *Decoding* adalah memilih gambar stego image dengan cara browse file yang telah disediakan aplikasi lalu memasukkan password beserta konfirmasi sesuai dengan password yang dimasukkan pada proses *Encoding*, jika password tidak sesuai maka proses *Decoding* akan gagal setelah itu menyimpan hasil steganografi ke lokasi yang diinginkan.

4.2.2 Form Dekripsi Kriptografi

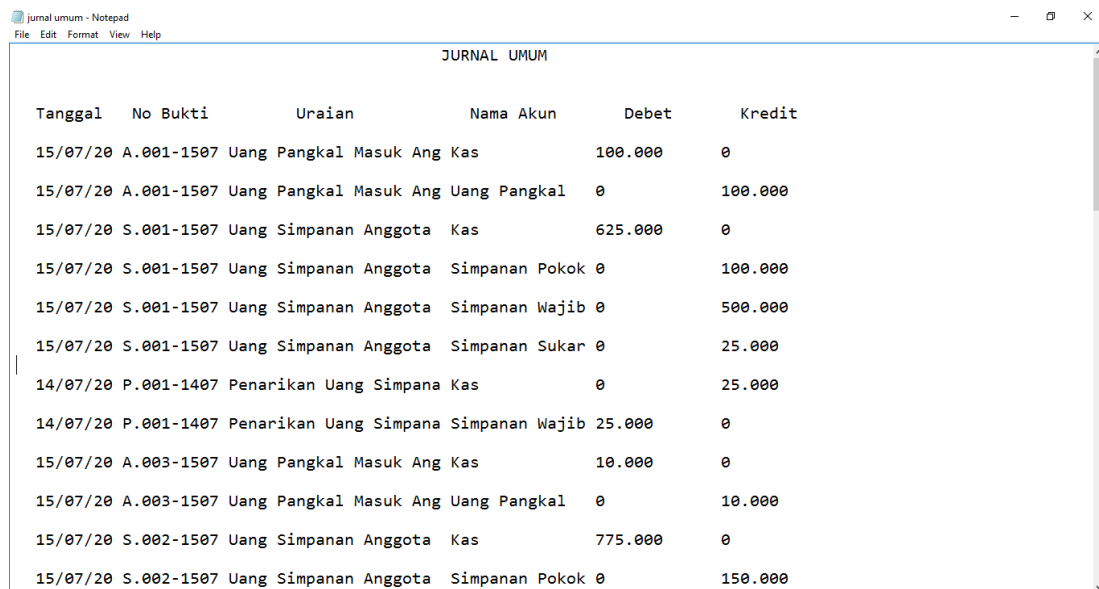
Form ini berfungsi untuk proses dekripsi data yaitu mengembalikan file *Ciphertext* menjadi file *Plaintext* kembali. Adapun form dekripsi kriptografi dapat dilihat pada gambar 7.

Gambar 7. Form Dekripsi Steganografi

Cara melakukan proses dekripsi adalah masukkan kunci yang sesuai pada proses enkripsi kriptografi, jika kunci salah maka proses dekripsi tidak akan dapat dilakukan. Selanjutnya pilih Tombol Ambil Laporan lalu pilih Tombol Dekripsi File dan Tombol Simpan Hasil Dekripsi, jika berhasil maka hasil dekripsi dapat ditampilkan dengan cara memilih Tombol Tampilkan Hasil Dekrip.

4.2.3 Hasil Dekripsi

Jika proses dekripsi berhasil maka hasilnya dapat ditampilkan sesuai dengan data aslinya sebelum dilakukan proses enkripsi. Adapun tampilan hasil dekripsi dapat dilihat pada gambar 8.



Tanggal	No Bukti	Uraian	Nama Akun	Debet	Kredit
15/07/20	A.001-1507	Uang Pangkal Masuk Ang Kas		100.000	0
15/07/20	A.001-1507	Uang Pangkal Masuk Ang Uang Pangkal		0	100.000
15/07/20	S.001-1507	Uang Simpanan Anggota Kas		625.000	0
15/07/20	S.001-1507	Uang Simpanan Anggota Simpanan Pokok		0	100.000
15/07/20	S.001-1507	Uang Simpanan Anggota Simpanan Wajib		0	500.000
15/07/20	S.001-1507	Uang Simpanan Anggota Simpanan Sukar		0	25.000
14/07/20	P.001-1407	Penarikan Uang Simpana Kas		0	25.000
14/07/20	P.001-1407	Penarikan Uang Simpana Simpanan Wajib		25.000	0
15/07/20	A.003-1507	Uang Pangkal Masuk Ang Kas		10.000	0
15/07/20	A.003-1507	Uang Pangkal Masuk Ang Uang Pangkal		0	10.000
15/07/20	S.002-1507	Uang Simpanan Anggota Kas		775.000	0
15/07/20	S.002-1507	Uang Simpanan Anggota Simpanan Pokok		0	150.000

Gambar 8. Hasil Dekripsi

Pada tampilan hasil dekripsi menunjukkan bahwa gambar 8 dan gambar 2 data *Plaintext* menghasilkan data yang sama persis, ini berarti menunjukkan bahwa proses pengujian kriptografi telah berhasil.

4.3 Pengujian Kualitas Citra Steganografi *End Of File*

Pengujian ini berfungsi untuk memastikan bahwa teknik steganografi *End Of File* yang telah disisipi pesan rahasia dalam bentuk *Ciphertext* mempunyai kualitas yang baik. Berikut ini adalah hasil pengujian kualitas citra steganografi *End Of File* berdasarkan aplikasi yang telah dibuat.

4.3.1 Pengujian MSE (*Mean Square Error*) PNSR (*Peak Signal to Noise Ratio*)

Pengujian MSE dan PNSR digunakan untuk mengukur kualitas citra yang dihasilkan. Metode PNSR adalah ukuran perbandingan antara nilai piksel *cover image* dengan nilai piksel pada citra *stego* yang dihasilkan. Tabel 1 dan 2 menunjukkan hasil pengujian MSE dan PNSR.

Tabel 1. Hasil Pengujian MSE

No	Nama Gambar	Ukuran Gambar Asli	Ukuran Stego Image	MSE
1	Tengkorak.jpg	17.1 KB	107 KB	27 db
2	Mobil.jpg	68 KB	1.021 KB	29 db
3	Sepeda.jpg	6.3 KB	2.52 KB	0.0003 db
4	Hitam.jpg	34.4 KB	327 KB	2.6 db
5	Putih.jpg	6.36 KB	2.51 KB	0.000008 db

Tabel 2. Hasil Pengujian PNSR

No	Nama Gambar	Ukuran Gambar Asli	Ukuran Stego Image	PSNR
1	Tekngkorak.jpg	210 KB	371 KB	27db
2	Mobil.jpg	210 KB	371 KB	21 db
3	Sepeda.jpg	210 KB	374 KB	133 db
4	Hitam.jpg	207 KB	322 KB	22 db
5	Putih.jpg	340 KB	575 KB	104 db

Berdasarkan pengujian MSE dan PNSR didapatkan bahwa nilai MSE yang dihasilkan kurang dari 1 dB dan PNSR di atas 50, berarti perubahan kualitas warna antara citra asli dengan *stego image* tidak mengalami perubahan yang signifikan, sehingga keberadaan dari file yang tersembunyi tidak mudah di deteksi oleh indra penglihatan manusia.

4.3.2 Pengujian Recovery

Pengujian *recovery* dilakukan untuk menguji apakah pesan rahasia yang disisipi pada sebuah citra harus dapat dipisahkan kembali dari stego-image-nya. Pengujian dapat dilakukan dengan melihat keutuhan pesan yang diekstraksi dari sejumlah citra uji. Berikut pengujian *recovery* yang dilakukan dapat dilihat pada tabel 3.

Tabel 3. Hasil Pengujian Recovery

No	Nama Gambar	Normal	Diputar	Dipotong	Recovery
1	Tengkorak.jpg	✓			Berhasil
2	Mobil.jpg		✓		Gagal
3	Sepeda.jpg		✓		Gagal
4	Hitam.jpg			✓	Gagal
5	Putih.jpg			✓	Gagal

Dari hasil pengujian *recovery* yang telah dijabarkan pada table diatas dapat dilihat bahwa pesan rahasia yang disisipi pada sebuah citra dapat dipisahkan kembali dari stego-image-nya. Namun saat *cover image* di putar dan di crop maka pesan tidak dapat dibuka.

5. SIMPULAN

1. Berdasarkan hasil pengujian kriptografi menggunakan metode *Data Encryption Standard*, didapatkan bahwa file *Plaintext* dapat diubah menjadi *Ciphertext* dengan kecepatan 61 Milidetik dengan panjang karakter 3.196 dan berhasil dikembalikan lagi menjadi file *Plaintext*
2. Aplikasi steganografi dengan metode *End Of File* yang telah dibuat didapatkan hasil keakuratan yang 100% akurat. Dalam proses *Encoding* dan *Decoding* akan menghasilkan output yang sama dengan inputannya.

3. Hasil pengujian nilai PSNR terhadap image atau file citra digital yang dihasilkan dari aplikasi Steganografi inipun menunjukkan nilai yang cukup baik bergantung pada besar ukuran file citra yang digunakan dan besarnya jumlah karakter yang disisipkan pada file citra tersebut
4. Penggabungan metode Kriptografi Data Encryption Standard dan Steganografi *End Of File* dapat menjadi salah satu alternative dalam pengamanan data dan informasi

6. UCAPAN TERIMA KASIH

Ucapan terimakasih kepada Ibu Dra. Wamiliana, M.A., Ph.D. sebagai promotor yang telah membimbing penulis dalam melakukan proses penelitian.

KEPUSTAKAAN

- [1] Munir, Rinaldi. (2006). Diktat Kuliah IF2153 Matematika Diskrit. Program Studi Teknik Informatika, Institut Teknologi Bandung.
- [2] Sitohang, Ernita., (2013). Perangkat Aplikasi Keamanan Data Text Menggunakan Electronic Codebook Dengan Algoritma DES, Pelita Informatika Budi Darma, ISSN 2301-9425
- [3] Ariyus, Dony., (2006). *Kriptografi Keamanan Data dan Komunikasi*. Graha Ilmu, Yogyakarta.
- [4] Wandany, Heny., dan Budiman, (2012). Implementasi Sistem Keamanan Data dengan Menggunakan Teknik Steganografi *End Of File* (EOF) dan Rabin Public Key Cryptosystem, Jurnal Alkhawarizmi. Vol.1, No 1.
- [5] Primartha, Rifkie., (2011). Penerapan Enkripsi Dan Dekripsi File Menggunakan Algoritma *Data Encryption Standard* (DES). Jurnal Sistem Informasi. Vol. 3 No 2. ISSN : 2085 – 1588
- [6] Darwis, Dedi., dan Kisworo (2017). Teknik Steganografi Untuk Penyembunyian Pesan Teks Menggunakan Algoritma *End Of File*. Jurnal Explore Sistem Informasi dan Telematika. ISSN : 2087 – 2062
- [7] Ariyus, Dony, (2008). *Pengantar Ilmu Kriptografi Teori, Analisis, dan Implementasi*. Andi Offset, Yogyakarta.
- [8] Mufida Khairani, Sajadin Sembiring, (2013). *Analisis dan Implementasi Steganografi Pada Citra GIF Menggunakan Algoritma GifShuffle*, SNASTIKOM, ISBN 978-602-19837-3-7.
- [9] Iswahyudi, C., Setyaningsih, E., (2012). *Pengamanan Kunci Enkripsi Citra Pada Algoritma Super Enkripsi Menggunakan Metode End Of File*. Jurnal Prosiding Nasional Aplikasi Sains & Teknologi (SNAST) Periode III.
- [10] Cheddad, A., Condell, J., Curran, K., Kevitt, P.Mc., (2010). *Digital Image Steganography : Survey and Analysis of Current Methods*. Signal Processing, Elsevier. Northern Ireland, UK.

Bayesian Inference of Poisson Distribution using Conjugate and Non-Informative Priors

Misgiyati, Khoirin Nisa, Warsono

Department of Mathematics University of Lampung
Jl. Prof. Dr. Soemantri Brodjonegoro No. 1 Gedung Meneng Bandar Lampung

ABSTRACT

Poisson distribution is one of the most important and widely used statistical distributions. It is commonly used to describe the frequency probability of specific events when the average probability of a single occurrence within a given time interval is known. In this paper, Bayesian inference of Poisson distribution parameter (μ) is presented. Two Bayesian estimators of μ using two different priors are derived, one by using conjugate prior by applying gamma distribution, and the other using non-informative prior by applying Jeffery prior. The two priors yield the same posterior distributions namely gamma distribution. Comparison of the two Bayesian estimators is conducted through their bias and mean square error evaluation.

Keywords: Bayesian, conjugate prior, Jeffery prior, marginal distribution, asymptotic variance.

1.INTRODUCTION

Parameter estimation is a method to estimate parameter value of population using the statistical values of sample [1]. To estimate the population parameter, a representative sample should be taken. Before making any predictions, the population of the random variable must be known such as the distribution form and its parameters [2]. In the theory of statistical inference, the estimation can be done by the classical method or Bayesian method. The classical method relies entirely on inference process on the sample data, while the Bayesian method not only relying on the sample data obtained from the population but also taking into account an initial distribution [3].

The Bayesian method assumes the parameter(s) as a variable that describes the initial knowledge of the parameter before the observation is performed and expressed in a distribution called the prior distribution. Prior selection is generally made on the basis of whether or not the information about parameter is available. If information about parameter is known then it is called as informative prior, meaning that the prior affects the posterior distribution and is subjective. Whereas if the information about the parameter is unknown then the non-informative prior is used and does not give significant effect to the posterior distribution so that the information obtained is more objective [4]. If the posterior distribution is in the same family as the prior probability distribution, the prior and posterior are then called conjugate distributions, and the prior is called a conjugate prior for the likelihood function.

Studies on Bayesian inference for distribution parameter estimation have been done by many authors. For example Fikhri *et al.* discussed the Bayesian estimator for Poisson distribution and compared it to the maximum likelihood estimator [7]. Noor and Soehardjoepri investigated the Bayesian method approach for the assessment of log-normal distribution parameter estimates using non-informative prior [8]. Hazhiah *et al.* discussed the estimation of two parameters Weibull distribution using the Bayesian method [9]. Prahutama *et al.* discussed the statistical inference from normal distribution using Bayesian method with non-informative prior [10].

In this paper, we study the estimation of the Poisson distribution parameters using Bayesian method with conjugate and non-informative prior. The conjugate prior is a special case of informative prior that leads to the posterior having the same form of the prior distribution. The conjugate prior of Poisson distribution is gamma distribution. For the non-informative prior, the well-known Jeffrey's prior is applied.

2. BAYESIAN INFERENCE FOR POISSON DISTRIBUTIONS

Poisson distribution was introduced by Siméon Denis Poisson in his book which explains the applications of probability theory. The book was published in 1837 with the title *Recherchessur la probabilité des jugements en matiêrecriminelleet en matiêrecivile* (Investigations into the Probability of Verdicts in Criminal and Civil Matters) [5].

Let $X_1, X_2, \dots, X_n$, be random sample from Poisson distributed population with parameter μ . The probability density function for Poisson distribution with parameter μ is given by:

$$f(x; \mu) = \begin{cases} \frac{e^{-\mu} \mu^x}{x!} & ; x = 0, 1, 2, \dots \\ 0 & ; x \text{ lainnya} \end{cases}$$

Now consider the general problem of inferring a Poisson distribution with parameter μ given some data x . From Bayes' theorem, the posterior distribution is equal to the product of the likelihood function $L(\mu) = f(x|\mu)$ and prior $f(\mu)$, divided by the probability of the data $f(x)$:

$$f(\mu|x) = \frac{f(\mu)f(x|\mu)}{\int_0^\infty f(\mu)f(x|\mu)} \quad (1)$$

Let the likelihood function be considered fixed; the likelihood function is usually well-determined from a statement of the data-generating process. It is clear that different choices of the prior distribution $f(\mu)$ may make the integral more or less difficult to calculate, and the product $f(\mu) \times f(x|\mu)$ may take one algebraic form or another. For certain choices of the prior, the posterior has the same algebraic form as the prior (generally with different parameter values). Such a choice is a *conjugate prior*.

The form of the conjugate prior can generally be determined by inspection of the probability density or probability mass function of a distribution. The usual conjugate prior of Poisson distribution is the gamma distribution with parameters (α, β) , namely:

$$f(\mu; \alpha, \beta) = \frac{1}{\beta^\alpha \Gamma(\alpha)} \mu^{\alpha-1} e^{-\frac{\mu}{\beta}} \quad (2)$$

where $\Gamma(\cdot)$ is the gamma function defines as: $\Gamma(\alpha) = \int_0^\infty x^{\alpha-1} e^{-x} dx$.

There has been a desire for a prior distribution that plays a minimal in the posterior distribution. These are sometime referred to a non-informative prior. While it may seem that picking a non-informative prior distribution might be easy, (e.g. just use a uniform distribution), it is not quite that straight forward. One approach of non-informative prior is by using Jeffrey's method. Jeffrey's principle states that any rule for determining the prior density $f(\mu)$ should yield an equivalent result if applied to the transformed parameter

(posterior). Applying this principle, Jeffrey's prior can be expressed as the square root of Fisher information as follow:

$$f(\mu) = \sqrt{I(\mu)}$$

where $I(\mu)$ is the Fisher information for μ , i.e.

$$I(\mu) = -E \left(\frac{\partial^2 \log L(\mu)}{\partial \mu^2} \right).$$

Then for the Poisson parameter μ we calculate the Jeffery's prior as the following. First we calculate the likelihood function

$$\begin{aligned} L(\mu) &= \prod_{i=1}^n f(x_i) \\ &= \prod_{i=1}^n \frac{e^{-\mu} \mu^{x_i}}{x_i!} \\ &= \frac{e^{-n\mu} \mu^{\sum_{i=1}^n x_i}}{\prod_{i=1}^n x_i!} \end{aligned}$$

The log of the likelihood function is given by

$$\begin{aligned} \log L(\mu) &= \log \frac{e^{-n\mu} \mu^{\sum_{i=1}^n x_i}}{\prod_{i=1}^n x_i!} \\ &= \log e^{-n\mu} + \log \mu^{\sum_{i=1}^n x_i} - \log \prod_{i=1}^n x_i! \\ &= -n\mu + \sum_{i=1}^n x_i \log \mu - \log \prod_{i=1}^n x_i! \end{aligned}$$

The first and second derivatives of $\log L(\mu)$ with respect to x are

$$\begin{aligned} \frac{\partial \log L(\mu)}{\partial \mu} &= \frac{\partial}{\partial \mu} (-n\mu) + \frac{\partial}{\partial \mu} (\sum_{i=1}^n x_i \log \mu) - \frac{\partial}{\partial \mu} (\log \prod_{i=1}^n x_i!) \\ &= -n + \frac{1}{\mu} \sum_{i=1}^n x_i - 0 \\ &= -n + \frac{1}{\mu} \sum_{i=1}^n x_i, \end{aligned}$$

and

$$\begin{aligned} \frac{\partial^2 \log L(\mu)}{\partial \mu^2} &= \frac{\partial}{\partial \mu} (-n + \frac{1}{\mu} \sum_{i=1}^n x_i) \\ &= \sum_{i=1}^n x_i \frac{\partial}{\partial \mu} (\mu^{-1}) \\ &= \frac{-\sum_{i=1}^n x_i}{\mu^2}. \end{aligned}$$

Fisher information for μ is given by:

$$\begin{aligned} I(\mu) &= -E \left(\frac{\partial^2 \log L(\mu)}{\partial \mu^2} \right) \\ &= -E \left(\frac{-\sum_{i=1}^n x_i}{\mu^2} \right) \\ &= \frac{1}{\mu^2} E(\sum_{i=1}^n x_i) \\ &= \frac{1}{\mu^2} (n\mu) \end{aligned}$$

$$= \frac{n}{\mu}$$

The Jeffrey's prior then is obtained as

$$\begin{aligned} f(\mu) &= \sqrt{I(\mu)} = \sqrt{\frac{n}{\mu}}, n \text{ constant} \\ &= \frac{1}{\sqrt{\mu}} \end{aligned} \quad (3)$$

In the next section we use the gamma and Jeffrey's priors in Equation (2) and (3) respectively for obtaining the Bayesian estimator of Poisson parameter μ .

3. RESULT AND DISCUSSION

A. Bayesian Inference of Poisson using Conjugate prior

Applying the Bayesian law of probability as expressed in Equation (1), we can calculate the posterior of μ as follow. The joint distribution and the marginal distribution are given by

$$\begin{aligned} f(\mu)f(x|\mu) &= \frac{\mu^{\alpha-1} e^{-\frac{\mu}{\beta}} e^{-n\mu} \mu^{\sum_{i=1}^n x_i}}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!} \\ &= \frac{\mu^{\alpha-1+\sum_{i=1}^n x_i} e^{-\mu(n+\frac{1}{\beta})}}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!}, \end{aligned}$$

and

$$\begin{aligned} \int_0^\infty f(\mu)f(x|\mu) d\mu &= \int_0^\infty \frac{\mu^{\alpha-1+\sum_{i=1}^n x_i} e^{-\mu(n+\frac{1}{\beta})}}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!} d\mu \\ &= \frac{1}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!} \int_0^\infty \mu^{\alpha-1+\sum_{i=1}^n x_i} e^{-\mu(n+\frac{1}{\beta})} d\mu \\ &= \frac{1}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!} \Gamma(\alpha+\sum x_i) \left(n+\frac{1}{\beta}\right)^{-(\alpha+\sum x_i)} \\ &= \frac{\Gamma(\alpha+\sum x_i) \left(n+\frac{1}{\beta}\right)^{-(\alpha+\sum x_i)}}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!}. \end{aligned}$$

Then the posterior distribution can be written as

$$\begin{aligned} f(\mu|x) &= \frac{f(\mu)f(x|\mu)}{\int_0^\infty f(\mu)f(x|\mu) d\mu} \\ &= \frac{\mu^{\alpha-1+\sum_{i=1}^n x_i} e^{-\mu(n+\frac{1}{\beta})}}{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!} \cdot \frac{\beta^\alpha \Gamma(\alpha) \prod_{i=1}^n x_i!}{\Gamma(\alpha+\sum x_i) \left(n+\frac{1}{\beta}\right)^{-(\alpha+\sum x_i)}} \\ &= \frac{\mu^{\alpha-1+\sum_{i=1}^n x_i} e^{-\mu(n+\frac{1}{\beta})}}{\Gamma(\alpha+\sum x_i) \left(n+\frac{1}{\beta}\right)^{-(\alpha+\sum x_i)}} \end{aligned} \quad (4)$$

The posterior distribution expressed in Equation (4) is gamma distribution with parameter $(\alpha + \sum x_i)$ and $\left(n + \frac{1}{\beta}\right)^{-1}$, or $\mu \sim \text{gamma} \left(\alpha + \sum x_i, \left(n + \frac{1}{\beta}\right)^{-1} \right)$. The Bayesian estimate is given by the mean of the

posterior. For the final step of Bayesian Estimation we need to recall the following property of gamma distribution.

Proposition 1

If X is a gamma distributed random variable with parameters (α, β) , the expected value and variance of X are given by $E(X) = \alpha\beta$ and $\text{Var}(X) = \alpha\beta^2$ respectively. (The proof of this property can be seen in e.g. [6]).

Applying Proposition 1 we obtain the Bayesian estimate of μ using gamma prior as follow

$$\hat{\mu}_g = \frac{\alpha + \sum x_i}{n + \frac{1}{\beta}}.$$

B. Bayesian Inference using non-Informative prior

To obtain Jeffrey's posterior we apply the prior in Equation (3) to Equation (1). First we calculate the joint and the marginal distributions as follow

$$\begin{aligned} f(x; \mu) &= f(x|\mu)f(\mu) \\ &= \frac{e^{-n\mu} \mu^{\sum_{i=1}^n x_i} \frac{1}{\sqrt{\mu}}}{\prod_{i=1}^n x_i!} \\ &= \frac{e^{-n\mu} \mu^{\sum_{i=1}^n x_i} \mu^{-1/2}}{\prod_{i=1}^n x_i!} \\ &= \frac{e^{-n\mu} \mu^{-\frac{1}{2} + \sum_{i=1}^n x_i}}{\prod_{i=1}^n x_i!} \end{aligned}$$

and

$$\begin{aligned} \int f(x; \mu) d\mu &= \int_0^\infty \frac{\mu^{-1/2 + \sum_{i=1}^n x_i} e^{-n\mu}}{\prod_{i=1}^n x_i!} d\mu \\ &= \frac{1}{\prod_{i=1}^n x_i!} \int_0^\infty \mu^{-1/2 + \sum_{i=1}^n x_i} e^{-n\mu} d\mu \\ &= \frac{1}{\prod_{i=1}^n x_i!} \Gamma\left(\frac{1}{2} + \sum_{i=1}^n x_i\right) (n^{-1})^{\sum_{i=1}^n x_i + 1/2} \end{aligned}$$

The Jeffrey's posterior is given by

$$\begin{aligned} f(\mu|x) &= \frac{f(x; \mu)}{\int f(x; \mu)} \\ &= \frac{e^{-n\mu} \mu^{-\frac{1}{2} + \sum_{i=1}^n x_i}}{\prod_{i=1}^n x_i!} \frac{\prod_{i=1}^n x_i!}{\Gamma\left(\frac{1}{2} + \sum_{i=1}^n x_i\right) (n^{-1})^{\sum_{i=1}^n x_i + 1/2}} \\ &= \frac{e^{-n\mu} \mu^{-\frac{1}{2} + \sum_{i=1}^n x_i}}{\Gamma\left(\frac{1}{2} + \sum_{i=1}^n x_i\right) (n^{-1})^{\sum_{i=1}^n x_i + 1/2}} \\ &= \frac{e^{-n\mu} \mu^{\sum_{i=1}^n x_i + 1/2}}{\Gamma\left(\frac{1}{2} + \sum_{i=1}^n x_i\right) (n^{-1})^{\sum_{i=1}^n x_i + 1/2}} \end{aligned} \tag{5}$$

The posterior distribution expressed in Equation (4) is gamma distribution with parameter $\left(\frac{1}{2} + \sum x_i\right)$ and $\frac{1}{n}$, or $\mu \sim \text{gamma}\left(\frac{1}{2} + \sum x_i, \frac{1}{n}\right)$. By Proposition 1 we obtain the Bayesian estimate of μ using Jeffrey's prior as follow

$$\hat{\mu}_J = \frac{\frac{1}{2} + \sum x_i}{n}$$

C. Evaluation of the estimators

The resulted Bayesian estimates using gamma and Jeffrey's priors are compared by evaluating their bias and variance. The bias of an estimator is the difference between this expected value of the estimator and the true value of the parameter being estimated. Below we show that the two estimates are bias but asymptotically unbiased.

The expected value of Bayesian estimate using gamma prior is

$$\begin{aligned} E(\hat{\mu}_g) &= E\left(\frac{\alpha + \sum x_i}{n + \frac{1}{\beta}}\right) = \frac{1}{n + \frac{1}{\beta}} E(\alpha + \sum x_i) \\ &= \frac{1}{n + \frac{1}{\beta}} E(\alpha) + E(\sum x_i) \\ &= \frac{1}{n + \frac{1}{\beta}} (\alpha + \sum E(x_i)) \\ &= \frac{1}{n + \frac{1}{\beta}} (\alpha + n\mu) \end{aligned} \quad (6)$$

Since $E(\hat{\mu}_g) \neq \mu$ then $\hat{\mu}_g$ is a biased estimator of μ , now we calculate the bias for $n \rightarrow \infty$ as follow

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{1}{n + \frac{1}{\beta}} (\alpha + n\mu) &= \lim_{n \rightarrow \infty} \frac{(\alpha + n\mu)}{n + \frac{1}{\beta}} \\ &= \lim_{n \rightarrow \infty} \frac{(\alpha/n + n\mu/n)}{n/n + \frac{1}{n\beta}} \\ &= \lim_{n \rightarrow \infty} \frac{(\alpha/n + \mu)}{1 + \frac{1}{n\beta}} \\ &= \mu \end{aligned}$$

Since $\lim_{n \rightarrow \infty} E(\hat{\mu}_g) = \mu$ then $\hat{\mu}_g$ an asymptotically unbiased estimator of μ .

Now consider the Bayesian estimate of μ using Jeffrey's prior $\hat{\mu}_J = \frac{\frac{1}{2} + \sum x_i}{n}$, the expected value is given by

$$\begin{aligned} E(\hat{\mu}_J) &= E\left(\frac{\frac{1}{2} + \sum x_i}{n}\right) = \frac{1}{n} E\left(\frac{1}{2} + \sum x_i\right) \\ &= \frac{1}{n} E\left(\frac{1}{2}\right) + E(\sum x_i) \\ &= \frac{1}{n} \left(\frac{1}{2} + \sum E(x_i)\right) \end{aligned}$$

$$\begin{aligned}
 &= \frac{1}{n} \left(\frac{1}{2} + n\mu \right) \\
 &= \frac{1}{2n} + \mu
 \end{aligned} \tag{7}$$

Since $E(\hat{\mu}_J) \neq \mu$ then $\hat{\mu}_g$ is an unbiased estimator of μ . For $n \rightarrow \infty$ as we obtain

$$\begin{aligned}
 \lim_{n \rightarrow \infty} \frac{1}{n} + \mu &= \lim_{n \rightarrow \infty} \frac{(1+2n\mu)}{2n} \\
 &= \lim_{n \rightarrow \infty} \frac{(1/n+2\mu)}{2} \\
 &= \mu,
 \end{aligned}$$

then $\hat{\mu}_J$ is an asymptotically unbiased estimator of μ .

The consistency analysis of the two estimators is described through their means square errors (MSEs) as follow

$$MSE = Var(\hat{\mu}) + (Bias(\hat{\mu}))^2$$

Where $Bias(\hat{\mu}) = E(\hat{\mu}) - \mu$ and $Var(\hat{\mu})$ is the variance of $\hat{\mu}$.

For $\hat{\mu}_g$,

$$Bias(\hat{\mu}_g) = \frac{\alpha + n\mu}{n + \frac{1}{\beta}} - \mu = \frac{\alpha + n\mu - (n + \frac{1}{\beta})\mu}{n + \frac{1}{\beta}} = \frac{\alpha - \frac{\mu}{\beta}}{n + \frac{1}{\beta}},$$

and

$$\begin{aligned}
 Var(\hat{\mu}_g) &= Var\left(\frac{\alpha + \sum x_i}{n + \frac{1}{\beta}}\right) = \frac{1}{(n + \frac{1}{\beta})^2} Var(\alpha + \sum x_i) \\
 &= \frac{1}{(n + \frac{1}{\beta})^2} [Var(\alpha) + Var(\sum x_i)] \\
 &= \frac{1}{(n + \frac{1}{\beta})^2} [0 + Var(\sum x_i)] \\
 &= \frac{1}{(n + \frac{1}{\beta})^2} [\sum Var(x_i)] \\
 &= \frac{1}{(n + \frac{1}{\beta})^2} n\alpha\beta^2 \\
 &= \frac{n\alpha\beta^2}{(n + \frac{1}{\beta})^2}.
 \end{aligned}$$

Then

$$MSE(\hat{\mu}_g) = \frac{n\alpha\beta^2}{(n + \frac{1}{\beta})^2} + \left[\frac{\alpha - \frac{\mu}{\beta}}{n + \frac{1}{\beta}} \right]^2 = \frac{n\alpha\beta^2 + (\alpha - \frac{\mu}{\beta})^2}{(n + \frac{1}{\beta})^2}$$

For $\hat{\mu}_J$,

$$Bias(\hat{\mu}_J) = E(\hat{\mu}_J) - \mu = \frac{1}{2n} + \mu - \mu = \frac{1}{2n},$$

and

$$Var(\hat{\mu}_J) = Var\left(\frac{\frac{1}{2} + \sum x_i}{n}\right)$$

$$\begin{aligned}
 &= \frac{1}{n^2} \text{Var} (1/2 + \sum x_i) \\
 &= \frac{1}{n^2} [\text{Var} (1/2) + \text{Var} (\sum x_i)] \\
 &= \frac{1}{n^2} [0 + \sum \text{Var} (x_i)] \\
 &= \frac{1}{n^2} [\sum \text{Var} (x_i)] \\
 &= \frac{1}{n^2} n\alpha\beta^2 \\
 &= \frac{n\alpha\beta^2}{n^2}.
 \end{aligned}$$

Then,

$$\begin{aligned}
 \text{MSE}(\hat{\mu}_J) &= \text{Var} (\hat{\mu}_J) + (\text{Bias} (\hat{\mu}_J))^2 \\
 &= \frac{n\alpha\beta^2}{n^2} + \left[\frac{1}{2n}\right]^2 \\
 &= \frac{4n\alpha\beta^2+1}{4n^2}
 \end{aligned}$$

For comparing the MSE of $\hat{\mu}_g$ and $\hat{\mu}_J$, we conducted a simulation study and we present the result in the following section.

4. SIMULATION STUDY

The objective of the data simulation is to compare empirically the MSE of Bayesian estimates using gamma and Jeffrey's priors. The simulation was done using R software. We generated data from Poisson distributions with parameter $\mu = 3, 5, 10$ and sample sizes $n=20, 50, 100, 300, 500, 1000, 5000$ dan 10000 . For the Bayesian estimate using gamma prior we fix gamma parameters: $\alpha =1$ and $\beta = 5$. The MSE of the estimates of the Bayesian estimators is provided in the following tables.

Table 1. MSE of $\hat{\mu}$ obtained using the conjugate (gamma) and non-informative (Jeffrey's) priors for $\mu = 3$

n	MSE	
	Conjugate Prior (gamma)	Non-informative Prior (Jeffrey's)
20	7.839709	8.235576
50	1.250809	1.276125
100	0.270299	0.272759
300	0.032169	0.032266
500	0.011520	0.011539
1000	0.002854	0.002861
5000	0.000128	0.000127
10000	0.000033	0.000033

Tabel 2. MSE of $\hat{\mu}$ obtained using the conjugate (gamma) and non-informative (Jeffrey's) priors $\mu = 5$

n	MSE	
	Conjugate Prior (gamma)	Non-informative Prior (Jeffrey's)

20	12.469016	13.091211
50	2.135610	2.177006
100	0.468095	0.472870
300	0.056265	0.056477
500	0.019575	0.019611
1000	0.005426	0.005412
5000	0.000197	0.000197
10000	0.000050	0.000050

Table 3. MSE of $\hat{\mu}$ obtained using the conjugate (gamma) and non-informative (Jeffrey's) priors $\mu = 10$

n	MSE	
	Conjugate Prior (gamma)	Non-informative Prior (Jeffrey's)
20	26.030662	27.275180
50	3.902299	3.970509
100	1.009183	1.017402
300	0.116299	0.116489
500	0.037545	0.037528
1000	0.010257	0.010206
5000	0.000384	0.000383
10000	0.000098	0.000097

The MSEs presented in Table 1-3 of the two Bayesian estimators are very similar. When the sample size (n) increases the MSEs become smaller and closer to zero. We notice that for small sample size ($n=20$) the gamma prior produces much better estimates than the Jeffrey's prior (the differences of the MSEs are around one unit), but they become more similar for larger n (>20). From the simulation result we indicate that the two estimators are consistent estimators for Poisson parameter μ .

5. CONCLUSIONS

From the analytical and empirical results presented in this paper we obtain the following conclusions :

1. The Bayesian estimate of Poisson parameter μ using gamma prior (notated by $\hat{\mu}_g$) and Jeffrey's prior (notated by $\hat{\mu}_J$) are :

$$\hat{\mu}_g = \frac{\alpha + \sum x_i}{n + \frac{1}{\beta}} \quad \text{and} \quad \hat{\mu}_J = \frac{\frac{1}{2} + \sum x_i}{n}$$

2. Theoretically, $\hat{\mu}_g$ and $\hat{\mu}_J$ are biased estimators of μ , but we proved that they are asymptotically unbiased.
3. Based on the simulation result, the larger the sample size the smaller the MSE value of the two estimators (and closer to zero) then they both are consistent.

REFERENCES

- [1] Supranto, J. (2001). *Statistika Teori dan Aplikasi*, edisi ke-6. Jakarta: Erlangga.
- [2] Roussas, G. (2003). *An Introduction to Probability and Statistical Inference*. California: Academic Press.

- [3] Brain, L. J & Engelhardt, M. (1992). *Introduction to Probability and Mathematical Statistics*. USA: Duxbury Thomson Learning.
- [4] Bolstad, W.M. (2007). *Introduction to Bayesian Statistics Second Edition*. America: A John Wiley & Sons Inc.
- [5] Ross, S. (2010). *A First Course in Probability, Eighth Edition*. New Jersey: Pearson.
- [6] Herrhyanto, N. & Gantini, T. (2011). *Pengantar Statistika Matematis*. Bandung :Yrama Widya.
- [7] Fikhri, M., Yanuar, F. & Asdi, Y. (2014). Pendugaan Parameter dari Distribusi Poisson dengan Menggunakan Maximum Likelihood Estimation (MLE) dan Metode Bayes. *Jurnal Matematika Unand*. Universitas Andalas.
- [8] Noor, D. & Soehardjoepri, E. (2016). Pendekatan Metode Bayesian untuk Kajian Estimasi Parameter Distribusi Log-Normal untuk Non-Informatif Prior. *Jurnal Sains dan Seni ITS*. Vol. 5 N0.2 (2016) 2337-3520 (2301-928X Print). Institut Teknologi Sepuluh November.
- [9] Hazhiah, T., Sugito, I & Rahmawati, R. (2012). Estimasi Parameter Distribusi Weibull Dua Parameter Menggunakan Metode Bayes. *Jurnal Gaussian*. Vol. 1 halaman 103-112. Universitas Diponegoro.
- [10] Prahutama, A., Sugito & Rusgiyono, A. (2012), Inferensi Statistik dari Distribusi Normal dengan Metode Bayes untuk Non-informatif Prior. *Media Statistika*. Vol. 5 No. 2. Desember 2012 : 95-104. Institut Teknologi Sepuluh November.

ANALISIS KLASIFIKASI MENGGUNAKAN METODE REGRESI LOGISTIK ORDINAL DAN KLASIFIKASI NAÏVE BAYES PADA DATA ALUMNI UNILA TAHUN 2016

Shintia Faramudhita, Rudi Ruswandi, Subian Saidi

Jurusan Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lampung

Email : shintiafaramudhita@gmail.com, ruswandir@ymail.com, subian.saidi@gmail.com

ABSTRAK

Klasifikasi adalah proses pencarian sekumpulan pola, model atau fungsi yang menggambarkan dan membedakan objek data untuk dikelompokkan kedalam kelas tertentu dari sejumlah kelas yang tersedia. Pada penelitian ini akan dilakukan klasifikasi tingkat kelancaran alumni unila tahun 2016 yang dibagi menjadi tiga kelas bertingkat, yaitu kelas Tidak Lancar (TL), kelas Kurang Lancar (KL) dan kelas Lancar (L). Metode analisis yang digunakan adalah Regresi Logistik Ordinal yang merupakan metode pengklasifikasi dengan teknik statistika dan Naïve Bayes yang merupakan metode pengklasifikasi dengan menggabungkan teknik statistika dan data mining. Penelitian ini bertujuan untuk mengetahui metode mana yang mempunyai tingkat akurasi yang lebih baik dalam mengklasifikasi tingkat kelancaran alumni unila tahun 2016 dalam mencari pekerjaan. Dengan melakukan sepuluh kali pengulangan klasifikasi dari masing-masing metode, didapat bahwa Regresi Logistik Ordinal mempunyai rata-rata tingkat error yang lebih kecil yaitu sebesar dibandingkan dengan Naïve Bayes.

Kata Kunci : Klasifikasi, Regresi Logistik Ordinal, Klasifikasi Naïve Bayes

I. PENDAHULUAN

Menurut [1] , klasifikasi adalah proses pencarian sekumpulan model, pola atau fungsi yang menggambarkan dan membedakan objek data untuk dikelompokkan ke dalam kelas tertentu dari sejumlah kelas yang tersedia. Variabel yang digunakan dalam klasifikasi terdiri dari variabel prediktor yang merupakan faktor-faktor yang mempengaruhi atau dapat menggambarkan variabel respon. Dalam hal ini variabel respon berupa variabel kategorik baik yang mempunyai pengurutan dalam penomoran (ordinal) maupun tidak (nominal).

Beberapa metode statistika yang dapat digunakan dalam klasifikasi melalui analisis data kategorik, yaitu Regresi Logistik Biner, Regresi Logistik Multinomial, Regresi Logistik Ordinal dan Model Log Linier. Pada penelitian ini data yang digunakan adalah klasifikasi tingkat kelancaran alumni Universitas Lampung (Unila) tahun 2016 dalam mendapatkan pekerjaan. Klasifikasi alumni dibagi menjadi kelompok tidak lancar, kelompok kurang lancar dan kelompok lancar berdasarkan lamanya waktu yang dibutuhkan untuk mendapatkan pekerjaan . Variabel respon yang digunakan berbentuk kategorik bertingkat atau ordinal. Sehingga metode yang dapat digunakan adalah Regresi Logistik Ordinal.

Metode pengklasifikasi lain yang dapat digunakan adalah Klasifikasi Naïve Bayes. Naïve Bayes merupakan metode penggolongan berdasarkan probabilitas sederhana dan dirancang untuk dipergunakan dengan asumsi bahwa antar satu kelas dengan kelas yang lain tidak saling tergantung (independen) meskipun asumsi ini tidak terpenuhi, dalam prakteknya masih berfungsi dengan baik. Naïve Bayes merupakan gabungan dari teknik statistika yang didasari oleh Teorema Bayes dan data mining *machine larning* yang mampu memberikan informasi berguna tentang teknik klasifikasi untuk menentukan alumni Unila akan bergabung dengan kelompok mana dalam mencari pekerjaan.

Karena kedua metode klasifikasi diatas memiliki perbedaan, dimana Regresi Logistik Ordinal merupakan metode klasifikasi yang menggunakan teknik statistika, sedangkan Naïve Bayes adalah metode klasifikasi yang menggabungkan teknik statistika dengan data mining. Sehingga menarik untuk mengkaji perbandingan penggunaan kedua metode tersebut pada data alumni Unila lulusan 2016. Untuk mengetahui metode mana yang mempunyai tingkat akurasi yang lebih baik. Akan tetapi, penelitian ini hanya dikhususkan untuk pengklasifikasian data Alumni Unila tahun 2016.

Berdasarkan latar belakang diatas maka penelitian ini akan menganalisis tingkat ketepatan metode Regresi Logistik Ordinal dan Naive Bayes dalam mengklasifikasi tingkat kelancaran alumni Unila tahun 2016 dalam mencari pekerjaan berdasarkan rata-rata tingkat error dari masing-masing metode.

II. TINJAUAN PUSTAKA

a. Regresi Logistik Ordinal

Regresi Logistik Ordinal merupakan salah satu metode statistika untuk menganalisis variabel respon yang mempunyai skala data ordinal yang memiliki tiga kategori atau lebih. Sedangkan variabel prediktor yang digunakan berupa data katgorik dan atau kuantitatif. Pada regresi logistik ordinal model berupa Model Logit Kumulatif (*Cumulative Logit Models*). Model logit kumulatif ini diperoleh dengan membandingkan peluang kumulatif yaitu peluang kurang dari atau sama dengan kategori respon ke-j pada i variabel prediktor yang dinyatakan dalam vektor x_i adalah $P(Y \leq j | X_i)$, dengan peluang lebih dari kategori respon ke-j pada i variabel prediktor vektor X_i $P(Y > j | X_i)$. Peluang kumulatif $P(Y \leq j | X_i)$ didefinisikan sebagai berikut :

$$\begin{aligned} P(Y \leq j | X_i) = \pi(x_i) &= \frac{e^{g_j(x)}}{1 + e^{g_j(x)}} \\ &= \frac{e^{\beta_{j0} - \sum_{i=1}^I \beta_i x_i}}{1 + e^{\beta_{j0} - \sum_{i=1}^I \beta_i x_i}} \end{aligned} \quad (1)$$

Model logit kumulatif didefinisikan dengan :

$$\begin{aligned} g_j(x) &= \ln \left[\frac{\pi_j(x)}{1 - \pi_j(x)} \right] \\ &= \beta_{j0} - \beta_1 x_1 - \beta_2 x_2 - \dots - \beta_i x_i \end{aligned} \quad (2)$$

Dengan, j adalah jumlah kategori variabel respon $j = 1, 2, \dots, J$ dan i adalah jumlah variabel prediktor. Jika terdapat kategori respon dimana $j=1,2,3$, maka Odds Ratio atau nilai peluang untuk tiap kategori respon dapat dihitung dengan menggunakan persamaan dibawah ini :

$$\pi_1(x) = \frac{e^{g_1(x)}}{1 + e^{g_1(x)}} \quad (3)$$

$$\pi_2(x) = \frac{e^{g_2(x)} - e^{g_1(x)}}{(1 + e^{g_2(x)})(1 + e^{g_1(x)})} \quad (4)$$

.

$$\pi_k(x) = 1 - \pi_{k-1}(x) \quad (5)$$

Untuk menentukan kelas dari suatu objek dapat dilihat dari nilai peluang kategori yang paling besar. Metode yang digunakan untuk menaksir parameter-parameter pada Regresi Logistik Ordinal adalah *Maximum Likelihood Estimation* (MLE) dengan fungsi likelihood :

$$l(\beta) = \prod_{i=1}^n [\pi_1(x_i)^{y_{1i}} \pi_2(x_i)^{y_{2i}} \pi_3(x_i)^{y_{3i}}] \quad (6)$$

Dengan,

$$y_{ji} = \begin{cases} 1 & \text{untuk } y=j \\ 0 & \text{untuk } y \neq j \end{cases}$$

Dari persamaan di atas didapatkan fungsi *ln- Likelihood* sebagai berikut :

$$\begin{aligned} \ln l(\beta) &= \sum_{i=1}^n \{y_{1i} \ln[\pi_1(x_i)] + y_{2i} \ln[\pi_2(x_i)] + y_{3i} \ln[\pi_3(x_i)]\} \\ &= \sum_{i=1}^n \left\{ y_{1i} \ln \left[\frac{e^{g_1(x)}}{1 + e^{g_1(x)}} \right] + y_{2i} \ln \left[\frac{e^{g_2(x)} - e^{g_2(x)}}{(1 + e^{g_2(x)})(1 + e^{g_1(x)})} \right] + y_{3i} \ln \left[1 - \frac{e^{g_2(x)}}{1 + e^{g_2(x)}} \right] \right\} \end{aligned} \quad (7)$$

Untuk mendapatkan nilai pendugaan parameter dari fungsi *ln-likelihood* pada regresi logistik ordinal dilakukan metode iterasi *Newton Raphson* [2]. Persamaan *Newton Raphson*.

$$\beta^{(t+1)} = \beta^t - (H^t)^{-1} u^t \quad (8)$$

Model yang telah terbentuk dari parameter-parameter yang di estimasi menggunakan *maximum likelihood estimation* selanjutnya diuji signifikansinya baik secara keseluruhan maupun secara parsial.

1. Uji Ratio Likelihood

Uji Ratio Likelihood dilakukan untuk menguji kesesuaian model dengan variabel-variabel prediktor secara keseluruhan [3].

Adapun hipotesis yang digunakan dalam uji ratio likelihood

$$H_0 : \beta_1 = \beta_2 = \dots = \beta_p = 0$$

H_1 : paling sedikit salah satu dari $\beta_r \neq 0$ dengan $r = 1, 2, \dots, p$

Dengan statistik uji

$$G = -2 \ln \left[\frac{\left(\frac{n_1}{n}\right)^{n_1} \left(\frac{n_0}{n}\right)^{n_0}}{\prod_{i=1}^n \hat{\pi}_i^{y_i} (1 - \hat{\pi}_i)^{(1-y_i)}} \right] \quad (9)$$

$$G = 2 \{ \sum_{i=1}^n [y_i \ln(\hat{\pi}_i) + (1-y_i) \ln(1-\hat{\pi}_i)] - [n_1 \ln(n_1) + n_0 \ln(n_0) - n \ln(n)] \} \quad (10)$$

dengan :

n_1 = banyaknya observasi berkategori 1

n_0 = banyaknya observasi berkategori 0

n = banyaknya observasi ($n_1 + n_0$)

Statistik uji G mengikuti distribusi *chi-square*. Sehingga untuk mengambil keputusan dilakukan perbandingan dengan χ^2 tabel. Kriteria penolakan tolak H_0 jika $\chi^2_{hit} > \chi^2_{(db,a)}$

2. Uji Wald

Uji Wald dilakukan untuk mengetahui variabel-variabel prediktor mempengaruhi variabel respon secara individu dengan kata lain apakah suatu variabel prediktor layak dimasukkan kedalam model [3].

Hipotesis yang digunakan dalam uji wald :

$$H_0 : \beta_i = 0;$$

$$H_1 : \beta_i \neq 0, \text{ dengan } i = 1, 2, \dots, n$$

Dengan statistik uji

$$W_r = \left[\frac{\hat{\beta}_i}{SE(\hat{\beta}_i)} \right]^2 \quad (11)$$

Statistik uji W_r mengikuti sebaran normal baku. Sehingga untuk memperoleh keputusan dilakukan perbandingan dengan distribusi normal baku (Z). Dengan kriteria pengambilan keputusan tolak H_0

Jika $W_i > Z_{\alpha/2}$.

b. Naive Bayes

Menurut [1], metode Klasifikasi Bayesian merupakan pengklasifikasian secara statistika yang memprediksi peluang anggota kelas tertentu berdasarkan database yang memenuhi syarat keanggotaan kelas tersebut. Dengan formula Naive Bayes :

$$P(Y_j|X) = \frac{P(Y) \prod_{i=1}^q P(X_i|Y_j)}{P(X)} \quad (12)$$

Dengan,

$P(Y_j|X_i)$ = Probabilitas posterior kelas Y pada semua fitur vektor

$P(Y)$ = Probabilitas prior kelas Y (*prior probability*)

$P(X_i|Y_j)$ = probabilitas posterior semua fitur vektor X pada kelas Y

$\prod_{i=1}^n P(X_i|Y_j)$ = $P(X_1|Y_j) \times P(X_2|Y_j) \times \dots \times P(X_n|Y_j)$

$P(X)$ = Probabilitas X

Pada Klasifikasi Naive Bayes hasil klasifikasi ditentukan dengan melihat nilai $P(Y_j|X)$ paling besar dari setiap variabel Y.

c. Error Rate

Untuk mengetahui tingkat akurasi hasil klasifikasi, maka dilakukan uji ketepatan hasil klasifikasi dengan menggunakan APER (*Apparent Error Rate*) atau yang disebut juga dengan laju error. Untuk menghitung

nilai APER, langkah yang harus dilakukan membentuk tabel perbandingan hasil klasifikasi berdasarkan observasi dengan hasil klasifikasi berdasarkan prediksi suatu metode yang disebut dengan matrik konfusi hasil klasifikasi [1].

Tabel 1. Matrik Konfusi Hasil Klasifikasi

<i>Fij</i>		Kelas Asli (Hasil Observasi)		
		Kelas = 1	Kelas = 2	Kelas = 3
Kelas Prediksi (Hasil Prediksi)	Kelas = 1	<i>F11</i>	<i>F12</i>	<i>F13</i>
	Kelas = 2	<i>F21</i>	<i>F22</i>	<i>F23</i>
	Kelas = 3	<i>F31</i>	<i>F32</i>	<i>F33</i>

Dengan,

F11=Jumlah alumni kelas 1 pada kelas asli dan kelas 1 pada kelas prediksi

F12 = Jumlah alumni kelas 1 pada kelas asli dan kelas 2 pada kelas prediksi

F13 = Jumlah alumni kelas 1 pada kelas asli dan kelas 3 pada kelas prediksi

.
.

.

Fij = Jumlah alumni kelas i pada kelas asli dan kelas j pada kelas prediksi

Selanjutnya dapat dilakukan perhitungan nilai APER dengan formulasi sebagai berikut

$$\text{APER} = \frac{f_{12}+f_{13}+f_{21}+f_{23}+f_{31}+f_{32}}{f_{11}+f_{12}+f_{13}+f_{21}+f_{22}+f_{23}+f_{31}+f_{32}+f_{33}} \times 100\% \quad (13)$$

$$\text{Tingkat Akurasi} = 1 - \text{APER} \quad (14)$$

Suatu metode dikatakan memiliki tingkat akurasi yang baik jika mempunyai nilai APER yang kecil dan tingkat akurasi yang tinggi.

d. Uji Dua Rata-Rata

Uji hipotesis dua rata-rata digunakan mengetahui ada atau tidaknya perbedaan (kesamaan) antara dua data .

Hipotesis :

$$H_0: \mu_1 = \mu_2$$

$$H_a: \mu_1 \neq \mu_2$$

Statistik Uji :

$$t = \frac{\bar{x}_1 - \bar{x}_2 - d_0}{\sqrt{S_p^2(1/n + 1/m)}} \quad (15)$$

Dengan

$$S_p^2 = \frac{(n-1)S_1^2 + (m-1)S_2^2}{n+m-2}$$

Keputusan :

Tolak H_0 jika $t < -t_{(\alpha, df)}$ atau $t > t_{(\alpha, df)}$

III. METODOLOGI PENELITIAN

a. Sumber Data dan Variabel Penelitian

Data yang digunakan dalam penelitian ini adalah data *Tracer Study* Universitas Lampung (Unila) 2016 yang diperoleh dari UPT. Pengembangan Karir dan Kewirausahaan Universitas Lampung. . Data yang digunakan adalah alumni yang mempunyai riwayat mencari pekerjaan, sedangkan yang tidak pernah mencari pekerjaan sama sekali baik melanjutkan studi maupun berwirausaha tidak dimasukan dalam objek penelitian. Software yang digunakan yaitu Rstudio (3.3.3).

Tabel.2 Variabel Prediktor

Variabel Prediktor (X)	Kategori	
Program Studi (X_1)	$X_{1=(1)}$ = Eksak	
	$X_{1=(2)}$ = Non Eksak	
Jenjang Pendidikan (X_2)	$X_{2=(1)}$ = D3	
	$X_{2=(2)}$ = S1	
IPK (X_3)	$X_{3=(1)}$ = ≤ 2.75	
	$X_{3=(2)}$ = $>2.75- \leq 3.5$	
	$X_{3=(3)}$ = >3.5	
Lama Study (X_4)	$X_{4=(1)}$ = ≤ 3.5 tahun	≤ 3 tahun
	$X_{4=(2)}$ = $> 3.5-4.5$ tahun	$>3 - \leq 4$ tahun
	$X_{4=(3)}$ = >4.5 tahun	>4 tahun
Cara Mencari Pekerjaan (X_5)	$X_{5=(1)}$ = Melalui Media	
	$X_{5=(2)}$ = Secara Mandiri	
	$X_{5=(3)}$ = Relasi	
Penguasaan Pengetahuan diluar Program studi (X_6)	$X_{6=(1)}$ = Sangat Tidak Menguasai	
	$X_{6=(2)}$ = Tidak Menguasai	
	$X_{6=(3)}$ = Cukup Menguasai	
	$X_{6=(4)}$ = Menguasai	
	$X_{6=(5)}$ = Sangat Menguasai	

Langkah-langkah yang dilakukan pada penelitian ini adalah sebagai berikut:

1. Menentukan klasifikasi awal

Pada studi kasus klasifikasi tingkat kelancaran alumni dalam mendapatkan pekerjaan dapat dibuat klasifikasi awal dengan kriteria waktu yang dibutuhkan untuk mendapatkan pekerjaan pertama :

Tabel.3 Variabel Respon

Variabel Respon (Y)	Kriteria
Tidak Lancar (TL) = 1	Mendapatkan pekerjaan dalam waktu > 12 bulan setelah wisuda
Kurang Lancar (KL) = 2	Mendapatkan pekerjaan dalam waktu > 6 bulan -12 bulan setelah wisuda
Lancar (L) = 3	Mendapatkan pekerjaan dalam ≤ 6 setelah wisuda

2. Membagi data menjadi dua, yaitu data training 75% dan data testing 25%. Dengan data yang sama dilakukan sepuluh kali pengulangan pembagian data training dan data testing, proporsi jumlah data testing data training yang sama tetapi kombinasi data yang berbeda.
3. Membuat model regresi logistik ordinal
 - a. Membentuk model awal regresi logistik ordinal dengan menggunakan data training.
 - b. Menguji signifikansi parameter secara keseluruhan dengan menggunakan Uji Ratio Likelihood
 - c. Menguji parameter secara parsial dengan Uji Wald, pengujian ini dilakukan untuk mengetahui variabel-variabel prediktor mempengaruhi variabel respon secara individu.
 - d. Pembentukan model akhir regresi logistik ordinal
 - e. Menentukan klasifikasi data testing menggunakan model akhir. Dalam regresi logistik ordinal kelas hasil prediksi adalah kelas yang memiliki nilai peluang paling tinggi.
 - f. Menghitung nilai APER dan akurasi dari model yang terbentuk.
4. Naive Bayes

Adapun tahapan klasifikasi Naive Bayes sebagai berikut :

 - a. Menghitung probabilitas awal (*prior probability*) peluang kelas Y $P(Y)$
 - b. Menghitung nilai probabilitas X_i bersyarat Y_j $P(X_i|Y_j)$
 - c. Menghitung *posterior probability*
 - d. Menentukan hasil klasifikasi, pada Klasifikasi Naive Bayes hasil klasifikasi ditentukan dengan melihat nilai $P(Y_j|X)$ paling besar dari setiap variabel Y dan data yang digunakan adalah data testing.
 - e. Menghitung nilai APER dan akurasi dari model yang terbentuk.
5. Membandingkan nilai error rate dari masing-masing metode, untuk mengetahui metode yang lebih baik dapat dilihat dari rata-rata akurasi yang lebih tinggi.
6. Menguji kesamaan dua rata-rata kedua metode tersebut.

IV. HASIL DAN PEMBAHASAN

a. Deskripsi Data

Setelah dilakukan klasifikasi awal dengan indikator pada tabel.3 maka diperoleh bahwa jumlah alumni yang mendapatkan pekerjaan > 12 bulan sebanyak 169 alumni, Mendapatkan pekerjaan dalam waktu > 6 bulan -12 bulan setelah wisuda sebanyak 228 alumni dan Mendapatkan pekerjaan dalam ≤ 6 setelah wisudah sebanyak 709 alumni. Dengan tabulasi variabel X terhadap variabel Y pada Tabel.4

Tabel.4. Tabulasi Variabel Prediktor Terhadap Variabel Respon

Variabel Prediktor (X)		Variabel Respon (Y)		
		TL	KL	L
Program Studi (X_1)	$X_{1=1}$ Eksak	58	73	245
	$X_{1=2}$ Non Eksak	111	155	464
Jenjang Pendidikan (X_2)	$X_{2=1}$ D3	25	18	51
	$X_{2=2}$ S1	144	210	658
IPK (X_3)	$X_{3=1} \leq 2.75$	5	8	9
	$X_{3=2} > 2.75 \text{ s.d } \leq 3.5$	107	138	432
	$X_{3=3} > 3.5$	57	82	268
Lama Studi (X_4)	$X_{4=1}$ Lama	34	49	143
	$X_{4=2}$ Tepat Waktu	95	132	428
	$X_{4=3}$ Cepat	40	47	138
Cara Mencari Pekerjaan (X_5)	$X_{5=1}$ Media	80	101	309
	$X_{5=2}$ Mandiri	49	70	238
	$X_{5=3}$ Relasi	40	57	162
Penguasaan Pengetahuan Diluar Bidang Studi (X_6)	$X_{6=1}$ STM	2	1	14
	$X_{6=2}$ Tidak Menguasai	13	25	62
	$X_{6=3}$ Cukup Menguasai	60	85	293
	$X_{6=4}$ Menguasai	74	89	268
	$X_{6=5}$ Sangat Menguasai	20	28	72

b. Regresi Logistik Ordinal

Salah satu contoh analisis klasifikasi tingkat kelancaran alumni Unila tahun 2016 dalam mencari pekerjaan dengan menggunakan regresi logistik ordinal.

Tabel.5 Output R

Parameter	Estimate	Std. Value	Z Value
1 2	-0.7544	0.6828	-1.105
2 3	0.4244	0.6816	0.623

Berdasarkan hasil output R 'log Lik.' -735.7309 (df=8) dapat dilihat parameter dari masing-masing variabel dengan menggunakan MLE. Selanjutnya akan dilakukan pengujian parameter secara serentak dengan ratio likelihood test menggunakan persamaan (10).

Tabel.6 Uji Ratio Likelihood Test

G	Chi Square ($\chi^2_{(0.05, 8)}$)	Df
19.8779183012	15.507	8

Berdasarkan uji ratio likelihood dapat dilihat bahwa $G (19.8779183012) > \chi^2_{(0.05, 8)} = 15.507$ maka keputusannya tolak H_0 artinya paling tidak terdapat satu variabel prediktor yang mempengaruhi model. Kemudian dilakukan pengujian parameter secara individu sebagai berikut :

Table.7 Uji Wald

Variabel	Value ($\hat{\beta}_i$)	Std. Error $\widehat{SE}(\hat{\beta}_i)$	$[\frac{\hat{\beta}_i}{\widehat{SE}(\hat{\beta}_i)}]$	$W = [\frac{\hat{\beta}_i}{\widehat{SE}(\hat{\beta}_i)}]^2$	Keputusan
X_1	-0.14543	0.15661	-0.929	0.863041	Tidak Tolak H_0
X_2	0.62080	0.25246	2.459	6.046681	Tolak H_0
X_3	0.29365	0.14715	1.996	3.865153	Tolak H_0
X_4	-0.14704	0.11714	-1.255	1.5750	Tidak Tolak H_0
X_5	0.06980	0.09036	0.772	0.595987	Tidak Tolak H_0
X_6	-0.12881	0.08437	-1.527	2.331729	Tolak H_0

Dengan menggunakan taraf signifikansi $\alpha = 0.05$ dari tabel Z, maka di peroleh nilai $Z_{\alpha/2} = 1.96$ dengan kriteria pengambilan keputusan tolak H_0 Jika $W_i > Z_{\alpha/2}$. Berdasarkan tabel diatas maka terdapat tiga variabel yang signifikan terhadap model yaitu X_2, X_3, X_6 .

Karena secara parsial hanya tiga parameter variabel yang signifikan, maka dibutuhkan pembentukan model baru dengan menggunakan variabel yang signifikan tersebut, sehingga model logit yang terbentuk sebagai berikut :

$$g_1(x) = \left(\frac{\pi_1(x)}{1-\pi_1(x)} \right) = -0.5897 - 0.50272(x_2) - 0.30652(x_3) + 0.14366(x_6)$$

$$g_2(x) = \left(\frac{\pi_2(x)}{1-\pi_2(x)} \right) = 0.5855 - 0.50272(x_2) - 0.30652(x_3) + 0.14366(x_6)$$

Untuk melakukan penerapan dari model diatas maka dihitung peluang dari masing-masing kelas dengan persamaan (3), (4) dan (5). Untuk ketepatan klasifikasi salah satu data testing menggunakan Regresi logistik Ordinal dapat menggunakan matrik konfusi sebagai berikut :

Tabel.8 Matrik Konfusi Regresi Logistik Ordinal Pertama

Kelas Prediksi	Kelas Observasi		
	Y ₁ = Tidak Lancar	Y ₂ = Kurang Lancar	Y ₃ = Lancar
Y ₁ = Tidak Lancar	0	0	0
Y ₂ = Kurang Lancar	0	0	0
Y ₃ = Lancar	45	54	177

Diperoleh tingkat error dan tingkat akurasi dari Regresi Logistik Ordinal :

$$APPER = \frac{45+54}{276} \times 100\% = 0.3586\% = 0.3586$$

$$Akurasi = 1 - APPEP = 1 - 0.35507 = 0.6413$$

Selanjutnya dilakukan pengulangan klasifikasi sebanyak sepuluh kali dengan kombinasi data yang berbeda pada testing dan training.

c. Klasifikasi Naive Bayes

Dalam klasifikasi Naive Bayes yang pertama dilakukan menghitung P(Y), P(X_i|Y_j) dan P(X) dengan data training yang sama seperti Regresi Logistik Ordinal, kemudian menghitung P(Y_j|X) dengan data testing, sehingga diperoleh hasil klasifikasi naive Bayes sebagai berikut :

Tabel.9 Matrik Konfusi Naive Bayes Klasifikasi Pertama

Kelas Prediksi	Kelas Observasi		
	Y ₁ = Tidak Lancar	Y ₂ = Kurang Lancar	Y ₃ = Lancar
Y ₁ = Tidak Lancar	4	3	9
Y ₂ = Kurang Lancar	0	0	0
Y ₃ = Lancar	41	51	168

Tingkat error dan Tingkat akurasi dari Klasifikasi Naive Bayes pertama :

$$APPER = \frac{41+51+3+9}{276} \times 100\% = 37.68\% = 0.3768$$

$$Akurasi = 1 - APPEP = 1 - 0.3768 = 0.6232$$

Selanjutnya dilakukan pengulangan klasifikasi sebanyak sepuluh kali dengan kombinasi data yang berbeda pada testing dan training.

d. Error Rate

Setelah dilakukan sepuluh kali pengulangan analisis pada Regresi Logistik Ordinal dan Naive Bayes (Lampiran .1) menghasilkan data error rate sebagai berikut :

Tabel.10 Perbandingan Tingkat Error Kedua Metode

Analisis Ke-	Regresi Logistik Ordinal	Naive Bayes
1	0.3586	0.3768
2	0.2935	0.2971
3	0.4094	0.434782
4	0.3732	0.39492
5	0.3514	0.3768
6	0.4239	0.45289
7	0.3768	0.39855
8	0.3297	0.35144
9	0.3696	0.37318
10	0.3225	0.3369
Rata-Rata	0.36086	0.37890664

Untuk mengetahui apakah terdapat perbedaan dalam penggunaan kedua metode dalam mengklasifikasi tingkat kelancaran alumni Unila tahun 2016 dalam mencari pekerjaan maka dilakukan Uji Kesamaan dua rata-rata pada tingkat error diatas.

Hipotesis :

$$H_0: \mu_1 = \mu_2$$

$$H_a: \mu_1 \neq \mu_2$$

Dengan menggunakan persamaan (17) diperoleh t hitung sebesar -0.20087 dengan derajat bebas 18, berdasarkan tabel $-2.306 < t_{\alpha/2, n+m-2} < 2.306$. Maka dapat diambil keputusan tidak tolak H_0 artinya tidak ada perbedaan dari penggunaan kedua metode dalam mengklasifikasi tingkat kelancaran alumni dalam mencari pekerjaan.

V. KESIMPULAN

Berdasarkan hasil dan pembahasan diatas maka dapat ditarik kesimpulan sebagai berikut

- Klasifikasi tingkat kelancaran alumni Unila tahun 2016 dalam mencari pekerjaan dengan menggunakan Naïve Bayes mempunyai rata-rata tingkat error lebih besar yaitu 0.37890664 dibandingkan dengan metode Regresi Logistik Ordinal sebesar 0.36086.
- Hasil dari uji dua rata-rata menunjukan tidak ada perbedaan secara signifikan penggunaan kedua metode dalam mengklasifikasi tingkat kelancaran alumni unila tahun 2016 dalam mencari pekerjaan,

KEPUSTAKAAN

- [1] Han,J., Kamber, M., Jian,P., 2012. *Data Mining Concepts and Techniques, Third Edition*. California : Morgan Kaufman.
- [2] Agresti, Alan. 2002. *An Introduction Categorical Data Analysis, Second Edition*. New Jersey : John Wiley and Sons Inc.
- [3] Hosmer, D.W., Lemeshow, S., 2000. *Applied Logistic Regression, Second Edition*. Canada : John Wiley and Sons Inc.

PEMODELAN MARKOV SWITCHING AUTOREGRESSIVE (MSAR) PADA DATA TIME SERIES

Aulianda Prasyanti¹⁾, Mustofa Usman²⁾, & Dorrah Aziz³⁾
Mahasiswa Jurusan Matematika FMIPA Universitas Lampung¹
Dosen Jurusan Matematika FMIPA Universitas Lampung^{2,3}
aulianda25.p@gmail.com¹

ABSTRAK

Model Markov Switching Autoregressive (MSAR) adalah suatu model yang menganalisis perubahan kondisi fluktuasi pada data time series. Data time series yang digunakan pada penelitian ini adalah kurs dollar AS terhadap rupiah pada tanggal 15 Mei 2016 sampai 20 Februari 2017. Data kurs mempunyai pergerakan perubahan kondisi fluktuasi yakni apresiasi dan depresiasi. Kondisi apresiasi dan depresiasi dianggap suatu variabel yang tidak teramati yang disebut dengan state. Pada penelitian ini estimasi parameter dilakukan dengan metode Maximum Likelihood Estimation (MLE). Namun, pada model MSAR terdapat variabel state dan nilai peluang (p_{ij}) yang tidak dapat diketahui nilainya dengan metode MLE sehingga dilakukan proses filtering dan smoothing terlebih dahulu untuk mencari nilai peluang tersebut. Model terbaik yang diperoleh adalah MS(2)AR(2) dengan peluang perpindahan state disajikan dalam matriks transisi. Dalam model MSAR dihitung rata-rata durasi lama state apresiasi bertahan sebesar 15 hari dan rata-rata durasi state depresiasi bertahan sebesar 22 hari.

Kata kunci : MSAR, fluktuasi, filtering, smoothing, peluang

1. PENDAHULUAN

Dampak ekonomis terjadinya perdagangan internasional adalah perbedaan mata uang yang digunakan oleh negara-negara di dunia dan menimbulkan suatu perbedaan nilai tukar mata uang (kurs). Kurs merupakan alat yang digunakan sebagai salah satu tolok ukur kondisi ekonomi suatu negara. Salah satu mata uang yang mempengaruhi ekonomi di dunia dan mudah untuk diperdagangkan adalah dollar AS. Data kurs dollar AS terhadap rupiah disajikan dalam bentuk *time series*. Analisis *time series* yang terkenal yakni ARIMA dikembangkan oleh George EP Box dan Gwilym M. Jenkins. Pada model ARIMA, asumsi stasioneritas dan homoskedastisitas harus dipenuhi. Namun kenyataannya, asumsi homoskedastisitas pada beberapa data *time series* tidak terpenuhi. Sehingga untuk mengatasi hal ini diperlukan analisis yang sesuai yakni model ARCH diperkenalkan oleh Robert F. Engle dan perluasan model ARCH yakni model GARCH yang dikenalkan oleh T. Bollerslev. Akan tetapi, model ARIMA, ARCH, dan GARCH adalah model yang tidak memperpertimbangkan kondisi fluktuasi (naik dan turun nilai suatu data). Kondisi fluktuasi pada data kurs adalah peningkatan nilai kurs (apresiasi) dan penurunan nilai kurs depresiasi). Perubahan kondisi fluktuasi sering diabaikan akibatnya diperlukan metode yang sesuai untuk menganalisis data *time series* dengan mempertimbangkan fluktuasi yang terjadi.

[1] mengenalkan model Markov Switching Autoregressive sebagai model *time series* yang dapat menjelaskan perubahan fluktuasi yang terjadi pada data. Fluktuasi pada data merupakan suatu variabel yang tidak teramati yang disebut *state*. Selain itu, model Markov Switching Autoregressive (MSAR) dapat menghitung peluang transisi dan menghitung rata-rata lama durasi untuk masing-masing *state*. Berdasarkan uraian tersebut, pada penelitian ini akan dibahas model Markov Switching Autoregressive (MSAR) pada data kurs dollar AS terhadap

rupiah. Pada penelitian ini terdapat dua *state*, karena perubahan kondisi fluktuasi kurs ada 2 yakni apresiasi dan depresiasi, dan akan dicari model terbaik bagi data kurs dollar AS terhadap rupiah, kemudian ditentukan peluang berpindah atau bertahan suatu *state*, serta menentukan rata-rata lama durasi bertahan untuk masing-masing *state*.

2. LANDASAN TEORI

2.1 Analisis Runtun waktu (*Time Series*)

Time series merupakan rangkaian pengamatan pada suatu variabel yang cara pengambilannya beruntun berdasarkan dengan interval waktu yang tetap. Rangkaian data pengamatan *time series* dapat disimbolkan oleh variabel Y_t dengan t adalah indeks waktu urutan pengamatan.[2]

2.2 Autoregressive (AR)

Model autoregressive adalah representasi random proses yang bergantung secara linear pada nilai sebelumnya. Berikut ini adalah model Autoregressive orde :

$$Y_t = \phi_1 y_{t-1} + \phi_2 y_{t-2} + \dots + \phi_p y_{t-p} + \varepsilon_t \quad (1)$$

dengan

$\phi_1, \phi_2, \dots, \phi_p$: parameter *autoregressive*

ε_t : residual yang *white noise*

Selain itu, ε_t tidak diamati dan harus diperkirakan dari data, akan tetapi berdasarkan model yang di asumsikan untuk Y_t . [2]

2.3 Stasioneritas

Stasioner adalah asumsi pada data yang menyatakan bahwa tidak ada perubahan yang drastis pada data *time series*. Stasioneritas dibedakan menjadi dua yakni stasioner dalam rata-rata dan ragam. Secara visual untuk melihat stasioneritas dibantu dengan menggunakan plot *time series*, yaitu dengan melihat fluktuasi data dari waktu ke waktu. [3]

2.4 Autocorrelation Function dan Parsial Autocorrelation Function

Dalam analisis *time series*, salah satu alat yang digunakan untuk mengecek stasioneritas data dengan menggunakan fungsi autokorelasi dan fungsi autokorelasi parsial. ACF dan PACF disajikan dalam bentuk plot yang sering disebut korelogram.[2]

2.5 Augmented Dickey-Fuller (ADF)

Kestasioneran juga dapat diuji dengan menggunakan uji *Augmented Dickey-Fuller* (ADF). Misalkan kita punya persamaan regresi :

$$\Delta y_t = \phi y_{t-1} + \sum_{j=1}^{p-1} \alpha_j^* \Delta y_{t-j} + u_t \quad (2)$$

Dimana $\phi = -\alpha(1)$ dan $\alpha_j^* = -(\alpha_{j+1} + \dots + \alpha_p)$

Uji ADF dilakukan dengan menghitung nilai τ statistik dengan rumus:

$$\tau = \frac{\hat{\rho}_k}{SE(\hat{\rho}_k)}$$

Pada model ini hipotesis yang diuji adalah :

$H_0 : \phi = 0$ (data *time series* tidak stationer)

$H_1 : \phi < 0$ (data *time series* stationer)

dengan kriteria tolak H_0 jika $|\tau| > |\tau_{\alpha,df}|$. [4]

2.6 State

State adalah suatu kondisi yang merupakan peubah acak x_t , dan jika suatu peubah acak berada pada *state* tertentu maka dapat berpindah ke *state* lain. [1]

Suatu *state* adalah apresiasi atau depresiasi dilihat pada dari μ_{s_t} dengan syarat $\mu_1 < \mu_0$. [1]

2.7 Markov Switching

Switching (perubahan) dapat terjadi pada rata-rata dan varian. Model markov *switching* dapat dituliskan :

$$y_t = \mu_{s_t} + e_t \quad (3)$$

dengan $e_t \sim N(0, \sigma_{s_t}^2)$ sedangkan s_t adalah *state* yang dipengaruhi oleh waktu t . [5]

2.8 Model Markov Switching Autoregressive (MSAR)

Model MSAR dapat dituliskan [6]:

$$(y_t - \mu_{s_t}) = \sum_{i=1}^p \Phi_i (y_{t-i} - \mu_{s_{t-i}}) + e_t \quad (4)$$

atau dapat ditulis sebagai berikut :

$$(y_t - \mu_{s_t}) = \Phi_1 (y_{t-1} - \mu_{s_{t-1}}) + \dots + \Phi_p (y_{t-p} - \mu_{s_{t-p}}) + e_t \quad (5)$$

dengan $e_t \sim iid N(0, \sigma_{s_t}^2)$.

Keterangan:

$y_t, y_{t-1}, \dots, y_{t-p}$: data pengamatan

$\Phi_1, \Phi_2, \dots, \Phi_p$: koefisien *Autoregressive*

$\mu_{s_t}, \mu_{s_{t-1}}, \dots, \mu_{s_{t-p}}$: rata-rata dipengaruhi perubahan *state* dan waktu

$\sigma_{s_t}^2$: varian dipengaruhi perubahan *state*

e_t : residual pada saat t

2.8.1 Peluang Perpindahan *State*

Pada Model MSAR, *state* tidak teramati dan menghitung nilai peluang dengan proses *filtering* dan *smoothing*. Peluang perpindahan *state* dibentuk dalam matriks transisi karena rantai markov pada matriks transisi menyatakan nilai sekarang dipengaruhi nilai masa lalu dengan jumlah entri pada baris matriks transisi bernilai 1. [6]

2.8.2 Rata-rata Jangka Waktu pada Masing-masing *State*

Model MSAR juga dapat menghitung rata-rata lama durasi dari masing-masing *state*. Elemen diagonal dari matriks peluang transisi mengandung informasi penting mengenai durasi rata-rata yang diharapkan dari suatu *state* akan bertahan. Durasi rata-rata *state* dihitung dengan persamaan $E(D) = \frac{1}{1-P_{jj}}$. [6]

2.9 Matriks Transisi (*Transition Matrix*)

Evolusi acak suatu rantai Markov $(Z_n)_{n \in N}$ ditentukan oleh data yang menyatakan bahwa nilai sekarang dipengaruhi oleh nilai masa lalu

$$P_{i,j} = P(Z_1 = j | Z_0 = i), \quad i, j \in S \quad (6)$$

yang bertepatan dengan peluang $P(Z_{n+1} = j | Z_n = i)$ yang independen

$n \in N$. Data ini bisa dibuat matriks yang disebut dengan matriks transisi yakni

$$[P_{i,j}]_{i,j \in S} = P(Z_1 = j | Z_0 = i) \quad (7)$$

dapat juga ditulis untuk keadaan yang memiliki 2 kondisi (*state*)

$$[P_{i,j}]_{i,j \in S} = \begin{bmatrix} P_{0,0} & P_{0,1} \\ P_{1,0} & P_{1,1} \end{bmatrix}$$

Hubungan baris pada matriks transisi memiliki kondisi

$$\sum_{j \in S} P_{i,j} = 1 \quad (8)$$

untuk setiap $i \in S$. [7]

2.10 Estimasi Parameter Model Markov Switching Autoregressive

Estimasi parameter dilakukan untuk menduga nilai dari masing-masing parameter pada model. Estimasi parameter model markov switching autoregressive (MSAR) menggunakan metode Pendugaan kemungkinan Maksimum (*Maximum Likelihood Estimation*). Langkah yang dilakukan adalah dengan menentukan fungsi densitas kemudian dibentuk menjadi fungsi log likelihood.

Model MSAR memiliki fungsi densitas : [8]

$$f(y_t | s_t, s_{t-1}, \dots, s_{t-p}, \Omega_{t-1}; \theta) = \frac{1}{\sigma_{s_t} \sqrt{2\pi}} \exp \left[-\frac{((y_t - \mu_{s_t}) - \Phi_1(y_{t-1} - \mu_{s_{t-1}}) - \Phi_2(y_{t-2} - \mu_{s_{t-2}}) - \dots - \Phi_p(y_{t-p} - \mu_{s_{t-p}}))^2}{2\sigma_{s_t}^2} \right]$$

dengan keterangan :

$\Omega_{t-1} = (y_{t-1}, y_{t-2}, \dots, y_{t-p})$: data pengamatan pada masa lalu

$\theta = (\mu_{s_t}, \sigma_{s_t}^2, \Phi_p, p_{ij})$: parameter model MSAR

Menghitung fungsi densitas dari y_t yang diberikan informasi masalalu Ω_{t-1} dan membutuhkan nilai $s_t, s_{t-1}, \dots, s_{t-p}$ yang merupakan variabel tidak teramati (*state*), untuk menyelesaikan masalah ini, hal yang harus dilakukan adalah mempertimbangkan fungsi densitas bersama dari y_t dan $s_t, s_{t-1}, \dots, s_{t-p}$: [6]

1. Memperoleh fungsi densitas bersama dari y_t dan $s_t, s_{t-1}, \dots, s_{t-p}$ bersyarat informasi masalalu.

$$f(y_t, s_t, s_{t-1}, \dots, s_{t-p} | \Omega_{t-1}; \theta) = f(y_t | s_t, s_{t-1}, \dots, s_{t-p}, \Omega_{t-1}; \theta) P[s_t, s_{t-1}, \dots, s_{t-p}, \Omega_{t-1}; \theta] \quad (9)$$

2. Untuk mendapatkan $f(y_t | \Omega_{t-1}; \theta)$, maka menggabungkan densitas bersama dari $s_t, s_{t-1}, \dots, s_{t-p}$ dengan menjumlahkan semua kemungkinan densitas bersama dari $s_t, s_{t-1}, \dots, s_{t-p}$:

$$\begin{aligned} f(y_t | \Omega_{t-1}; \theta) &= \sum_{s_t=0}^1 \sum_{s_{t-1}=0}^1 \dots \sum_{s_{t-p}=0}^1 f(y_t, s_t, s_{t-1}, \dots, s_{t-p} | \Omega_{t-1}; \theta) \\ &= \sum_{s_t=0}^1 \sum_{s_{t-1}=0}^1 \dots \sum_{s_{t-p}=0}^1 f(y_t | s_t, s_{t-1}, \dots, s_{t-p}, \Omega_{t-1}; \theta) P[s_t, s_{t-1}, \dots, s_{t-p} | \Omega_{t-1}] \end{aligned}$$

Nilai peluang $P[s_t, s_{t-1}, \dots, s_{t-p} | \Omega_{t-1}]$ dihitung dengan menggunakan proses *filtering* dan *smoothing*.

2.10.1 Filtering

Filtering adalah proses yang digunakan untuk mendapatkan nilai peluang suatu *state* pada saat t . Nilai *filtered state probability* merupakan hasil dari proses *filtering*. [8]

proses *filtering* dimulai dengan :

$$P[s_t = s_t, s_{t-1} = s_{t-1}, \dots, s_{t-p+1} = s_{t-p+1} | \Omega_t]$$

dengan

$$P[S_0 = 1 | \Omega_0] = \frac{1 - p_{22}}{2 - p_{22} - p_{11}}$$

$$P[S_0 = 2 | \Omega_0] = \frac{1 - p_{11}}{2 - p_{22} - p_{11}}$$

selanjutnya, menghitung :

$$\begin{aligned} & \mathcal{P}[S_t = s_t = j, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} = i | \Omega_t] \\ &= P[S_t = s_t = j, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} = i | \Omega_t, y_t] \\ &= \frac{f(S_t = s_t = j, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} = i, y_t | \Omega_{t-1})}{f(y_t | \Omega_{t-1})} \\ &= \frac{f(y_t | S_t = s_t, S_{t-1} = s_{t-1}, \dots, S_{t-p} = s_{t-p}, \Omega_{t-1}) \times P[S_t = s_t, S_{t-1} = s_{t-1}, \dots, S_{t-p} = s_{t-p} | \Omega_{t-1}]}{\sum_{s_t=0}^1 \sum_{s_{t-1}=0}^1 \dots \sum_{s_{t-p}=0}^1 f(y_t, S_t = s_t, S_{t-1} = s_{t-1}, \dots, S_{t-p} = s_{t-p} | \Omega_{t-1})} \end{aligned}$$

dan hasil dari proses *filtering* :

$$\begin{aligned} & P[S_t = s_t = j, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} = i | \Omega_t] \\ &= \sum_{s_{t-p}=0}^1 P[S_t = s_t, S_{t-1} = s_{t-1}, \dots, S_{t-p} = s_{t-p} | \Omega_t] \end{aligned} \quad (10)$$

2.10.2 Smoothing

Proses *smoothing* merupakan lanjutan dari proses *filtering*. Peluang hasil dari proses *filtering* adalah $P[S_t = s_t = j, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} = i | \Omega_t]$ dimulai dari $t=T-1, T-2, \dots, 1$ dengan $T > t$ maka diperoleh hasil *smoothed state probability* adalah

$$P[S_t = s_t, S_{t-1} = s_{t-1}, \dots, S_{t-p+1} = s_{t-p+1} | \Omega_T]$$

kemudian didapat :

$$f(y_t | \Omega_T; \theta) = \sum_{s_t=0}^1 \sum_{s_{t-1}=0}^1 \dots \sum_{s_{t-p}=0}^1 f(y_t | s_t, s_{t-1}, \dots, s_{t-p}, \Omega_{t-1}; \theta) P[s_t, s_{t-1}, \dots, s_{t-p} | \Omega_T] \quad (11)$$

Selanjutnya didapat fungsi likelihood [6]

$$L(\theta) = \prod_{t=0}^T f(y_t | \Omega_T; \theta) \quad (12)$$

dan fungsi log likelihood

$$\ln L(\theta) = \sum_{t=0}^T \ln f(y_t | \Omega_T; \theta) \quad (13)$$

dengan EM algoritma didapatkan penaksiran yang mendekati nilai maksimum sebagai estimator. Fungsi log likelihood didiferensialkan terhadap masing-masing parameter dan disamadengankan nol [9].

Hasil estimasi parameter adalah sebagai berikut:

$$\hat{\mu}_{s_t} = \frac{\sum_{t=1}^T y_t \times p(s_t = j | y_t; \theta)}{\sum_{t=1}^T p(s_t = j | y_t; \theta)} \quad (14)$$

$$\hat{\sigma}_{s_t}^2 = \frac{\sum_{t=1}^T (y_t - \hat{\mu}_{s_t})(y_t - \hat{\mu}_{s_t})' P(s_t = j | y_t; \theta)}{\sum_{t=1}^T P(s_t = j | y_t; \theta)} \quad (15)$$

$$\hat{p}_{ij} = \frac{\sum_{t=1}^T P(s_t = j, s_{t-1}, \dots, s_{t-p} = i | \Omega_T; \theta)}{\sum_{t=0}^T P(s_{t-1} = i | \Omega_T; \theta)} \quad (16)$$

$$\hat{\Phi}_p = \frac{\sum_{t=1}^T \{ \sum_{s_t=0}^1 (y_t - \hat{\mu}_{s_t})(y_t - \hat{\mu}_{s_t})' \times P(s_t = j | \Omega_T; \theta) \}}{\sum_{t=1}^T \{ \sum_{s_t=0}^1 P(s_t = j | \Omega_T; \theta) \}} \quad (17)$$

2.11 AIC (Akaike Info Criterion)

Kriteria informasi digunakan untuk pemilihan model terbaik yang dipilih berdasarkan *Akaike Info Criterion* (AIC) karena kriteria ini konsisten untuk menduga parameter model. Tujuan AIC adalah menentukan prediksi terbaik.

dirumuskan dengan :

$$AIC = -2 \left(\frac{l}{T} \right) + 2 \left(\frac{k}{T} \right)$$

dengan l adalah fungsi log-likelihood, k adalah jumlah parameter yang diestimasi,

T adalah jumlah observasi. [10]

2.12 Uji Jarque-Berra

Uji normalitas residual adalah uji yang digunakan untuk mengetahui kenormalan residual pada suatu model univariat. Tujuan dilakukannya uji ini adalah untuk mengetahui apakah residual pada model tersebut berdistribusi normal atau tidak. Uji normalitas dilakukan dengan menggunakan *Jarque-Bera (JB) Test of Normality*. Uji ini menggunakan ukuran skewness dan kurtosis. Perhitungan JB adalah sebagai berikut:

$$JB = \frac{n}{6} \left(S^2 + \frac{(K - 3)^2}{4} \right)$$

Dengan :

n = jumlah sampel

$$S = \text{Skewness} = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^3}{\left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 \right)^{\frac{3}{2}}}$$

$$K = \text{Kurtosis} = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^4}{\left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 \right)^2}$$

Dimana Jarque-Bera (JB) Test of Normality berdistribusi chi-square (χ^2) dengan derajat kebebasan 2. [11]

3. METODOLOGI PENELITIAN

Pada penelitian ini data yang digunakan adalah data *time series* kurs beli dollar AS terhadap rupiah dalam frekuensi harian dari 15 Mei 2016 sampai dengan 20 Februari 2017 yang didapat dari website resmi Bank Indonesia yakni www.bi.go.id sebanyak 282 data dan menggunakan studi literatur secara sistematis yang diperoleh dari buku-buku untuk mendapatkan informasi. Analisis data dilakukan dengan menggunakan software Eviews 9 dan Python 3.6.

Langkah-langkah yang dilakukan pada penelitian ini adalah sebagai berikut :

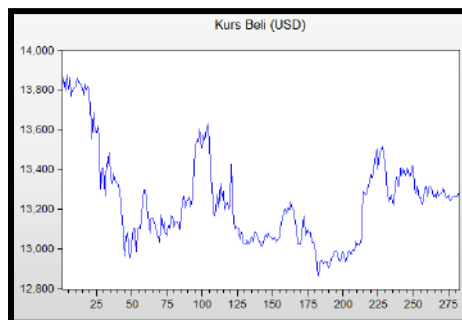
1. Menguji stasioneritas data kurs dollar AS terhadap rupiah.
2. Penentuan *state* dengan mempertimbangkan nilai μ_{s_t} . Pada penelitian ini ada dua *state* yakni *state* apresiasi dan *state* depresiasi sehingga μ_{s_t} ada dua yakni μ_0 dan μ_1 .

3. Estimasi parameter model MSAR dengan bantuan *software* Python 3.6 untuk memperoleh orde yang sesuai.
4. Pengujian diagnostik model.
5. Pengujian normalitas residual digunakan uji Jarque-Bera.
6. Memilih model terbaik yang didapat dari nilai AIC yang minimum.

4. HASIL DAN PEMBAHASAN

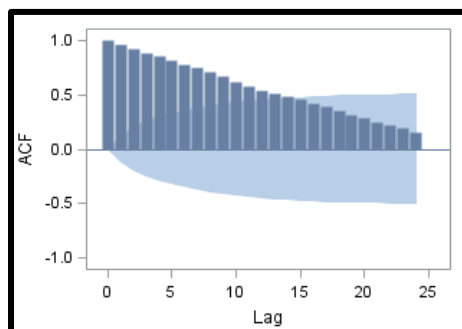
4.1 Uji Stasioneritas

Berikut ini adalah plot *time series* yang dihasilkan dari output E-views 9 :



Gambar 1. Plot data kurs dollar AS terhadap rupiah

Secara visual dapat dilihat bahwa plot time series data kurs dollar AS terhadap rupiah tidak stasioner terhadap rata-rata dan ragam. Hal ini karena fluktuasi rata-rata dan ragam dari data kurs dollar AS terhadap rupiah tidak konstan. Untuk itu, dapat dilihat kolerogram seperti berikut :



Gambar 2. Plot Autocorrelation Function(ACF)

Terlihat bahwa plot ACF data kurs dollar AS terhadap rupiah cenderung menurun lambat sehingga dapat dikatakan tidak stasioner terhadap rata-rata dan ragam.

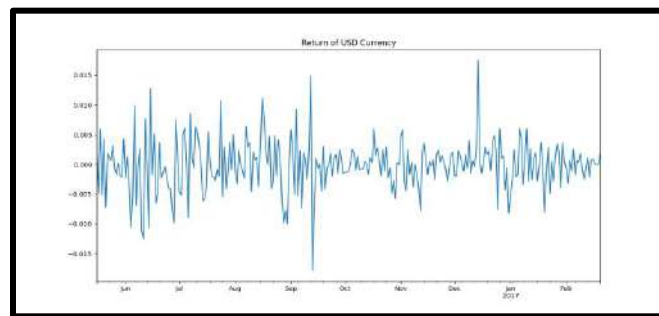
Untuk lebih memastikan bahwa data tidak stasioner, digunakan uji akar unit yakni *Augmented Dickey-Fuller* (ADF).

Tabel 1. Uji akar unit Augmented Dickey-Fuller (ADF) dari kurs dollar AS terhadap rupiah

	t-statistic	P-value
Augmented Dickey-Fuller Test statistic	-2.863511	0.0510

Pada Tabel 1, dapat diketahui bahwa p -value dari uji ADF adalah sebesar 0.0510, karena p -value $> \alpha$ hal ini mengakibatkan data kurs dollar AS terhadap rupiah tidak stasioner terhadap rata-rata dan ragam.

Sehingga diperlukan data *return* pada data kurs dollar AS terhadap rupiah dengan cara differensiasi dan transformasi supaya data kurs dollar AS terhadap rupiah stasioner terhadap rata-rata dan ragam. Setelah didapat data *return* hasil transformasi dan differensiasi maka langkah selanjutnya adalah menguji stasioneritas dari data *return*.



Gambar 3. Plot data return kurs dollar terhadap rupiah

Berdasarkan Gambar 3, plot *return* dapat diketahui bahwa data *return* stasioner terhadap ragam maupun rata-rata. Secara visual hal ini dapat ditunjukkan oleh fluktuasi rata-rata data *return* kurs dollar berada pada nilai yang konstan dan ragamnya berada pada nilai yang konstan.

Untuk memastikan bahwa data *return* kurs dollar AS terhadap rupiah stasioner maka dilakukan uji ADF yang disajikan dalam Tabel berikut:

Tabel 2. Uji ADF dari data return kurs dollar AS terhadap rupiah

	t-statistic	Probability
Augmented Dickey-Fuller test statistic	-13.03277	0.0000

Pada Tabel 2, dapat diketahui bahwa p -value dari uji ADF adalah sebesar 0.0000, karena p -value $< \alpha$ hal ini mengakibatkan data *return* kurs dollar AS terhadap rupiah stasioner. Dengan kata lain fluktuasi rata-rata dan ragam data *return* kurs dollar AS terhadap rupiah konstan.

4.2 Penentuan State

Menentukan *state* dengan mempertimbangkan μ_{s_t} . Dari μ_{s_t} dapat ditentukan *state* 0 kurs mengalami apresiasi dan *state* 1 kurs mengalami depresiasi atau sebaliknya. Dari data *return* kurs dollar AS didapat μ_0 sebesar

0.000131 dan μ_1 sebesar -0.000277 karena $\mu_0 > \mu_1$ maka pada penelitian ini *state* 0 adalah *state* apresiasi dan *state* 1 adalah *state* depresiasi.

4.3 Estimasi Parameter

Setelah data stasioner dan *state* telah ditentukan, maka langkah selanjutnya adalah mengestimasi parameter.

Parameter yang di estimasi pada penelitian ini adalah $\mu_0, \mu_1, \sigma_0^2, \sigma_1^2, p_{00}, p_{10}, \Phi_p$.

Pada model Markov *Switching Autoregressive* estimasi dilakukan dari *Autoregressive* orde 1 hingga orde 5. orde AR dipilih sampai 5 karena jika terlalu banyak orde yang dipilih maka data semakin berkurang sehingga hal ini tidak efektif untuk mencari model terbaik. Estimasi dilakukan dengan menggunakan metode MLE (*Maximum Likelihood Estimation*) digabung dengan proses *filtering* dan *smoothing* menggunakan bantuan software Python 3.6, diperoleh hasil *output* estimasi parameter MSAR yang disajikan dalam Tabel 3 sebagai berikut :

Tabel 3. Estimasi parameter Markov Switching Autoregressive (MSAR)

Model	Parameter	Koefisien	Z-statistik	Probabilitas	AIC	Jarque-Bera
MS(2) AR(1)	$\hat{\mu}_0$	-0.0002	-0.009	0.2877	-2242.846	49.53385
	$\log(\hat{\sigma}_0)$	-6.2171	-60.216	0.0000		
	$\hat{\mu}_1$	-0.0003	-0.004	0.3171		
	$\log(\hat{\sigma}_1)$	-5.2246	-68.772	0.0000		
	$\hat{\Phi}_1$	-0.0579	-0.971	0.331		
	$\hat{p}_{00}$	0.4998				
	$\hat{p}_{10}$	0.4999				
MS(2) AR(2)	$\hat{\mu}_0$	0.0004	1.753	0.030	-2253.228	43.23959
	$\log(\hat{\sigma}_0)$	-6.2886	-55.616	0.000		
	$\hat{\mu}_1$	-0.0078	-5.878	0.031		
	$\log(\hat{\sigma}_1)$	-5.2321	-68.615	0.000		
	$\hat{\Phi}_1$	-0.1952	-2.698	0.007		
	$\hat{\Phi}_2$	-0.1006	-1.423	0.045		
	$\hat{p}_{00}$	0.9362				
	$\hat{p}_{10}$	0.0450				
MS(2) AR(3)	$\hat{\mu}_0$	-0.0002	-0.015	0.988	-2225.542	51.2428
	$\log(\hat{\sigma}_0)$	-5.2368	-68.541	0.000		
	$\hat{\mu}_1$	0.0002	1.2576	0.208		
	$\log(\hat{\sigma}_1)$	-6.2983	-53.134	0.000		
	$\hat{\Phi}_1$	-0.0532	-0.890	0.374		
	$\hat{\Phi}_2$	-0.0766	-1.286	0.119		
	$\hat{\Phi}_3$	0.0229	0.384	0.701		
	$\hat{p}_{00}$	0.5257				
	$\hat{p}_{10}$	0.4999				
	$\hat{\mu}_0$	0.0004	1.672	0.095		

MS(2) AR(4)	$\log(\hat{\sigma}_0)$	-6.0017	-41.223	0.000	-2235.403	59.78371
	$\hat{\mu}_1$	-0.0077	-5.669	0.395		
	$\log(\hat{\sigma}_1)$	-5.0778	-40.201	0.000		
	$\hat{\Phi}_1$	-0.1858	-2.502	0.012		
	$\hat{\Phi}_2$	-0.1066	-1.498	0.134		
	$\hat{\Phi}_3$	-0.0078	-0.135	0.892		
	$\hat{\Phi}_4$	0.0652	1.1405	0.254		
	$\hat{p}_{00}$	0.9657				
	$\hat{p}_{10}$	0.5056				

Model	Parameter	Koefisien	Z-statistik	Probabilitas	AIC	Jarque-Bera
MS(2) AR(5)	$\hat{\mu}_0$	-0.0002	-0.026	0.979	-2208.906	51.50372
	$\log(\hat{\sigma}_0)$	-5.3057	-100.16	0.000		
	$\hat{\mu}_1$	0.0003	0.902	0.366		
	$\log(\hat{\sigma}_1)$	-6.5170	-57.781	0.000		
	$\hat{\Phi}_1$	-0.0391	-0.654	0.513		
	$\hat{\Phi}_2$	-0.0799	-1.339	0.181		
	$\hat{\Phi}_3$	0.0269	0.451	0.652		
	$\hat{\Phi}_4$	0.0451	0.759	0.448		
	$\hat{\Phi}_5$	-0.0853	-1.437	0.151		
	$\hat{p}_{00}$	0.4998				
	$\hat{p}_{10}$	0.4999				

Dari Tabel 3 yakni estimasi parameter MSAR (Markov Switching Autoregressive), dapat diketahui bahwa model yang terbaik adalah MS(2)AR(2).

4.4 Uji diagnostik model

Uji signifikansi parameter diketahui melalui nilai probabilitas parameter pada model. Lalu, pengujian normalitas residual menggunakan uji Jarque-Bera.

4.4.1 Uji signifikansi parameter model

Taraf signifikansi (α) sebesar 5 % = 0.05 parameter $\mu_0, \mu_1, \sigma_0^2, \sigma_1^2, \Phi_1, \Phi_2$

probabilitasnya < p-value yang mengakibatkan parameter signifikan. Sehingga parameter pada model MS(2)AR(2) layak digunakan untuk model MS(2)AR(2).

4.4.2 Uji Jarque-Berra

Uji normalitas residual kita menggunakan uji Jarque-Berra. Setelah Model terpilih dan semua parameter dalam model signifikan. Dilanjutkan dengan uji normalitas residual model MS(2)AR(2) dengan Uji Jarque-Berra yang ditampilkan dalam Tabel sebagai berikut:

Tabel 4. Uji Normalitas Residual

Jarque-Bera	43.23959
Probabilitas	0.1572

Uji hipotesis Jarque-berra :

H_0 : Residual berdistribusi normal

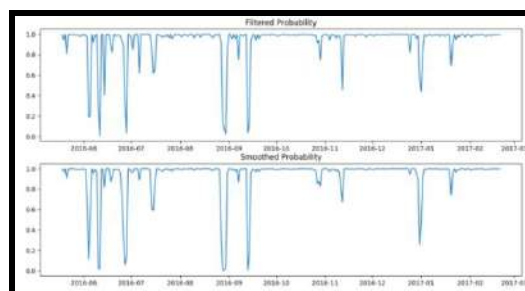
H_1 : Residual tidak berdistribusi normal

Taraf signifikansi: 5 %

Keputusan : p-value < α tolak H_0 .

Kesimpulan karena p-value > α maka tidak tolak H_0 maka dapat disimpulkan bahwa residual berdistribusi normal.

4.5 Hasil dari Proses *Filtering* dan *Smoothing* pada Estimasi Model MSAR



Gambar 4. Plot *filtering* dan *smoothing* data return kurs dollar AS

Dari Gambar 4 plot *filtering* dan *smoothing* dapat diketahui peluang setiap *state* berada pada kondisi apresiasi maupun depresiasi. Pada bulan Mei data return berfluktuasi pada *state* 0 yakni apresiasi. Selanjutnya, pada bulan Juli data return mengalami penurunan sehingga pada bulan ini terjadi depresiasi. Awal bulan Juli kedudukan masih berada pada kondisi depresiasi. Pertengahan hingga akhir bulan Juli *state* berada pada kondisi peningkatan nilai kurs. Pada bulan agustus data *return* pun masih berada pada kondisi apresiasi. Data *return* berada pada *state* 1 yakni depresiasi pada awal bulan September, pertengahan bulan September kondisi fluktuasi meningkat, kemudian pada akhir bulan September kembali pada kondisi depresiasi. Data *retun* pada bulan Oktober mengalami peningkatan kondisi fluktuasi (apresiasi) dan keadaan ini bertahan selama 3.5 bulan sampai awal bulan Januari dan mengalami depresiasi pada akhir bulan Januari. Pada bulan Februari kondisi meningkat dan pada bulan ini data *return* mengalami kondisi Apresiasi.

Berdasarkan estimasi parameter model MSAR, dapat diketahui matriks transisi yang dihasilkan dari model yang kita pilih yakni MS(2)AR(2) adalah sebagai berikut :

$$P_{ij} = \begin{bmatrix} 0.936210 & 0.063790 \\ 0.045036 & 0.954964 \end{bmatrix}$$

Peluang transisi kondisi apresiasi tetap bertahan adalah sebesar 0.936210, peluang dari apresiasi berpindah ke depresiasi adalah sebesar 0.063790, sebaliknya peluang dari depresiasi berpindah ke apresiasi adalah sebesar 0.45036 dan peluang kondisi depresiasi tetap bertahan adalah sebesar 0.954964.

Hasil *output* estimasi parameter juga dapat diketahui lama durasi *state* apresiasi bertahan adalah berkisar 15 hari, sedangkan lama durasi kondisi depresiasi bertahan adalah berkisar 22 hari. Perhitungan untuk lamanya kondisi apresiasi maupun depresiasi bertahan adalah dengan menggunakan persamaan $E(D) = \frac{1}{1-P_{jj}}$.

Untuk rata-rata durasi lamanya kondisi apresiasi bertahan adalah

$$E(D) = \frac{1}{1-P_{jj}} = \frac{1}{1-P_{00}} = \frac{1}{1-0.936210} = \frac{1}{0.06379} = 15.67$$

Sedangkan untuk rata-rata durasi lamanya kondisi depresiasi bertahan adalah

$$E(D) = \frac{1}{1-P_{jj}} = \frac{1}{1-P_{11}} = \frac{1}{1-0.954964} = \frac{1}{0.045036} = 22.20$$

sehingga model MS(2)AR(2) dapat ditulis :

$$(y_t - \mu_{s_t}) = -0.1952(y_{t-1} - \mu_{s_{t-1}}) - 0.1006(y_{t-2} - \mu_{s_{t-2}}) + e_t$$

$$e_t \sim N(0, \sigma_{s_t}^2)$$

dengan rata-rata dari masing-masing *state* adalah sebagai berikut :

$$\mu_{s_t} \begin{cases} \mu_0 = 0.0004 \\ \mu_1 = -0.0078 \end{cases}$$

dan ragam dari masing-masing *state* adalah :

$$\sigma_{s_t}^2 \begin{cases} \sigma_0^2 = (\text{antilog}(-6.2886))^2 = 2.6472 \times 10^{-13} \\ \sigma_1^2 = (\text{antilog}(-5.2321))^2 = 3.4339 \times 10^{-11} \end{cases}$$

untuk $s_t = 0$ (apresiasi)

untuk $s_t = 1$ (depresiasi)

5. KESIMPULAN

Berdasarkan penelitian yang dilakukan diperoleh kesimpulan sebagai berikut :

1. Data kurs dollar AS terhadap rupiah tidak stasioner dalam ragam maupun *mean*. Sehingga, dilakukan transformasi dan diferensiasi untuk memperoleh data yang stasioner pada ragam dan *mean*. Model Markov *Switching Autoregressive* (MSAR) yang terbaik adalah MS(2)AR(2), yakni :

$$(y_t - \mu_{s_t}) = -0.1952(y_{t-1} - \mu_{s_{t-1}}) - 0.1006(y_{t-2} - \mu_{s_{t-2}}) + e_t$$

dengan $e_t \sim N(0, \sigma_{s_t}^2)$ dan rata-rata masing-masing *state* adalah sebagai berikut :

$$\mu_{s_t} \begin{cases} \mu_0 = 0.0004 \\ \mu_1 = -0.0078 \end{cases}$$

serta ragam dari masing-masing *state* adalah :

$$\sigma_{s_t}^2 \begin{cases} \sigma_0^2 = 2.6472 \times 10^{-13} \\ \sigma_1^2 = 3.4339 \times 10^{-11} \end{cases}$$

untuk $s_t = 0$ (apresiasi)

untuk $s_t = 1$ (depresiasi)

2. Peluang *return* kurs dollar AS terhadap rupiah saat t bertahan pada kondisi apresiasi adalah sebesar 0.936210 dan peluang *return* kurs dollar AS dari kondisi apresiasi berpindah ke depresiasi adalah 0.063790.
Peluang *return* kurs dollar AS terhadap rupiah saat t bertahan pada kondisi depresiasi adalah 0.954964 dan peluang *return* kurs dollar AS terhadap rupiah dari depresiasi pindah ke apresiasi adalah 0.045036.
3. Rata-rata lama durasi kondisi *return* kurs dollar AS terhadap rupiah mengalami apresiasi adalah 15 hari dan rata-rata lama kondisi *return* kurs dollar AS terhadap rupiah yang mengalami depresiasi adalah 22 hari.

KEPUSTAKAAN

- [1] Hamilton, J.D. 1994. *Time Series Analysis*. Princeton University Press.
- [2] Wei, W.W. 2006. *Time Series Analysis : Univariate and Multivariate Methods (2nd ed)*. Pearson, New York.
- [3] Franses, P.H., Dijk, D.V., and Opschoor, A. 2014. *Time Series Model : for Bussiness and Economic Forecasting 2nd edition*. Cambridge, University Press.
- [4] Brockwell, P.J., and Davis, R.A. 2002. *Introduction to Time Series and Forecasting*. Springer : New York.
- [5] Cox, D.R., and Miller, H.D., 1965. *The Theory of Stochastic Process*. Chapman and Hall : London.
- [6] Kim, C.J and Nelson C.R, 1999. *State Space Models with Regime Switching, Classical and Gibbs Sampling Approaches with Applications*. Cambridge, MA: MIT Press.
- [7] Privault, N. 2013. *Understanding Markov chain : Examples and Applications*. Springer : Singapore.
- [8] Hamilton, J.D. 1989. *A New Approach to the Economic Analysis of Nonstationary Time Series and the Business Cycle*. *Econometrica* 57: 357-384.
- [9] Hamilton, J.D. 1996. *Specification Testing in Markov-Switching Time Series Model*. *Journal of Econometrics*, Vol 70: 127-157.
- [10] Azam, I. 2007. *The Effect of Model-Selection Uncertainty on Autoregressive Models Estimates*. *International Research Journal of Finance and Economics*, issue. 11, hal 80-93.
- [11] Jarque, C. M. and Berra, A.K. 1980. Efficient Tests For Normality, Homoskedasticity, and Serial Independence of Regression Residuals. *Economic Letters*. 6: 255-259.

PERBANDINGAN METODE ADAMS BASHFORTH-MOULTON DAN METODE MILNE-SIMPSON DALAM PENYELESAIAN PERSAMAAN DIFERENSIAL EULER ORDE-8

Faranika Latip¹⁾, Dorrah Azis²⁾, dan Suharsono³⁾

Universitas Lampung, Jalan Prof. Dr. Soemantri Brodjonegoro No. 1 Bandar Lampung 35145
latipfaranika@gmail.com

¹⁾Mahasiswa Jurusan Matematika FMIPA Universitas Lampung

^{2,3)}Dosen Jurusan Matematika FMIPA Universitas Lampung

ABSTRAK

Persamaan diferensial Euler adalah bentuk khusus dari persamaan diferensial linear dengan koefisien peubah. Dalam penyelesaian persamaan diferensial Euler orde-8 dimana akar-akar karakteristiknya riil berbeda, digunakan metode langkah banyak (multi steps) atau disebut juga dengan metode prediktor-korektor yaitu Adams-Bashforth Moulton dan metode Milne-Simpson sebagai metode dalam menghampiri solusi analitiknya. Bentuk umum dari persamaan diferensial Euler orde-8 :

$$\sum_{i=0}^8 a_i x^i y^{(i)} = 0$$

Ditransformasikan menjadi sistem persamaan diferensial biasa (SPDB) orde-1 dengan menggunakan transformasi $x = e^t$, dimana $x > 0$. Didapatkan sistem persamaan diferensial biasa (SPDB) orde-1 dengan nilai-nilai awal yang kemudian disubstitusikan ke dalam persamaan prediktor lalu dikoreksi pada persamaan korektor dengan pemilihan ukuran h yang tepat. Diperoleh kesimpulan bahwa kedua metode di atas dapat digunakan dalam menyelesaikan persamaan diferensial Euler orde-8. Metode Adams Bashforth-Moulton lebih akurat dalam menyelesaikan persamaan diferensial Euler orde-8 diketahui dari perbandingan jumlah galatnya, sebaliknya metode Milne-Simpson lebih efisien dalam melakukan iterasi karena lebih cepat dalam menyelesaikan persamaan diferensial Euler orde-8.

Kata kunci : Persamaan Diferensial Euler, Adams Bashforth-Moulton, Milne-Simpson

1. PENDAHULUAN

Persamaan Diferensial (*differential equation*) adalah persamaan yang melibatkan variabel-variabel tak bebas dan turunan-turunannya terhadap variabel-variabel bebas [1]. Salah satu bentuk persamaan diferensial biasa yaitu persamaan diferensial Euler. Persamaan diferensial Euler adalah salah satu bentuk khusus dari persamaan diferensial linear dengan koefisien peubah [2]. Setiap bentuk persamaan differensial mempunyai metode penyelesaian yang berbeda. Metode lain yang dapat digunakan dalam menyelesaikan persamaan diferensial biasa orde tinggi yaitu metode numerik.

Metode numerik merupakan metode yang digunakan untuk memformulasikan persoalan matematika sehingga dapat diselesaikan dengan operasi perhitungan atau aritmatika biasa (tambah, kurang, kali, dan bagi) [3]. Metode numerik dibagi menjadi dua yaitu metode langkah tunggal (*one step*) dan metode langkah banyak (*multi steps*). Metode yang termasuk metode langkah tunggal adalah metode Euler, metode Heun, dan metode Runge-Kutta. Sedangkan metode yang termasuk metode langkah banyak adalah metode Adams Bashforth-Moulton, metode Milne-Simpson, dan metode Hamming. Semakin tinggi orde yang muncul pada persamaan diferensial maka akan semakin sulit ditemukan solusinya secara analitik, sehingga penyelesaian dengan menggunakan metode numerik merupakan metode yang digunakan untuk memperoleh solusi pendekatannya .

Pada penelitian ini hanya difokuskan dalam menyelesaikan persamaan diferensial Euler orde-8 dimana akar-akar karakteristiknya riil berbeda lalu membandingkan solusi dari metode Adams Bashforth-Moulton dan metode Milne-Simpson sebagai metode lain dalam penyelesaiannya.

2. LANDASAN TEORI

PERSAMAAN DIFERENSIAL EULER CAUCHY ORDE TINGGI HOMOGEN

Persamaan diferensial linear orde tinggi homogen dengan koefisien peubah dikatakan sebagai persamaan diferensial Euler Cauchy, jika persamaan diferensial tersebut berbentuk [2]:

$$a_n x^n y^{(n)} + a_{n-1} x^{n-1} y^{(n-1)} + \dots + a_1 x y' + a_0 y = 0 \quad (1)$$

dimana $a_n, a_{n-1}, \dots, a_1, a_0$ merupakan konstanta-konstanta dan $a_n \neq 0$. Dengan menggunakan pendekatan yang dikembangkan pada persamaan diferensial Euler Cauchy homogen orde dua, untuk menentukan penyelesaian umumnya, ambil basis-basis penyelesaiannya berbentuk :

$$y = x^m$$

Dengan menurunkan persamaan $y = x^m$, terhadap x sebanyak n kali dan,

dengan mensubstitusikan basis $y = x^m$ dan turunan-turunannya $y', y'', \dots, y^{(n)}$ ke persamaan (1), didapatkan

$$y = c_1 x^{m_1} + c_2 x^{m_2} + c_3 x^{m_3} + \dots + c_n x^{m_n} \quad (2)$$

METODE ADAMS BASHFORTH-MOULTON

Metode prediktor-korektor Adams Bashforth-Moulton (ABM) adalah metode langkah banyak yang diturunkan dari teorema dasar kalkulus [4] :

$$y(t_{k+1}) = y(t_k) + \int_{t_k}^{t_{k+1}} f(t, y(t)) dt \quad (3)$$

Persamaan prediktor Adams Bashforth :

$$P_{k+1} = y_k + \frac{h}{24} (-9f_{k-3} + 37f_{k-2} - 59f_{k-1} + 55f_k) \quad (4)$$

Persamaan korektor Adams Moulton :

$$y_{k+1} = y_k + \frac{h}{24} (f_{k-2} - 5f_{k-1} + 19f_k + 9f_{k+1}) \quad (5)$$

METODE MILNE-SIMPSON

Metode prediktor-korektor lain yang terkenal adalah metode Milne-Simpson [4]. Prediktor berdasarkan integrasi dari $f(t, y(t))$ pada interval $[t_{k-3}, t_{k+1}]$:

$$y(t_{k+1}) = y(t_{k-3}) + \int_{t_{k-3}}^{t_{k+1}} f(t, y(t)) dt \quad (6)$$

Prediktor Milne:

$$P_{k+1} = y_{k-3} + \frac{4h}{3} (2f_{k-2} - f_{k-1} + 2f_k) \quad (7)$$

Persamaan Simpson :

$$y_{k+1} = y_{k-1} + \frac{h}{3}(f_{k-1} + 4f_k + f_{k+1}) \quad (8)$$

ANALISIS ERROR

Dalam melakukan analisis numerik perlu untuk memperhatikan bahwa solusi yang dihitung bukanlah solusi analitik. Ketelitian dari solusi numerik dapat mengurangi nilai *error*.

Definisi Misalkan $\hat{p}$ adalah aproksimasi (pendekatan) ke p . *Error* mutlak adalah $E_p = |p - \hat{p}|$ dan *error* relatif adalah $R_p = |p - \hat{p}|/|p|$, dinyatakan bahwa $p \neq 0$ [4].

3. METODOLOGI PENELITIAN

Adapun langkah-langkah yang dilakukan dalam penelitian adalah sebagai berikut :

1. Studi pustaka yaitu mencari referensi, menelaah, dan mengkaji yang berhubungan dengan penelitian ini.
2. Mentransformasikan bentuk umum persamaan diferensial Euler orde-8 ke dalam bentuk sistem persamaan diferensial biasa (SPDB) orde-1 dengan koefisien konstanta untuk menyelesaikan secara numerik.
3. Menyelesaikan contoh kasus dengan cara, yaitu :
 - a. Mentransformasikan contoh kasus dari persamaan diferensial Euler orde-8 ke dalam bentuk sistem persamaan diferensial orde-1 dengan koefisien konstanta untuk menyelesaikannya secara numerik.
 - b. Menentukan solusi umum analitik dari contoh kasus persamaan diferensial Euler orde-8 dan mentransformasikannya ke dalam bentuk $x = e^t$ sehingga didapatkan solusi khusus analitiknya.
 - c. Mencari solusi awal dari solusi khusus analitik contoh kasus persamaan diferensial Euler orde-8 dengan metode Runge-Kutta orde-4.
 - d. Menentukan algoritma dan mencari solusi numerik dari contoh kasus persamaan diferensial Euler orde-8 menggunakan metode banyak langkah (*multi steps*) yaitu metode Adams Bashforth-Moulton dan metode Milne-Simpson dengan *software* MATLAB R2013b.
 - e. Membandingkan solusi analitik dan solusi numerik dari persamaan diferensial Euler orde-8 menggunakan kedua metode dengan *software* MATLAB R2013b sehingga dapat diketahui metode manakah yang lebih baik dalam menyelesaikan persamaan diferensial Euler orde-8.

4. HASIL DAN PEMBAHASAN

HASIL

ALGORITMA METODE ADAMS BASHFORTH-MOULTON

Algoritma penyelesaian persamaan diferensial Euler orde-8 menggunakan metode Adams Bashforth-Moulton adalah sebagai berikut :

1. Diketahui masalah nilai awal dari persamaan diferensial Euler orde-8 yang telah ditransformasikan ke dalam sistem persamaan diferensial orde-1,

$$\bar{Y}^{(1)} = \frac{d\bar{Y}}{dt} = f(t, \bar{Y}) \quad , \text{ dengan nilai awal } \bar{Y}(0) = \hat{Y}_{(0)}$$

dimana $\bar{Y} = [Y_0, Y_1, \dots, Y_7]^T$, dengan ukuran langkah h dan $t_{(i+1)} = t_{(i)} + h$.

2. Hitung empat solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4.
3. Tentukan nilai-nilai $f_{(i)}, f_{(i-1)}, f_{(i-2)}$, dan $f_{(i-3)}$ dengan $i = 3, 4, \dots, 20$
4. Tentukan solusi numerik menggunakan metode Adams Bashforth :

$$\bar{P}_{(i+1)} = \hat{Y}_{(i)} + \frac{h}{24} (-9f_{(i-3)} + 37f_{(i-2)} - 59f_{(i-1)} + 55f_{(i)})$$

5. Hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ dan disubstitusikan pada metode Adams Moulton.
6. Hitung solusi numerik menggunakan metode Adams Moulton :

$$\hat{Y}_{(i+1)} = \hat{Y}_{(i)} + \frac{h}{24} (f_{(i-2)} - 5f_{(i-1)} + 19f_{(i)} + 9f_{(i+1)})$$

dengan perhitungan galat mutlak sebagai berikut :

$$E_{Y_{(i)}} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

ALGORITMA METODE MILNE-SIMPSON

Algoritma penyelesaian persamaan diferensial Euler orde-8 menggunakan metode Milne-Simpson adalah sebagai berikut :

1. Diketahui masalah nilai awal dari persamaan diferensial Euler orde-8 yang telah ditransformasikan ke dalam sistem persamaan diferensial orde-1 ,

$$\bar{Y}^{(1)} = \frac{d\bar{Y}}{dt} = f(t, \bar{Y}) \quad , \text{ dengan nilai awal } \bar{Y}(0) = \hat{Y}_{(0)}$$

dimana $\bar{Y} = [Y_0, Y_1, \dots, Y_7]^T$, dengan ukuran langkah h dan $t_{(i+1)} = t_i + h$.

2. Hitung empat solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4.
3. Tentukan nilai-nilai $f_{(i)}, f_{(i-1)}$, dan $f_{(i-2)}$ dengan $i = 3, 4, \dots, 20$
4. Tentukan solusi numerik menggunakan metode Milne :

$$\bar{P}_{(i+1)} = \hat{Y}_{(i-3)} + \frac{4h}{3} (2f_{(i-2)} - f_{(i-1)} + 2f_{(i)})$$

5. Hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ dan disubstitusikan pada metode Simpson.
6. Hitung solusi numerik menggunakan metode Simpson :

$$\hat{Y}_{(i+1)} = \hat{Y}_{(i-1)} + \frac{h}{3} (f_{(i-1)} + 4f_{(i)} + 9f_{(i+1)})$$

7. Kesalahan (*error*) yang terjadi adalah sebagai berikut :

$$E_{Y_{(i)}} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

Contoh Kasus 1 :

$$x^8 y^{(8)} + 9x^7 y^{(7)} + 7x^6 y^{(6)} - 15x^5 y^{(5)} + 13x^4 y^{(4)} - 6x^3 y^{(3)} + 6x^2 y^{(2)} - 7x^1 y^{(1)} + 14y = 0 \quad (9)$$

Sistem persamaan diferensial biasa (SPDB) orde-1 adalah sebagai berikut :

$$\begin{aligned} Y_0^{(1)} &= \frac{dY_0}{dt} = Y^{(1)} = Y_1 \\ Y_1^{(1)} &= \frac{dY_1}{dt} = Y^{(2)} = Y_2 \\ Y_2^{(1)} &= \frac{dY_2}{dt} = Y^{(3)} = Y_3 \\ Y_3^{(1)} &= \frac{dY_3}{dt} = Y^{(4)} = Y_4 \\ Y_4^{(1)} &= \frac{dY_4}{dt} = Y^{(5)} = Y_5 \\ Y_5^{(1)} &= \frac{dY_5}{dt} = Y^{(6)} = Y_6 \\ Y_6^{(1)} &= \frac{dY_6}{dt} = Y^{(7)} = Y_7 \\ Y_7^{(1)} &= \frac{dY_7}{dt} = Y^{(8)} \\ &= -\frac{k_7}{k_8} Y_7 - \frac{k_6}{k_8} Y_6 - \frac{k_5}{k_8} Y_5 - \frac{k_4}{k_8} Y_4 - \frac{k_3}{k_8} Y_3 - \frac{k_2}{k_8} Y_2 - \\ &\quad \frac{k_1}{k_8} Y_1 + \frac{k_0}{k_8} Y_0 \\ &= 19Y_7 - 140Y_6 + 505Y_5 - 912Y_4 + 700Y_3 - 27Y_2 \\ &\quad - 137Y_1 + 14Y_0 \end{aligned} \quad (10)$$

Sehingga diperoleh solusi umum analitik dari persamaan di atas sebagai berikut :

$$y(x) = c_1 x^{(-0,1133)} + c_2 x^{(-0,2823)} + c_3 x^{(1,1146)} + c_4 x^{(1,6452)} + c_5 x^{(2,6185)} + c_6 x^{(3,1288)} + c_7 x^{(4,7254)} + c_8 x^{(6,1629)} \quad (11)$$

Dimisalkan sebarang $c_1, c_2, \dots, c_8$, kemudian solusi umum analitik ditransformasikan ke dalam bentuk $x = e^t$, diperoleh solusi khusus analitik dari persamaan di atas menjadi :

$$y(x) = 2e^{t(-0,1133)} + 2e^{t(-0,2823)} + e^{t(1,1146)} + e^{t(1,6452)} + 2e^{t(2,6185)} + 2e^{t(3,1288)} + 2e^{t(3,1288)} + e^{t(4,7254)} + e^{t(6,1629)} \quad (12)$$

Dengan nilai-nilai awalnya :

$$\begin{aligned} Y_0(0) &= Y_{0(0)} = 12 & Y_1(0) &= Y_{1(0)} = 24.3515 \\ Y_2(0) &= Y_{2(0)} = 97.7366 & Y_3(0) &= Y_{3(0)} = 442.5461 \\ Y_4(0) &= Y_{4(0)} = 2235.7554 & Y_5(0) &= Y_{5(0)} = 12106.2351 \\ Y_6(0) &= Y_{6(0)} = 68467.3887 & Y_7(0) &= Y_{7(0)} = 397876.1827 \end{aligned}$$

Selanjutnya, digunakan metode Adams Bashforth-Moulton dan metode Milne-Simpson.

Penyelesaian dengan menggunakan metode Adams Bashforth-Moulton, telah diketahui bahwa sistem persamaan diferensial orde-1 pada persamaan (13) yaitu :

$$f(t, \bar{Y}) = \begin{bmatrix} Y_1, Y_2, \dots, Y_7, (19Y_7 - 140Y_6 + 505Y_5 - 912Y_4 + 700Y_3 - 27Y_2) \\ -137Y_1 + 14Y_0 \end{bmatrix}$$

Ditentukan solusi masalah nilai awal dengan nilai awal $t_{(0)} = 0$ dan nilai-nilai awal, yaitu :

$$\begin{aligned} Y_0(0) = Y_{0(0)} &= 12 & Y_1(0) = Y_{1(0)} &= 24.3515 \\ Y_2(0) = Y_{2(0)} &= 97.7366 & Y_3(0) = Y_{3(0)} &= 442.5461 \\ Y_4(0) = Y_{4(0)} &= 2235.7554 & Y_5(0) = Y_{5(0)} &= 12106.2351 \\ Y_6(0) = Y_{6(0)} &= 68467.3887 & Y_7(0) = Y_{7(0)} &= 397876.1827 \end{aligned}$$

Serta $h = 0.05$ dan $t_{(i+1)} = t_{(i)} + h$. Akan dicari solusi masalah nilai awal di atas pada interval $[0,1]$ untuk memperoleh solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4. Setelah diperoleh solusi awal menggunakan Runge-Kutta orde-4, tentukan nilai-nilai $f_{(i)}, f_{(i-1)}, f_{(i-2)}$, dan $f_{(i-3)}$ kemudian hasil tersebut disubstitusikan ke metode Adams Bashforth :

$$\begin{aligned} \bar{P}_{(i+1)} &= \hat{Y}_{(i)} + \frac{h}{24}(-9f_{(i-3)} + 37f_{(i-2)} - 59f_{(i-1)} + 55f_{(i)}) \\ \bar{P}_{(4)} &= \hat{Y}_{(3)} + \frac{h}{24}(-9f_{(0)} + 37f_{(1)} - 59f_{(2)} + 55f_{(3)}) \end{aligned}$$

Setelah didapatkan nilai prediksi, hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ kemudian disubstitusikan ke persamaan Adams Moulton untuk memperoleh nilai koreksi :

$$\begin{aligned} \hat{Y}_{(i+1)} &= \hat{Y}_{(i)} + \frac{h}{24}(f_{(i-2)} - 5f_{(i-1)} + 19f_{(i)} + 9f_{(i+1)}) \\ \hat{Y}_{(4)} &= \hat{Y}_{(3)} + \frac{h}{24}(f_{(1)} - 5f_{(2)} + 19f_{(3)} + 9f_{(4)}) \end{aligned}$$

Kemudian hitung galat mutlak :

$$E_{Y_{(i)}} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

didapatkan solusi analitik dan solusi numerik dari contoh kasus 1 persamaan diferensial Euler orde-8 sebagai berikut :

Tabel 4.1. Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 1 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Adams Bashforth-Moulton

i	$t_{(i)}$	Solusi Numerik ABM ($\hat{Y}_{0(i)}$)	Solusi Analitik ($Y(t_{(i)})$)	Galat (Error) ($E_{Y_{(i)}}$)
0	0	12.0000000000000	12.0000000000000	0.0000000000000
1	0.05	13.3495476883854	13.3495808684785	0.0000331800931
2	0.10	15.0079308991162	15.0080192175174	0.0000883184012
3	0.15	17.0570744923525	17.0572509854047	0.0001764930522
4	0.20	19.6034266711746	19.6036972309463	0.0002705597717
5	0.25	22.7857135452836	22.7861221169730	0.0004085716893

6	0.30	26.7855073680835	26.7861214272842	0.0006140592008
7	0.35	31.8412293426238	31.8421446816390	0.0009153390152
8	0.40	38.2669165390816	38.2682688100184	0.0013522709367
9	0.45	46.4773894767133	46.4793671623633	0.0019776856500
10	0.50	57.0220351495490	57.0248937002438	0.0028585506948
11	0.55	70.6302071411936	70.6342818966311	0.0040747554375
12	0.60	88.2723009897650	88.2780135424466	0.0057125526816
13	0.65	111.2419944936692	111.2498425446593	0.0078480509902
14	0.70	141.2670830583940	141.2775962084315	0.0105131500374
15	0.75	180.6589706343647	180.6726023887393	0.0136317543746
16	0.80	232.5144437667462	232.5313509731758	0.0169072064296
17	0.85	300.9881945375205	301.0078260613651	0.0196315238446
18	0.90	391.6611224609805	391.6814940757838	0.0203716148033
19	0.95	512.0383535036823	512.0548183893688	0.0164648856866
20	1.00	672.2230058080054	672.2262290757910	0.0032232677856

Penyelesaian dengan menggunakan metode Milne-Simpson, telah diketahui bahwa sistem persamaan diferensial orde-1 pada persamaan (13), yaitu :

$$f(t, \bar{Y}) = \begin{bmatrix} Y_1, Y_2, \dots, Y_7, (19Y_7 - 140Y_6 + 505Y_5 - 912Y_4 + 700Y_3 - 27Y_2) \\ -137Y_1 + 14Y_0 \end{bmatrix}$$

Ditentukan solusi masalah nilai awal dengan nilai awal $t_0 = 0$ dan nilai-nilai awal yaitu :

$$\begin{aligned} Y_0(0) = Y_{0(0)} &= 12 & Y_1(0) = Y_{1(0)} &= 24.3515 \\ Y_2(0) = Y_{2(0)} &= 97.7366 & Y_3(0) = Y_{3(0)} &= 442.5461 \\ Y_4(0) = Y_{4(0)} &= 2235.7554 & Y_5(0) = Y_{5(0)} &= 12106.2351 \\ Y_6(0) = Y_{6(0)} &= 68467.3887 & Y_7(0) = Y_{7(0)} &= 397876.1827 \end{aligned}$$

Serta $h = 0.05$ dan $t_{(i+1)} = t_{(i)} + h$. Akan dicari solusi masalah nilai awal di atas pada interval $[0,1]$ untuk memperoleh solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4. Setelah diperoleh solusi awal menggunakan Runge-Kutta orde-4, tentukan nilai-nilai $f_{(i)}, f_{(i-1)}$, dan $f_{(i-2)}$ kemudian hasil tersebut disubstitusikan ke metode Milne :

$$\bar{P}_{(i+1)} = \hat{Y}_{(i-3)} + \frac{4h}{3}(2f_{(i-2)} - f_{(i-1)} + 2f_{(i)})$$

$$\bar{P}_{(4)} = \hat{Y}_{(0)} + \frac{4h}{3}(2f_{(1)} - f_{(2)} + 2f_{(3)})$$

Setelah didapatkan nilai prediksi, hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ kemudian disubstitusikan ke persamaan korektor Simpson untuk memperoleh nilai koreksi :

$$\hat{Y}_{(i+1)} = \hat{Y}_{(i-1)} + \frac{h}{3}(f_{(i-1)} + 4f_{(i)} + 9f_{(i+1)})$$

$$\hat{Y}_{(4)} = \hat{Y}_{(2)} + \frac{h}{3}(f_{(2)} + 4f_{(3)} + 9f_{(4)})$$

Kemudian dihitung galat mutlak :

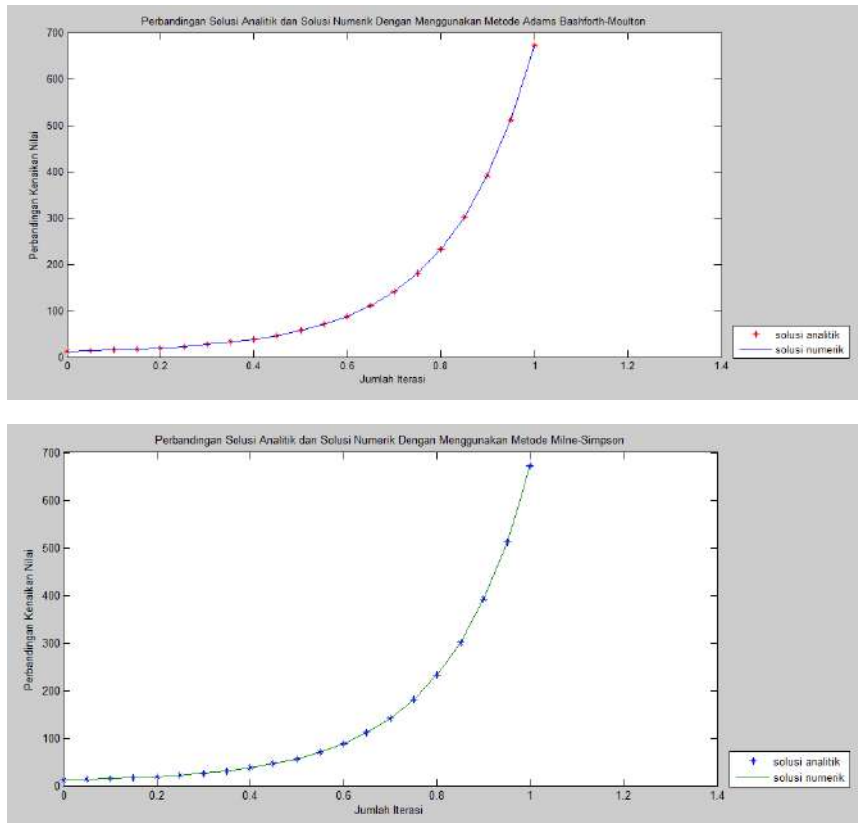
$$E_{Y(i)} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

didapatkan solusi analitik dan solusi numerik dari contoh kasus 1 persamaan diferensial Euler orde-8 sebagai berikut :

Tabel 4.2. Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 1 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Milne-Simpson

i	$t_{(i)}$	Solusi Numerik MS ($\hat{Y}_{0(i)}$)	Solusi Analitik ($Y(t_{(i)})$)	Galat (Error) ($E_{Y(i)}$)
0	0	12.0000000000000	12.0000000000000	0.0000000000000
1	0.05	13.3495476883854	13.3495808684785	0.0000331800931
2	0.10	15.0079308991162	15.0080192175174	0.0000883184012
3	0.15	17.0570744923525	17.0572509854047	0.0001764930522
4	0.20	19.6034178756786	19.6036972309463	0.0002793552677
5	0.25	22.7856570817454	22.7861221169730	0.0004650352276
6	0.30	26.7854006885009	26.7861214272842	0.0007207387833
7	0.35	31.8410253060456	31.8421446816390	0.0011193755934
8	0.40	38.2665831203740	38.2682688100184	0.0016856896443
9	0.45	46.4768479714010	46.4793671623633	0.0025191909623
10	0.50	57.0212000320491	57.0248937002438	0.0036936681947
11	0.55	70.6289305887431	70.6342818966311	0.0053513078880
12	0.60	88.2703927104533	88.2780135424466	0.0076208319933
13	0.65	111.2391632147506	111.2498425446593	0.0106793299087
14	0.70	141.2629302470732	141.2775962084315	0.0146659613583
15	0.75	180.6529141807399	180.6726023887393	0.0196882079994
16	0.80	232.5056702627636	232.5313509731718	0.0256807104122
17	0.85	300.9755390704375	301.0078260613651	0.0322869909276
18	0.90	391.6429468694513	391.6814940757838	0.0385472063325
19	0.95	512.0123332984838	512.0548183893688	0.0424850908851
20	1.00	672.1858688254897	672.2262290757910	0.0403602503013

Berikut adalah gambar grafik perbandingan antara solusi analitik dan solusi numerik untuk contoh kasus 1 persamaan diferensial Euler orde-8 dengan menggunakan metode Adams Bashforth-Moulton dan metode Milne-Simpson :



Gambar 4.1. Grafik Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 1 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Adams Bashforth-Moulton dan Metode Milne-Simpson

Contoh Kasus 2 :

$$\begin{aligned}
 x^8 y^{(8)} + 10x^7 y^{(7)} + 14x^6 y^{(6)} - 13x^5 y^{(5)} + 13x^4 y^{(4)} - x^3 y^{(3)} \\
 + x^2 y^{(2)} + 4x^1 y^{(1)} - 10y = 0
 \end{aligned}
 \tag{13}$$

Didapatkan sistem persamaan diferensial biasa (SPDB) orde-1 adalah sebagai berikut :

$$Y_0^{(1)} = \frac{dY_0}{dt} = Y^{(1)} = Y_1$$

$$Y_1^{(1)} = \frac{dY_1}{dt} = Y^{(2)} = Y_2$$

$$Y_2^{(1)} = \frac{dY_2}{dt} = Y^{(3)} = Y_3$$

$$Y_3^{(1)} = \frac{dY_3}{dt} = Y^{(4)} = Y_4$$

$$Y_4^{(1)} = \frac{dY_4}{dt} = Y^{(5)} = Y_5$$

$$Y_5^{(1)} = \frac{dY_5}{dt} = Y^{(6)} = Y_6$$

$$Y_6^{(1)} = \frac{dY_6}{dt} = Y^{(7)} = Y_7$$

$$\begin{aligned}
 Y_7^{(1)} &= \frac{dY_7}{dt} = Y^{(8)} \\
 &= -\frac{k_7}{k_8}Y_7 - \frac{k_6}{k_8}Y_6 - \frac{k_5}{k_8}Y_5 - \frac{k_4}{k_8}Y_4 - \frac{k_3}{k_8}Y_3 - \frac{k_2}{k_8}Y_2 - \\
 &\quad \frac{k_1}{k_8}Y_1 + \frac{k_0}{k_8}Y_0 \\
 &= 18Y_7 - 126Y_6 + 433Y_5 - 752Y_4 + 576Y_3 - \\
 &\quad 61Y_2 - 91Y_1 - 10Y_0
 \end{aligned} \tag{14}$$

Sehingga diperoleh solusi umum analitik dari persamaan di atas sebagai berikut :

$$\begin{aligned}
 y(x) &= c_1x^{(-0.3367)} + c_2x^{(0.1090)} + c_3x^{(0.8529)} + c_4x^{(1.9803)} + \\
 &\quad c_5x^{(2.0000)} + c_6x^{(2.9803)} + c_7x^{(4.9661)} + c_8x^{(5.4479)}
 \end{aligned} \tag{15}$$

Dimisalkan sebarang $c_1, c_2, \dots, c_8$, kemudian solusi umum analitik ditransformasikan ke dalam bentuk $x = e^t$, diperoleh solusi khusus analitik dari persamaan di atas menjadi :

$$\begin{aligned}
 Y(t) &= 2e^{t(-0.3367)} + e^{t(0.1090)} + 2e^{t(0.8529)} + e^{t(1.9803)} + \\
 &\quad 2e^{t(2.0000)} + e^{t(2.9803)} + 2e^{t(4.9661)} + e^{t(5.4479)}
 \end{aligned} \tag{16}$$

Dengan nilai-nilai awalnya :

$$\begin{aligned}
 Y_0(0) &= Y_{0(0)} = 12 & Y_1(0) &= Y_{1(0)} = 25.4821 \\
 Y_2(0) &= Y_{2(0)} = 101.5011 & Y_3(0) &= Y_{3(0)} = 458.0442 \\
 Y_4(0) &= Y_{4(0)} = 2224.6790 & Y_5(0) &= Y_{5(0)} = 11170.3964 \\
 Y_6(0) &= Y_{6(0)} = 57034.0952 & Y_7(0) &= Y_{7(0)} = 293878.8247
 \end{aligned}$$

Selanjutnya, digunakan metode Adams Bashforth-Moulton dan metode Milne-Simpson.

Penyelesaian dengan menggunakan metode Adams Bashforth-Moulton, telah diketahui bahwa sistem persamaan diferensial orde-1 pada persamaan (17) yaitu :

$$f(t, \bar{Y}) = \begin{bmatrix} Y_1, Y_2, \dots, Y_7, (18Y_7 - 126Y_6 + 433Y_5 - 752Y_4 + 576Y_3 - 61Y_2 - \\ -91Y_1 - 10Y_0) \end{bmatrix}$$

Ditentukan solusi masalah nilai awal dengan nilai awal $t_0 = 0$ dan nilai-nilai awal yaitu :

$$\begin{aligned}
 Y_0(0) &= Y_{0(0)} = 12 & Y_1(0) &= Y_{1(0)} = 25.4821 \\
 Y_2(0) &= Y_{2(0)} = 101.5011 & Y_3(0) &= Y_{3(0)} = 458.0442 \\
 Y_4(0) &= Y_{4(0)} = 2224.6790 & Y_5(0) &= Y_{5(0)} = 11170.3964 \\
 Y_6(0) &= Y_{6(0)} = 57034.0952 & Y_7(0) &= Y_{7(0)} = 293878.8247
 \end{aligned}$$

Serta $h = 0.05$ dan $t_{(i+1)} = t_{(i)} + h$. Akan dicari solusi masalah nilai awal di atas pada interval $[0,1]$ untuk memperoleh solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4. Setelah diperoleh solusi awal menggunakan Runge-Kutta orde-4, tentukan nilai-nilai $f_{(i)}, f_{(i-1)}, f_{(i-2)}$, dan $f_{(i-3)}$ kemudian hasil tersebut disubstitusikan ke metode Adams bashforth :

$$\bar{P}_{(i+1)} = \hat{Y}_{(i)} + \frac{h}{24}(-9f_{(i-3)} + 37f_{(i-2)} - 59f_{(i-1)} + 55f_{(i)})$$

$$\bar{P}_{(4)} = \hat{Y}_{(3)} + \frac{h}{24}(-9f_{(0)} + 37f_{(1)} - 59f_{(2)} + 55f_{(3)})$$

Setelah didapatkan nilai prediksi, hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ kemudian disubstitusikan ke persamaan Adams Moulton untuk memperoleh nilai koreksi :

$$\hat{Y}_{(i+1)} = \hat{Y}_{(i)} + \frac{h}{24}(f_{(i-2)} - 5f_{(i-1)} + 19f_{(i)} + 9f_{(i+1)})$$

$$\hat{Y}_{(4)} = \hat{Y}_{(3)} + \frac{h}{24}(f_{(1)} - 5f_{(2)} + 19f_{(3)} + 9f_{(4)})$$

Kemudian dihitung galat mutlak :

$$E_{Y_{(i)}} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

didapatkan solusi analitik dan solusi numerik dari contoh kasus 2 persamaan diferensial Euler orde-8 sebagai berikut :

Tabel 4.3. Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 2 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Adams Bashforth-Moulton

i	t _(i)	Solusi Numerik ABM ($\hat{Y}_{0(i)}$)	Solusi Analitik ($Y(t_{(i)})$)	Galat (Error) ($E_{Y_{(i)}}$)
0	0	12.00000000000000	12.00000000000000	0.00000000000000
1	0.05	13.4111033059896	13.4111337830808	0.0000304770912
2	0.10	15.1422635610499	15.1423424266375	0.0000788655877
3	0.15	17.2767083056203	17.2768613001565	0.0001529945362
4	0.20	19.9209849147927	19.9211942018391	0.0002092870464
5	0.25	23.2115424203424	23.2118253294148	0.0002829090724
6	0.30	27.3235128652981	27.3238973047029	0.0003844394048
7	0.35	32.4819128264705	32.4824360407517	0.0005232142812
8	0.40	38.9761610052408	38.9768757474937	0.0007147422529
9	0.45	47.1788782536437	47.1798613256007	0.0009830719570
10	0.50	57.5702295395952	57.5715961933444	0.0013666537492
11	0.55	70.7694526401717	70.7713812430724	0.0019286029007
12	0.60	87.5757073128212	87.5784811589605	0.0027738461393
13	0.65	109.0210141703044	109.0250915942196	0.0040774239152
14	0.70	136.4388783251231	136.4450087307618	0.0061304056387
15	0.75	171.5532656196127	171.5626788021488	0.0094131825361
16	0.80	216.5939930023279	216.6087037708262	0.0147107684983
17	0.85	274.4464056678120	274.4696974832635	0.0232918154516
18	0.90	348.8455672041557	348.8827504837690	0.0371832796133
19	0.95	444.6282481406327	444.6878355186185	0.0595873779858
20	1.00	568.0599749599790	568.1554834546089	0.0955084946298

Penyelesaian dengan menggunakan metode Milne-Simpson, telah diketahui bahwa sistem persamaan diferensial orde-1 pada persamaan (17), yaitu :

$$f(t, \bar{Y}) = \begin{bmatrix} Y_1, Y_2, \dots, Y_7, (18Y_7 - 126Y_6 + 433Y_5 - 752Y_4 + 576Y_3 - 61Y_2) \\ -91Y_1 - 10Y_0 \end{bmatrix}$$

Ditentukan solusi masalah nilai awal dengan nilai awal $t_0 = 0$ dan nilai-nilai awal yaitu :

$$\begin{aligned} Y_0(0) = Y_{0(0)} &= 12 & Y_1(0) = Y_{1(0)} &= 25.4821 \\ Y_2(0) = Y_{2(0)} &= 101.5011 & Y_3(0) = Y_{3(0)} &= 458.0442 \\ Y_4(0) = Y_{4(0)} &= 2224.6790 & Y_5(0) = Y_{5(0)} &= 11170.3964 \\ Y_6(0) = Y_{6(0)} &= 57034.0952 & Y_7(0) = Y_{7(0)} &= 293878.8247 \end{aligned}$$

Serta $h = 0.05$ dan $t_{(i+1)} = t_{(i)} + h$. Akan dicari solusi masalah nilai awal di atas pada interval $[0,1]$ untuk memperoleh solusi awal $\hat{Y}_{(0)}, \hat{Y}_{(1)}, \hat{Y}_{(2)}$, dan $\hat{Y}_{(3)}$ menggunakan metode Runge-Kutta orde-4. Setelah diperoleh solusi awal menggunakan Runge-Kutta orde-4, tentukan nilai-nilai $f_{(i)}, f_{(i-1)}$, dan $f_{(i-2)}$ kemudian hasil tersebut disubstitusikan ke metode Milne :

$$\begin{aligned} \bar{P}_{(i+1)} &= \hat{Y}_{(i-3)} + \frac{4h}{3}(2f_{(i-2)} - f_{(i-1)} + 2f_{(i)}) \\ \bar{P}_{(4)} &= \hat{Y}_{(0)} + \frac{4h}{3}(2f_{(1)} - f_{(2)} + 2f_{(3)}) \end{aligned}$$

Setelah didapatkan nilai prediksi, hitung $f_{(i+1)} = f(t_{(i+1)}, \bar{P}_{(i+1)})$ kemudian disubstitusikan ke persamaan Simpson untuk memperoleh nilai koreksi :

$$\begin{aligned} \hat{Y}_{(i+1)} &= \hat{Y}_{(i-1)} + \frac{h}{3}(f_{(i-1)} + 4f_{(i)} + 9f_{(i+1)}) \\ \hat{Y}_{(4)} &= \hat{Y}_{(2)} + \frac{h}{3}(f_{(2)} + 4f_{(3)} + 9f_{(4)}) \end{aligned}$$

Kemudian dihitung galat mutlak :

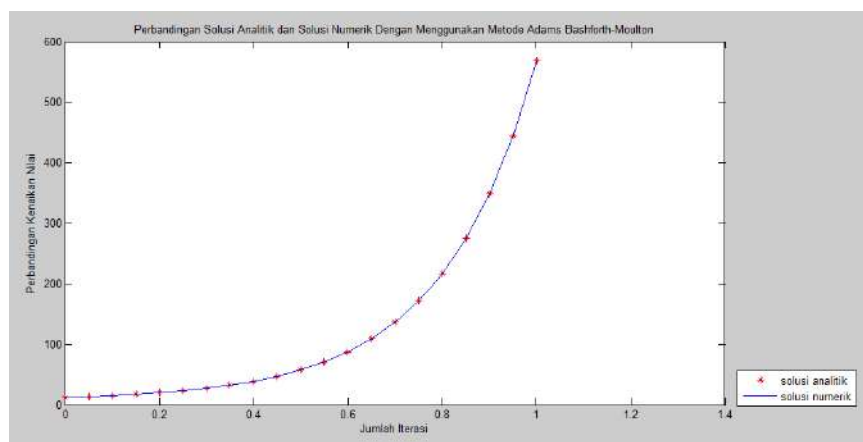
$$E_{Y_{(i)}} = |Y(t_{(i)}) - \hat{Y}_{0(i)}|$$

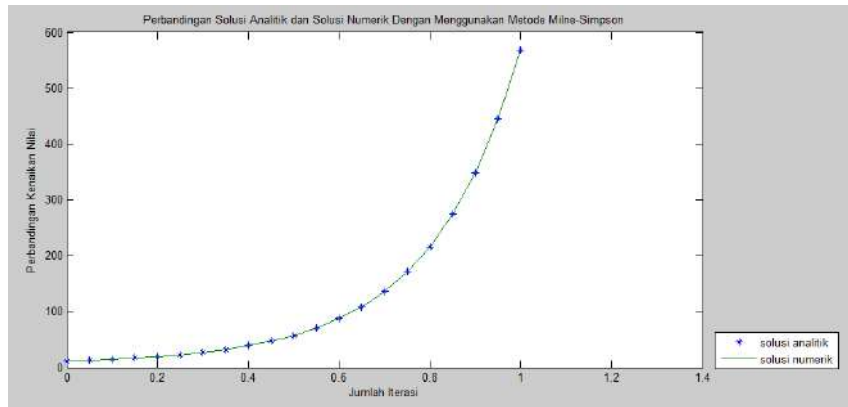
didapatkan solusi analitik dan solusi numerik dari contoh kasus 2 persamaan diferensial Euler orde-8 sebagai berikut :

Tabel 4.4. Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 2 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Milne-Simpson

i	$t_{(i)}$	Solusi Numerik MS ($\hat{Y}_{0(i)}$)	Solusi Analitik ($Y(t_{(i)})$)	Galat (Error) ($E_{Y(i)}$)
0	0	12.0000000000000	12.0000000000000	0.0000000000000
1	0.05	13.4111033059896	13.4111337830808	0.0000304770912
2	0.10	15.1422635610499	15.1423424266375	0.0000788655877
3	0.15	17.2767083056203	17.2768613001565	0.0001529945362
4	0.20	19.9209745040050	19.9211942018391	0.0002196978342
5	0.25	23.2114728188767	23.2118253294148	0.0003525105381
6	0.30	27.3233836126020	27.3238973047029	0.0005136921008
7	0.35	32.4816679091208	32.4824360407517	0.0007681316309
8	0.40	38.9757680798519	38.9768757474937	0.0011076676418
9	0.45	47.1782508172365	47.1798613256007	0.0016105083641
10	0.50	57.5692824286226	57.5715961933444	0.0023137647218
11	0.55	70.7680361636050	70.7713812430724	0.0033450794673
12	0.60	87.5736414406642	87.5784811589605	0.0048397182963
13	0.65	109.0180267070998	109.0250915942196	0.0070648871198
14	0.70	136.4346159593168	136.4450087307618	0.0103927714451
15	0.75	171.5472258240373	171.5626788021488	0.0154529781114
16	0.80	216.5855049800691	216.6087037708262	0.0231987907572
17	0.85	274.4345407001467	274.4696974832635	0.0351567831168
18	0.90	348.8290742070952	348.8827504837690	0.0536762766737
19	0.95	444.6054175209614	444.6878355186185	0.0824179976572
20	1.00	568.0284993485176	568.1554834546089	0.1269841060913

Berikut adalah gambar grafik perbandingan antara solusi analitik dan solusi numerik untuk contoh kasus 2 persamaan diferensial Euler orde-8 dengan menggunakan metode Adams Bashforth-Moulton dan metode Milne-Simpson :





Gambar 4.2. Grafik Perbandingan Antara Solusi Analitik dan Solusi Numerik dari Contoh Kasus 2 Persamaan Diferensial Euler Orde-8 dengan Menggunakan Metode Adams Bashforth-Moulton dan Metode Milne-Simpson

PEMBAHASAN

Metode Adams Bashforth-Moulton dan metode Milne-Simpson dapat digunakan untuk penyelesaian solusi analitik dan solusi numerik persamaan diferensial Euler orde-8 dengan akar-akar karakteristiknya riil berbeda. Terlihat pada Table 4.1, Tabel 4.2, Tabel 4.3 serta Tabel 4.4 bahwa nilai solusi numerik dari hasil metode Adams Bashforth-Moulton dan metode Milne-Simpson semakin mendekati nilai solusi analitiknya, serta dari Gambar 4.1 dan Gambar 4.2 bahwa gambar grafik perbandingan antara solusi analitik dan solusi numerik kedua metode terlihat berhimpitan.

Diketahui dari Table 4.5 bahwa metode Adams Bashforth-Moulton lebih mendekati nilai solusi analitik dibandingkan metode Milne-Simpson serta dilihat dari jumlah perbandingan galat yang didapatkan bahwa untuk kedua kasus galat metode Adams Bashforth-Moulton lebih kecil dibandingkan metode Milne-Simpson, sehingga dapat disimpulkan bahwa metode Adams Bashforth-Moulton lebih akurat dibandingkan metode Milne-Simpson untuk kedua contoh kasus tersebut. Berikut adalah perbandingan galat yang didapatkan dari kedua contoh kasus :

Tabel 4.5. Jumlah Perbandingan Galat (Error) Antara Solusi Analitik dan Solusi Numerik dengan Metode Adams Bashforth-Moulton dan Metode Milne-Simpson Untuk Contoh Kasus 1 dan Contoh Kasus 2

ABM CK 1	MS CK 1	ABM CK 2	MS CK 2
0.127073791	0.248146933	0.26	0.369678

Keterangan :

ABM CK1 = Jumlah galat metode Adams Bashforth-Moulton dari contoh kasus 1

MS CK 1 = Jumlah galat metode Milne-Simpson dari contoh kasus 1

ABM CK 2 = Jumlah galat metode Adams Bashforth-Moulton dari contoh kasus 1

MS CK 2 = Jumlah galat metode Milne-Simpson dari contoh kasus 2

Untuk efisiensi dapat dibandingkan dengan menghitung lama waktu program berjalan hingga diperoleh solusi numerik. Berikut adalah tabel perbandingan perhitungan dari kedua metode tersebut :

Tabel 4.6. Perbandingan Lama Waktu Proses Iterasi dengan Menggunakan Metode Adams Bashforth-Moulton dan Metode Milne-Simpson

ABM CK1	MS CK1	ABM CK 2	MS CK 2
0.335147 s	0.020519 s	0.307962 s	0.020490 s

Terlihat bahwa untuk kedua kasus tersebut metode Milne-Simpson dalam melakukan iterasi lebih cepat dibandingkan dengan metode Adams Bashforth-Moulton. Sehingga dapat disimpulkan bahwa metode Milne-Simpson lebih efisien daripada metode Adams Bashforth-Moulton.

5. SIMPULAN

Adapun kesimpulan yang didapatkan dari penelitian ini adalah sebagai berikut :

1. Persamaan diferensial Euler orde-8 dapat ditransformasikan kedalam bentuk sistem persamaan diferensial orde-1 dengan menggunakan transformasi $x = e^t, x > 0$.
2. Hasil perbandingan antara metode Adams Bashforth-Moulton dan metode Milne-Simpson dapat disimpulkan bahwa :
 - a. Metode Adams Bashforth-Moulton dan metode Milne-Simpson dapat digunakan dalam menyelesaikan persamaan diferensial Euler orde-8.
 - b. Solusi numerik dari kedua metode tersebut mendekati solusi analitik dari kedua contoh kasus yang telah diselesaikan.
 - c. Metode Milne-Simpson lebih efisien dibandingkan dengan metode Adams Bashforth-Moulton.
 - d. Metode Adams Bashforth-Moulton lebih baik dalam menyelesaikan persamaan diferensial Euler orde- 8 karena lebih akurat terlihat dari jumlah perbandingan galatnya.

KEPUSTAAKAN

- [1] Nugroho, Didit Budi. 2011. *Persamaan Diferensial Biasa dan Aplikasinya : Penyelesaian Manual dan Menggunakan Maple*. Graha Ilmu, Yogyakarta.
- [2] Prayudi. 2006. *Matematika Teknik*. Graha Ilmu, Yogyakarta.
- [3] Munir, Rinaldi. 2008. *Metode Numerik : Revisi Kedua*. Penerbit Informatika, Bandung.
- [4] Mathews, John H., and Kurtis D. Fink. 2004. *Numerical Methods Using MATLAB*. Pearson Education, Inc, United States of America.

PENGEMBANGAN EKOWISATA DENGAN MEMANFAATKAN MEDIA SOSIAL UNTUK MENGUKUR SELERA CALON KONSUMEN

Gustafika Maulana¹⁾, Gunardi Djoko Winarno²⁾, & Samsul Bakri³⁾

¹⁾ Mahasiswa PS Kehutanan Fakultas Pertanian, ²⁾ Dosen PS Kehutanan Fakultas Pertanian,

³⁾ Dosen PS Kehutanan Fakultas Pertanian, Universitas Lampung

Jl. Soemantri Brojonegoro No.1 Bandar Lampung

Email : gustafika11@gmail.com Mobile: 081326657857

ABSTRAK

Pengembangan ekowisata dengan memanfaatkan media sosial dapat mengukur selera calon konsumen. Perlu dilakukan eksplorasi untuk mencari focal point bagi percepatan pengembangan wisata. Penelitian ini bertujuan untuk mengembangkan potensi wisata terdapat di Desa Braja Harjosari, pengembangan rancangan integrasi potensi gajah liar dengan potensi ekowisata untuk menunjang desa wisata Braja Harjosari, mengetahui selera calon konsumen terhadap potensi objek wisata yang akan dikembangkan. Penelitian ini menggunakan metode observasi, kuisioner online dan wawancara. Data observasi yang telah diperoleh dianalisis dengan menggunakan pendekatan Recreation Opportunity Spectrum (ROS) dan juga dicatat variabel demografi, pendidikan dan waktu luang tiap responden. Untuk menguji pengaruh 3 kelompok variabel responden ini digunakan model log linier pada taraf nyata 15%. Hasil penelitian menunjukkan bahwa ekowisata yang semula terdiri dari 7 obyek ekowisata dapat dikembangkan menjadi 16 obyek ekowisata, aktivitas gajah liar yang selalu masuk kedalam wilayah desa yang berbatasan dengan kawasan Taman Nasional Way Kambas dapat dijadikan obyek ekowisata yang terintegrasi dengan potensi sumber daya wisata pedesaan dan potensi-potensi yang ada di Desa Braja Harjosari dari analisis hasil output Minitab 16 berbasis demografi, pendidikan dan rutinitas wisata dinilai sangat baik.

Kata kunci: gajah liar, media sosial, potensi ekowisata, Taman Nasional Way Kambas

1. PENDAHULUAN

Wisata satwa liar bukan hanya berada di dalam Taman Nasional atau kawasan dilindungi lainnya. Tetapi juga dapat di luar kawasan dilindungi seperti daerah penyangga. Kondisi ini terjadi karena untuk mengatasi permasalahan satwa liar yang sering kali keluar dari kawasan hutan menuju desa-desa di sekitar zona penyangga dan sering terjadi konflik karena satwa ingin mencari sumber pakan.

Salah satu desa penyangga yang berada diperbatasan kawasan Taman Nasional Way Kambas adalah Desa Braja Harjosari. Desa Braja Harjosari merupakan sebuah desa binaan dibawah program *Tropical Forest Conservation Action-Sumatera* (TFCA-Sumatera) Konsorsium AleRT-Universitas Lampung. Desa ini memiliki berbagai potensi wisata seperti budidaya jamur tiram, perkebunan jambu kristal (*Psidium guajava*), obyek wisata dermaga, pasar tradisional braja slebah, sentra tungku api, tanam padi tradisional dan tarian khas bali.

Pengembangan potensi sumber daya wisata dan pemanfaatan gajah untuk ekowisata kedepan diharapkan bisa mengatasi konflik yang sering terjadi dan bermanfaat bagi kedua belah pihak. Pengembangan ekowisata dengan memanfaatkan media sosial dapat juga mengukur selera calon konsumen atau calon wisatawan. Media sosial adalah sebuah istilah yang menggambarkan bermacam-macam teknologi yang digunakan untuk mengikat orang-orang ke dalam suatu kolaborasi, saling bertukar informasi dan berinteraksi melalui isi pesan yang berbasis web [1]. Selera konsumen merupakan variabel yang mempengaruhi besar kecilnya permintaan terhadap suatu barang. [2]. Itulah pentingnya penelitian ini dilakukan dengan melibatkan masyarakat lokal. Potensi ekowisata tidak hanya didukung satwa liar tetapi juga dengan potensi desa penyangga.

Tujuan dari penelitian ini untuk mengembangkan potensi wisata termasuk atraksi gajah liar yang terdapat di desa wisata Braja Harjosari sebagai objek wisata serta pengembangan ekowisata dilakukan secara terintegrasi antara potensi gajah liar dengan potensi desa untuk menunjang ekowisata desa wisata Braja Harjosari dan mengetahui selera calon konsumen terhadap potensi objek wisata yang akan dikembangkan di Desa Braja Harjosari.

2. LANDASAN TEORI

2.1 Definisi desa

Desa adalah sebagai kesatuan masyarakat hukum yang mempunyai susunan asli berdasarkan hak asal-usul yang bersifat istimewa. Landasan pemikiran dalam mengenai Pemerintahan Desa adalah keanekaragaman, partisipasi, otonomi asli, demokratisasi dan pemberdayaan masyarakat [3].

[4], desa adalah sekumpulan yang hidup bersama atau suatu wilayah, yang memiliki suatu serangkaian peraturan-peraturan yang ditetapkan sendiri, serta berada di wilayah pimpinan yang dipilih dan ditetapkan sendiri.

Menurut [5], desa sebagai salah satu jenis persekutuan hukum teritorial, persekutuan hukum teritorial adalah kelompok dimana anggota-anggotanya merasa terikat satu dengan yang lainnya karena merasa dilahirkan dan menjalani kehidupan di tempat atau wilayah yang sama.

2.2 Definisi Desa Wisata

Desa wisata yakni adat istiadat masyarakat setempat sebagai daya tarik wisata seperti: kehidupan sehari-hari, upacara adat, rumah adat, budaya dan kesenian asli daerah, makanan minuman tradisional, kekayaan alam, dan lain-lain [6].

Desa wisata adalah keterlibatan masyarakat desa dalam setiap aspek wisata yang ada di desa tersebut. Masyarakat terlibat langsung dalam kegiatan pariwisata dalam bentuk pemberian jasa dan pelayanan yang hasilnya dapat meningkatkan pendapatan masyarakat diluar aktifitas mereka sehari-hari [7].

2.3 Definisi Wisata dan Ekowisata

Ekowisata sebagai kegiatan perjalanan ke daerah- daerah yang masih alami dengan kepedulian terhadap lingkungan hidup dan masyarakat sekitar [8].

Ekowisata adalah kegiatan perjalanan wisata yang dikemas secara profesional, terlatih, dan memuat unsur pendidikan, sebagai suatu sektor usaha ekonomi, yang mempertimbangkan warisan budaya, partisipasi dan kesejahteraan penduduk lokal serta upaya-upaya konservasi sumberdaya alam dan Menurut [9].

Konsep dasar ekowisata menjadi lima prinsip inti. Mereka termasuk yang berbasis alam, berkelanjutan secara ekologis, lingkungan edukatif, dan lokal wisatawan bermanfaat dan menghasilkan kepuasan [10].

Tujuan pengembangan ekowisata menurut [11] sebagai berikut:

1. Membangun kesadaran lingkungan dan budaya di daerah tujuan wisata baik bagi wisatawan, masyarakat setempat maupun penentu kebijakan di bidang kebudayaan dan kepariwisataan.
2. Mengurangi dampak negatif berupa kerusakan atau pencemaran lingkungan dan budaya lokal akibat kegiatan ekowisata.
3. Memberikan keuntungan ekonomi secara langsung bagi konservasi melalui kontribusi atau pengeluaran wisatawan.
4. Mengembangkan ekonomi masyarakat dan pemberdayaan masyarakat setempat dengan menciptakan produk wisata alternatif yang mengedepankan nilai-nilai dan keunikan lokal.

Jasa ekowisata menurut [12] meliputi enam jenis: (i) pemandangan dan atraksi lingkungan dan budaya, misalnya titik pengamatan atau sajian budaya; (ii) manfaat lansekap, misalnya jalur pendakian atau trekking; (iii) akomodasi, misalnya pondok wisata, restoran; (iv) peralatan dan perlengkapan, misalnya sewa alat penyelam dan camping; (v) pendidikan dan ketrampilan, dan (vi) penghargaan, yakni prestasi di dalam upaya konservasi.

2.4 Pemanfaatan Gajah untuk Wisata

Status Tangkahan sebagai Kawasan Ekowisata Tangkahan (KET) dengan gajah sumatera sebagai ikon yang paling melekat dengannya, maka sudah seharusnya dibarengi dengan persepsi ataupun wawasan masyarakat yang baik dalam konservasi gajah sumatera terlebih lagi dengan kegiatan-kegiatan penyuluhan konservasi yang pernah dilakukan oleh pihak manajemen CRU kepada masyarakat sekitar KET. Karena, gajah sumatera yang menjadikan

Tangkahan sebagai obyek wisata yang menarik perhatian baik dari kalangan wisatawan asing maupun dalam negeri [13].

Salah satu aktivitas TNWK antara lain melakukan pembinaan anak gajah, yaitu anak-anak gajah yang berasal dari gajah domestikasi tetap jinak sedangkan yang berasal dari gajah liar turut menjadi jinak. Gajah liar yang menjadi pelaku pengrusakan lahan pertanian masyarakat dapat diminimalisir jumlahnya dengan gajah jinak dan gajah jinak juga dapat dimanfaatkan sebagai penunjang ekowisata di PKG maupun ERU [14].

2.5 Recreation Opportunity Spectrum (ROS)

Recreation Opportunity Spectrum (ROS) adalah sebuah *planning framework* yang diterapkan pada *landscape* dan *seascape* dengan tujuan menangani terjadinya *landscape conflict* melalui inventarisasi, perencanaan dan manajemen. Tujuan dari penerapan ROS adalah untuk mendapatkan keseimbangan antara pemanfaatan dan

konservasi. ROS mendukung zonasi dan pembangunan *recreation experience* dimana wilayah diklasifikasikan dan dibagi berdasarkan kondisi lingkungan dan aktivitas rekreasi [15].

2.6 Media Sosial

Penyebaran informasi melalui media sosial dapat memberikan keuntungan maupun kerugian tergantung dari cara penggunaannya. Dengan menggunakan media sosial secara tepat, berpotensi dalam meningkatkan minat wisata bagi para pengguna internet yang membaca dan mengikuti media sosial tersebut [16].

3. METODOLOGI PENELITIAN

Penelitian dilakukan pada bulan Januari 2017-Maret 2017 di Desa Wisata Braja Harjosari Kecamatan Braja Slebah Kabupaten Lampung Timur. Pengumpulan data dalam penelitian ini dengan metode *Recreation Opportunity Spectrum* (ROS), kuisioner *online*, observasi dan wawancara. *Recreation Opportunity Spectrum* (ROS) adalah sebuah *planning framework* yang diterapkan pada *landscape* dan *seascape* dengan tujuan menangani terjadinya *landuse conflict* melalui inventarisasi, perencanaan dan manajemen agar mendapatkan keseimbangan antara pemanfaatan dan konservasi. ROS mendukung zonasi dan pembangunan *recreation experience* dimana wilayah diklasifikasikan dan dibagi berdasarkan kondisi lingkungan dan aktivitas rekreasi [17].

Observasi terhadap potensi sumber daya wisata yang tersebar di desa wisata Braja Harjosari. Wawancara dilakukan untuk mengkaji karakteristik obyek ekowisata di wilayah pedesaan dan area yang berbatasan dengan kawasan Taman Nasional Way Kambas. Analisis data disajikan secara deskriptif dan didukung dengan gambar serta tabel.

Kuisioner juga disiapkan untuk mengukur ketertarikan calon wisatawan terhadap potensi wisata yang akan dikembangkan di Desa Braja Harjosari. Responden terdiri dari calon wisatawan domestik maupun wisatawan asing. Responden berjumlah minimal 30 orang. Pengumpulan data dilakukan melalui media sosial instagram. Menurut [18] instagram adalah media sosial yang digunakan untuk membagi foto sehingga membuat banyak pengguna yang terjun ke bisnis *online* turut mempromosikan produknya dan dimanfaatkan sebagai media pemasaran langsung. Di dalam instagram tersedia *link* kuisioner *online* yang akan diisi oleh responden.

4. HASIL DAN PEMBAHASAN

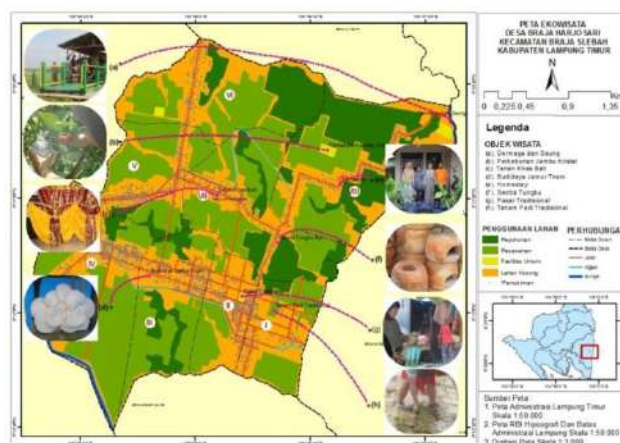
4.1 Sebaran Obyek Ekowisata Desa Braja Harjosari

Desa wisata merupakan bentuk pariwisata, yang sekelompok kecil wisatawan tinggal di dalam atau di dekat kehidupan tradisional atau di desa-desa terpencil dan mempelajari kehidupan desa dan lingkungan setempat [19]. Menurut [20], desa wisata merupakan suatu bentuk integrasi antara atraksi, akomodasi dan fasilitas pendukung yang disajikan dalam suatu struktur kehidupan masyarakat yang menyatu dengan tata cara dan tradisi yang berlaku. Desa wisata biasanya memiliki kecenderungan kawasan pedesaan yang memiliki daya tarik sebagai tujuan wisata

Menurut [21], wisatawan tidak hanya datang untuk melakukan aktifitas rekreasi tetapi juga mencari lokasi rekreasi yang spesifik disamping menikmati pengalaman rekreasi tertentu sekaligus mendapatkan manfaat (*benefit*) baik secara individu, komunitas, ekonomi dan lingkungan. Sehingga dikatakan empat komponen yang membentuk peluang rekreasi (*recreation opportunity*) adalah aktifitas, lokasi, pengalaman dan manfaat. Desa Braja Harjosari sebagai desa wisata sudah memiliki obyek-obyek ekowisata yang menjadi daya tarik wisata. Obyek wisata yang ada saat ini di Desa Braja Harjosari memiliki sebanyak 7 lokasi yang tersebar di berbagai dusun (Gambar 1) yaitu wisatawan bisa belajar cara budidaya jamur tiram yang dikelola secara tradisional dan mengolah jamur tiram yang sudah dipanen. Lalu memetik, mengkonsumsi dan mengemas jambu kristal yang siap panen langsung dari kebunnya. Wisatawan juga bisa membeli alat-alat dan jajanan tradisional dan melihat proses jual beli di pasar tradisional Braja Slebah. Membeli dan belajar membuat tungku api secara tradisional di sentra tungku api. Menanam dan membajak padi secara tradisional. Menikmati dan belajar tarian khas bali. Obyek yang terakhir yaitu wisatawan melihat panorama kawasan Taman Nasional Way Kambas di dermaga. Sebaran potensi obyek ekowisata yang berada di Desa Braja Harjosari disajikan pada Tabel 1.

Tabel 1. Deskripsi sebaran Obyek Wisata Desa Braja Harjosari

No	Nama Obyek Wisata	Potensi	Waktu dan Lokasi
1	Budidaya Jamur Tiram	Jamur tiram ini dibudidayakan dengan cara tradisional. Wisatawan akan praktik membuat media tanam jamur tiram. Selain itu wisatawan juga mengolah jamur tiram crispy dan langsung bisa mencicipi hasil olahannya sendiri.	Pagi Dusun IV
2	Obyek Wisata Perkebunan Jambu Kristal	Perkebunan jambu kristal ini memiliki luas wilayah 3 hektar. Wisata yang ditawarkan kepada para pengunjung yaitu pengunjung dapat memetik langsung buah jambu kristal dari pohonnya. Hanya buah yang terlihat besar dan matang yang sudah bisa dipetik. Selain memetik buah, wisatawan juga diajari cara mengemas jambu kristal dengan baik dan benar.	Siang Dusun VIII
3	Obyek Wisata Dermaga	Lokasi obyek wisata ini terletak di dekat sungai yang merupakan perbatasan antara Desa Braja Harjosari dengan kawasan Taman Nasional Way Kambas. Wisatawan juga dapat melihat panorama kawasan dengan keindahan burung-burung dan suasana alam yang ada disekitarnya.	Malam Dusun VIII
4	Obyek Wisata Pasar Tradisional Braja Slebah	Pasar tradisional Braja Slebah ini merupakan pasar kecamatan Braja Slebah. Kecamatan Braja Slebah memiliki 7 desa yang salah satunya adalah Desa Braja Harjosari. Wisatawan akan disuguhkan dengan proses jual beli yang ada di pasar tradisional ini. Pasar tradisional ini menjual berbagai alat tradisional, jajanan tradisional dll.	Pagi Dusun II
5	Obyek Wisata Sentra Tungku Api	Wisatawan dapat melihat cara membuat tungku api di obyek wisata ini. Selain melihat, wisatawan juga akan praktik langsung dalam membuat tungku api ini seperti pada obyek wisata jamu tiram. Tungku api dapat dibawa pulang dengan membayar 20 ribu per tungku api.	Siang Dusun VII
6	Obyek Wisata Tanam Padi Tradisional	Wisatawan akan diajak untuk membajak sawah secara tradisional. Wisatawan juga akan diajak untuk nandur disawah, dengan dipandu oleh petani lokal agar proses nandur baik dan benar. Sensasi yang dirasakan akan berbeda, karena semuanya dibuat konsep "back to nature".	Siang Dusun I
7	Obyek Wisata Tarian Khas Bali	Wisatawan bisa menikmati tarian asli bali seperti tari kencak, tari ngibing dll. Selain itu wisatawan bisa belajar bermain gamelan maupun belajar menari.	Malam Dusun VII



Gambar 1. Peta Sebaran Obyek Wisata saat ini Desa Braja Harjosari

4.2 Sebaran Obyek Ekowisata Desa Braja Harjosari

Berdasarkan metode *Recreation Opportunity Spectrum* (ROS), diperoleh berbagai spektrum peluang rekreasi dengan karakteristiknya masing-masing. Menurut [22] pengembangan merupakan proses atau langkah untuk mengembangkan suatu produk baru untuk menyempurnakan produk yang sudah ada yang bisa dipertanggung jawabkan.

Ada berbagai obyek ekowisata wisata di Desa Braja Harjosari yang bisa dikembangkan agar meningkatkan daya tarik menjadi lebih baik. Obyek wisata merupakan salah satu komponen yang penting dalam industri pariwisata dan salah satu alasan pengunjung melakukan perjalanan (*something to see*). Di luar negeri obyek wisata disebut *tourist attraction* (atraksi wisata), sedangkan di Indonesia lebih dikenal dengan objek wisata. Obyek wisata adalah tempat atau keadaan alam yang memiliki sumber daya wisata yang dibangun dan dikembangkan sehingga mempunyai daya tarik dan diusahakan sebagai tempat yang dikunjungi wisatawan [23]. Terdapat penambahan sebanyak 9 obyek wisata yang dapat dikembangkan yaitu obyek wisata gajah liar (*elephant watching*), obyek wisata sungai, obyek wisata air, menunggang kuda, menunggang gajah jinak, sentra kuliner, sentra souvenir, tarian kuda lumping dan kebun nangkada (nangka cempeda) adapun deskripsinya disajikan pada Tabel 2.

Tabel 2. Deskripsi Pengembangan Obyek Wisata Desa Braja Harjosari

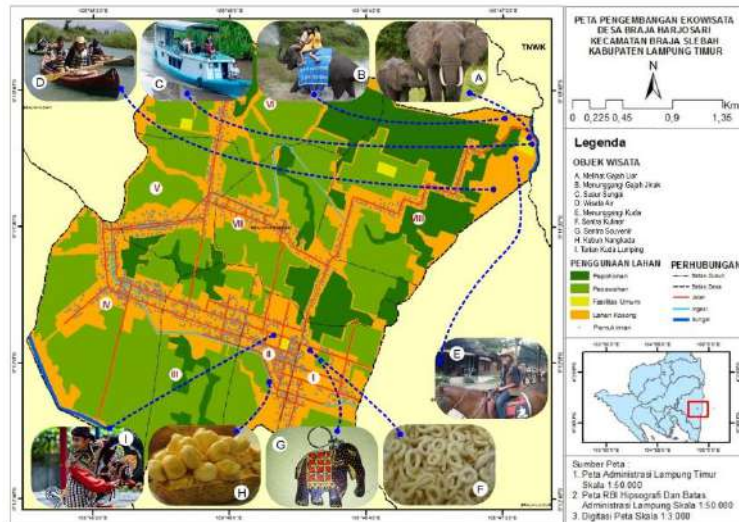
No	Nama Obyek Wisata	Potensi	Waktu dan Lokasi
1	Obyek Wisata Melihat Gajah Liar (Elephant Watching)	Lokasi obyek wisata ini berdekatan dengan lokasi dermaga. Dilokasi ini berbatasan dengan kawasan Taman Nasional Way Kambas, dimana gajah liar sering terlihat. Wisatawan bisa melihat langsung gajah liar mulai petang hingga dini hari. Lokasi terjadinya konflik gajah dengan masyarakat. Gajah-gajah mencoba masuk ke desa untuk mencari pakan melalui jalur dilokasi ini.	Malam Dusun VIII
2	Obyek Wisata Sungai	Pemanfaatan potensi sungai yang berada di perbatasan antara kawasan TNWK dengan Desa Braja Harjosari, banyak aktivitas yang dapat dilakukan untuk meningkatkan daya tarik wisata. Kegiatan wisata yang bisa dilakukan seperti susur sungai dengan perahu, mandi di sungai dan memancing ikan.	Siang Dusun VIII
3	Obyek Wisata Air	Pemanfaatan area padang rumput dan view yang terbuka Desa Braja Harjosari yang berada di dekat dermaga. Akan didesain sungai buatan di lahan tersebut yang dimanfaatkan untuk berbagai macam wisata air seperti bermain bebek air, bermain sampan dan bermain ban air.	Malam Dusun VIII
4	Obyek Wisata Menunggang Kuda	Pemanfaatan area padang rumput dan view yang terbuka Desa Braja Harjosari juga yang berada di dekat dermaga. Akan dimanfaatkan untuk area menunggang kuda dan berfoto dengan kuda. Kita juga dapat menikmati keindahan alam karena berbatasan langsung dengan kawasan TNWK. Kuda yang dipersiapkan adalah kuda terlatih.	Pagi dan Sore Dusun VIII
5	Obyek Wisata Menunggang Gajah Jinak	Pemanfaatan area padang rumput dan view yang terbuka Desa Braja Harjosari yang berada di dekat dermaga. Akan dimanfaatkan untuk area menunggang gajah jinak dan berfoto dengan gajah jinak. Kita juga dapat menikmati keindahan alam karena berbatasan langsung dengan kawasan TNWK perlu proses perizinan agar bisa membawa gajah jinak keluar dari kawasan.	Pagi dan Sore Dusun VIII
6	Obyek Wisata Sentra Kuliner	Masyarakat setempat mengolah produk makanan yang bisa dibawa sebagai oleh-oleh wisatawan seperti keripik, kelanting, tiwul dan mokaf.	Siang Dusun I
7	Obyek Wisata Sentra Souvenir	Masyarakat setempat mengolah kerajinan tangan menjadi oleh-oleh wisatawan seperti gantungan kunci, patung, tas, kaos, boneka, stiker dll	Siang Dusun I
8	Obyek Wisata Tarian Kuda Lumpung	Masyarakat yang didominasi oleh suku Jawa seharusnya juga bisa menampilkan tarian khas suku mereka seperti tarian kuda lumping dll	Siang Dusun VIII
9	Obyek Wisata Kebun Nangkada (Nangka Cempeda)	Perkebunan Nangkada ini memiliki luas wilayah 2 hektar. Kebun Nangkada ini milik perseorangan yang bernama bapak Kembang. Wisata yang ditawarkan kepada para pengunjung yaitu pengunjung dapat memetik langsung buah Nangkada dari pohonnya.	Siang Dusun II

Menurut [24] destinasi pariwisata dalam pengertian sistem kepariwisataan adalah kawasan geografis yang berada dalam 1 (satu) atau lebih wilayah administrasi yang di dalamnya terdapat daya tarik wisata, fasilitas umum, fasilitas pariwisata, aksesibilitas, serta masyarakat yang saling terkait dan melengkapi terwujudnya kepariwisataan. Pengembangan obyek ekowisata di Desa Braja Harjosari tersebar di berbagai dusun seperti dusun I, dusun II dan dusun VIII. Pengembangan ini juga memanfaatkan dalam rangka pemanfaatan lahan-lahan pedesaan yang berupa padang rumput dengan view yang indah selengkapnya dapat dilihat pada Gambar 2.

Menurut [25], suatu daerah untuk menjadi daerah tujuan wisata (DTW) yang baik, harus mengembangkan tiga hal agar daerah tersebut menarik untuk dikunjungi, yakni adanya sesuatu yang dapat dilihat (*something to see*), adanya sesuatu yang dapat dibeli (*something to buy*), adanya sesuatu yang dapat dilakukan (*something to do*).

Potensi sumber daya wisata yang akan dikembangkan memiliki 3 hal yang membuat wisatawan menjadi menarik Tabel 2.

Potensi wisata adalah berbagai sumber daya yang terdapat di sebuah daerah tertentu yang bisa dikembangkan menjadi daya tarik wisata. Dengan kata lain, potensi wisata adalah berbagai sumber daya yang dimiliki oleh suatu tempat dan dapat dikembangkan menjadi suatu atraksi wisata (*tourist attraction*) yang dimanfaatkan untuk kepentingan ekonomi dengan tetap memperhatikan aspek-aspek lainnya [26].



4.3 Analisa Data Hasil Pemodelan

4.3.1 Potensi Daya Tarik Obyek Wisata Melihat Gajah Liar (*Elephant Watching*)

Hasil optimasi parameter model potensi obyek wisata gajah liar (*elephant watching*) berbasis demografi, pendidikan dan rutinitas berwisata pada Tabel 3.

Tabel 3. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Gajah Liar (*Elephant Watching*) Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		α_0	-36,8939	0,993	
Demografi					
Umur					
- Dewasa Awal ($I=DAW$)	DAW	α_1	17,1048	0,999	26823382,98
-Dewasa Akhir ($I=DAR$)	DAR	α_2	3,85732	1,000	47,34
-Lansia Awal ($I=LAW$)	LAW	α_3	20,6599	0,999	9,38625E+08
Kelamin ($I=Lk$)	KLM	α_4	-0,116445	0,937	0,89
Alamat					
-Kota dalam Provinsi ($I=KTDP$)	KTDP	α_5	3,15247	0,069*	23,39
-Kabupaten luar Provinsi ($I=KBLP$)	KBLP	α_6	0,560757	0,845	1,75
-Kota luar Provinsi ($I=KTLP$)	KTLP	α_7	4,93570	0,059*	139,17
-Luar Negeri ($I=LN$)	LN	α_8	5,68253	0,099*	293,69
Pekerjaan					
-Pegawai Negeri Sipil ($I=PNS$)	PNS	α_9	17,8617	0,999	57180605,60
-Non Pegawai Negeri Sipil ($I=NPNS$)	NPNS	α_{10}	2,13719	0,245	8,48
Pendidikan					
-Diploma ($I=DP$)	DP	α_{11}	16,8934	0,996	21711609,90
-Sarjana ($I=SJ$)	SJ	α_{12}	3,20865	0,204	24,75
Penghasilan					
-Rendah ($I=RDH$)	RDH	α_{13}	25,6777	0,995	1,41799E+11
-Sedang ($I=SDG$)	SDG	α_{14}	29,3980	0,995	5,85289E+12
-Tinggi ($I=TGI$)	TGI	α_{15}	30,5697	0,994	1,88911E+13
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	α_{16}	1,83798	0,502	6,28
-Lebih dari 5 kali ($I=LBDL$)	LBDL	α_{17}	1,08101	0,584	2,95

Log-Likelihood = -11,351

Test that all slopes are zero: G = 102,339, DF = 17, P-Value = 0,0004

Keterangan: -*nyata pada taraf 15% ; ** nyata pada taraf 5%

- P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Dari hasil output Minitab 16 pada Tabel 3 memperlihatkan bahwa potensi obyek wisata gajah liar (*elephant watching*) dapat dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y1] memberikan hasil P-value = 0,0004 (0,04%). Responden yang beralamat di luar negeri dengan $\alpha=5,68253$, P-value=0,099 dan OR=293,69, ini memberikan makna bahwa responden yang berada di luar negeri mengapresiasi obyek wisata gajah liar (*elephant watching*) lebih besar 293,69 kali dari pada responden yang berada di dalam kabupaten.

4.3.2 Potensi Daya Tarik Objek Wisata Menunggang Gajah Jinak

Hasil analisis kuisioner potensi objek wisata menunggang gajah jinak pada Tabel 4.

Tabel 4. Hasil Optimasi Parameter Model Potensi Potensi Obyek Menunggang Gajah Jinak Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		β_0	-1,88491	0,520	
Demografi					
Umur					
- Dewasa Awal ($I=DAW$)	DAW	β_1	20,3949	0,998	7,20127E+08
-Dewasa Akhir ($I=DAR$)	DAR	β_2	18,9301	0,999	1,66440E+08
-Lansia Awal ($I=LAW$)	LAW	β_3	19,3490	0,999	2,53036E+08
Kelamin ($I=Lk$)	KLM	β_4	-1,95117	0,008	0,14
Alamat					
-Kota dalam Provinsi ($I=KTDP$)	KTDP	β_5	1,31419	0,103	3,72
-Kabupaten luar Provinsi ($I=KBLP$)	KBLP	β_6	-20,6935	0,999	0,00
-Kota luar Provinsi ($I=KTLP$)	KTLP	β_7	0,401302	0,696	1,49
-Luar Negeri ($I=LN$)	LN	β_8	2,02664	0,097	7,59
Pekerjaan					
-Pegawai Negeri Sipil ($I=PNS$)	PNS	β_9	-2,17951	0,255	0,11
-Non Pegawai Negeri Sipil ($I=NPNS$)	NPNS	β_{10}	0,937842	0,315	2,55
Pendidikan					
-Diploma ($I=DP$)	DP	β_{11}	3,12816	0,066	22,83
-Sarjana ($I=SJ$)	SJ	β_{12}	1,05021	0,299	2,86
Penghasilan					
-Rendah ($I=RDH$)	RDH	β_{13}	1,15495	0,199	3,17
-Sedang ($I=SDG$)	SDG	β_{14}	1,74634	0,160	5,73
-Tinggi ($I=TGI$)	TGI	β_{15}	2,94239	0,011	18,96
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	β_{16}	0,555807	0,835	1,74
-Lebih dari 5 kali ($I=LBDL$)	LBDL	β_{17}	-1,51322	0,153	0,22

Log-Likelihood = -34,434

Test that all slopes are zero: G = 59,772, DF = 17, P-Value = 0,0004

Keterangan: *nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 4 memperlihatkan bahwa potensi objek wisata menunggang gajah jinak dapat dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y2] memberikan hasil P-value=0,0004 (0,04 %). Hasil analisis pada Tabel 7, variabel jenis kelamin menghasilkan nilai duga sebesar ($\beta_4= -1,95117$) dengan (P-value=0,008) dan (Odd ratio=0,14), bermakna bahwa responden laki-laki dalam mengapresiasi objek wisata menunggang gajah jinak lebih rendah sebesar 0.14 dibandingkan responden perempuan. Hal ini bertentangan dengan [27] yang meyakini bahwa wisata menunggang gajah di PKG dapat dimanfaatkan oleh semua orang baik yang berjenis kelamin laki-laki maupun perempuan.

4.3.3 Potensi Daya Tarik Objek Wisata Sungai

Hasil analisis tentang potensi daya tarik wisata sungai disajikan Tabel 5.

Tabel 5. Hasil Optimasi Parameter Model Potensi Daya Tarik Obyek Wisata Sungai Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		γ_0	-41,3254	0,996	
Demografi					
Umur					

-Dewasa Awal (<i>I=DAW</i>)	DAW	γ_1	-42,2681	0,997	0,00
-Dewasa Akhir (<i>I=DAR</i>)	DAR	γ_2	-23,9445	0,999	0,00
-Lansia Awal (<i>I=LAW</i>)	LAW	γ_3	13,0090	1,000	446435,18
Kelamin (<i>I=Lk</i>)	KLM	γ_4	2,33684	0,009	10,35
Alamat					
-Kota dalam Provinsi (<i>I=KTDP</i>)	KTDP	γ_5	-0,263289	0,774	0,77
-Kabupaten luar Provinsi (<i>I=KBLP</i>)	KBLP	γ_6	20,9347	0,999	1,23542E+09
-Kota luar Provinsi (<i>I=KTLP</i>)	KTLP	γ_7	0,646803	0,653	1,91
-Luar Negeri (<i>I=LN</i>)	LN	γ_8	19,5170	0,998	2,99317E+08
Pekerjaan					
-Pegawai Negeri Sipil (<i>I=PNS</i>)	PNS	γ_9	3,98643	1,000	53,86
-Non Pegawai Negeri Sipil (<i>I=NPNS</i>)	NPNS	γ_{10}	-0,631957	0,701	0,53
Pendidikan					
-Diploma (<i>I=DP</i>)	DP	γ_{11}	-0,780167	0,725	0,46
-Sarjana (<i>I=SJ</i>)	SJ	γ_{12}	0,552220	0,697	1,74
Penghasilan					
-Rendah (<i>I=RDH</i>)	RDH	γ_{13}	1,60316	0,107	4,97
-Sedang (<i>I=SDG</i>)	SDG	γ_{14}	2,68010	0,095	14,59
-Tinggi (<i>I=TGI</i>)	TGI	γ_{15}	20,7065	0,997	9,83404E+08
Rutinitas Berwisata					
-Secara berkelompok (<i>I=KLMP</i>)	KLMP	γ_{16}	-15,3180	0,998	0,00 1,64556E+24
-Lebih dari 5 kali (<i>I=LBDL</i>)	LBDL	γ_{17}	55,7601	0,996	

Log-Likelihood = -22,673

Test that all slopes are zero: G = 74,646, DF = 17, P-Value = 0,0004

Keterangan: *nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 5 memperlihatkan bahwa potensi daya tarik obyek wisata sungai dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y3] mem-berikan hasil P-value = 0,0004 (0,04%). Hasil analisis Tabel 8, menghasilkan nilai duga sebesar ($\gamma_4=2,33684$) dengan (P-value=0,009) dan (OR=10,35), yang bermakna bahwa responden laki-laki dalam mengapresiasi obyek wisata sungai meningkat sebesar 10,35 dibandingkan responden perempuan. Hal ini bertentangan dengan [28] bahwa laki-laki dan perempuan memiliki kecenderungan yang sama dalam kebutuhan rekreasi atau wisata.

4.3.4 Potensi Daya Tarik Objek Wisata Air

Hasil analisis kuisioner terhadap potensi daya tarik wisata air disajikan dalam Tabel 6.

Tabel 6. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Air Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		δ_0	-1,94616	0,329	
Demografi					
Umur					
-Dewasa Awal (<i>I=DAW</i>)	DAW	δ_1	-22,6769	0,998	0,00
-Dewasa Akhir (<i>I=DAR</i>)	DAR	δ_2	-20,8091	0,999	0,00
-Lansia Awal (<i>I=LAW</i>)	LAW	δ_3	-18,4307	0,999	0,00
Kelamin (<i>I=Lk</i>)	KLM	δ_4	-0,972662	0,094	0,38
Alamat					
-Kota dalam Provinsi (<i>I=KTDP</i>)	KTDP	δ_5	-0,606811	0,345	0,55
-Kabupaten luar Provinsi (<i>I=KBLP</i>)	KBLP	δ_6	20,3811	0,999	7,10196E+08
-Kota luar Provinsi (<i>I=KTLP</i>)	KTLP	δ_7	0,132559	0,875	1,14
-Luar Negeri (<i>I=LN</i>)	LN	δ_8	-0,087835	0,943	0,92
Pekerjaan					
-Pegawai Negeri Sipil (<i>I=PNS</i>)	PNS	δ_9	0,335902	0,867	1,40
-Non Pegawai Negeri Sipil (<i>I=NPNS</i>)	NPNS	δ_{10}	-0,007760	0,993	0,99
Pendidikan					
-Diploma (<i>I=DP</i>)	DP	δ_{11}	0,682178	0,617	1,98
-Sarjana (<i>I=SJ</i>)	SJ	δ_{12}	0,425758	0,649	1,53
Penghasilan					
-Rendah (<i>I=RDH</i>)	RDH	δ_{13}	1,56271	0,023	4,77
-Sedang (<i>I=SDG</i>)	SDG	δ_{14}	0,708035	0,467	2,03
-Tinggi (<i>I=TGI</i>)	TGI	δ_{15}	1,99635	0,039	7,36
Rutinitas Berwisata					
-Secara berkelompok (<i>I=KLMP</i>)	KLMP	δ_{16}	-1,23528	0,484	0,29
-Lebih dari 5 kali (<i>I=LBDL</i>)	LBDL	δ_{17}	2,98447	0,026	19,78

Log-Likelihood = -44,547

Test that all slopes are zero: G = 42,594, DF = 17, P-Value = 0,001

Keterangan: *nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,001 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 6 memperlihatkan bahwa potensi daya tarik wisata air dapat dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y4] memberikan hasil P-value = 0,001 (0,1%). Responden

yang berpenghasilan tinggi dengan $\delta_{15} = 1,99635$, $P\text{-value} = 0,039$ dan $OR = 7,36$, ini memberikan makna bahwa responden yang berpenghasilan tinggi mengapresiasi potensi daya tarik wisata air lebih besar 7,36 kali dari pada responden yang berpenghasilan sangat rendah.

4.3.5 Potensi Daya Tarik Obyek Wisata Menunggang Kuda

Hasil optimasi parameter model potensi daya tarik obyek wisata pada Tabel 7.

Tabel 7. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Menunggang Kuda Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		ϵ_0	-19,2084	0,999	
Demografi					
Umur					
- Dewasa Awal ($I = DAW$)	DAW	ϵ_1	-0,637592	0,768	0,53
- Dewasa Akhir ($I = DAR$)	DAR	ϵ_2	-21,6297	0,999	0,00
- Lansia Awal ($I = LAW$)	LAW	ϵ_3	-21,5823	0,999	0,00
Kelamin ($I = Lk$)	KLM	ϵ_4	-3,32665	0,000	0,04
Alamat					
- Kota dalam Provinsi ($I = KTDP$)	KTDP	ϵ_5	-0,000404	1,000	1,00
- Kabupaten luar Provinsi ($I = KBLP$)	KBLP	ϵ_6	-22,8938	0,999	0,00
- Kota luar Provinsi ($I = KTLP$)	KTLP	ϵ_7	-4,84845	0,001	0,01
- Luar Negeri ($I = LN$)	LN	ϵ_8	-0,871474	0,667	0,42
Pekerjaan					
- Pegawai Negeri Sipil ($I = PNS$)	PNS	ϵ_9	-23,0806	0,999	0,00
- Non Pegawai Negeri Sipil ($I = NPNS$)	NPNS	ϵ_{10}	-0,061440	0,954	0,94
Pendidikan					
- Diploma ($I = DP$)	DP	ϵ_{11}	1,42426	0,438	4,15
- Sarjana ($I = SJ$)	SJ	ϵ_{12}	-0,81441	0,464	0,44
Penghasilan					
- Rendah ($I = RDH$)	RDH	ϵ_{13}	2,01032	0,040	7,47
- Sedang ($I = SDG$)	SDG	ϵ_{14}	1,65684	0,198	5,24
- Tinggi ($I = TGI$)	TGI	ϵ_{15}	3,58081	0,007	35,90
Rutinitas Berwisata					
- Secara berkelompok ($I = KLMP$)	KLMP	ϵ_{16}	19,2015	0,999	2,18318E+08
- Lebih dari 5 kali ($I = LBDL$)	LBDL	ϵ_{17}	0,547115	0,741	1,73

Log-Likelihood = -28,571

Test that all slopes are zero: $G = 61,353$, $DF = 17$, $P\text{-Value} = 0,0004$

Keterangan: *nyata pada taraf 15%; ** nyata pada taraf 5%

$P\text{-value} = 0,0004$ menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Dari hasil output Minitab 16 pada Tabel 7 memperlihatkan bahwa potensi daya tarik obyek wisata menunggang kuda dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y5] memberikan hasil $P\text{-value} = 0,0004$ (0,04%). Berdasarkan hasil penelitian Tabel 10, responden yang berpenghasilan tinggi pun dalam mengapresiasi potensi daya tarik obyek wisata menunggang kuda lebih tinggi yaitu 0,01 kali dari responden yang berpenghasilan sangat rendah. Ini dapat dilihat dari hasil analisis ($\epsilon_{15} = 3,58081$, $P\text{-value} = 0,007$, $OR = 35,90$).

4.3.6 Potensi Daya Tarik Objek Wisata Sentra Kuliner

Hasil analisis kuisioner terhadap potensi objek wisata sentra kuliner pada Tabel 8.

Tabel 8. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Sentra Kuliner Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		η_0	-59,7603	0,997	
Demografi					
Umur					
- Dewasa Awal ($I = DAW$)	DAW	η_1	19,0649	0,998	1,90442E+08
- Dewasa Akhir ($I = DAR$)	DAR	η_2	2,00278	1,000	7,41
- Lansia Awal ($I = LAW$)	LAW	η_3	16,0405	1,000	9252961,02
Kelamin ($I = Lk$)	KLM	η_4	-2,04978	0,016	0,13
Alamat					
- Kota dalam Provinsi ($I = KTDP$)	KTDP	η_5	1,95833	0,052	7,09
- Kabupaten luar Provinsi ($I = KBLP$)	KBLP	η_6	-1,72389	0,190	0,18
- Kota luar Provinsi ($I = KTLP$)	KTLP	η_7	-0,696461	0,501	0,50
- Luar Negeri ($I = LN$)	LN	η_8	1,70997	0,242	5,53
Pekerjaan					
- Pegawai Negeri Sipil ($I = PNS$)	PNS	η_9	19,4043	0,998	2,67401E+08
- Non Pegawai Negeri Sipil ($I = NPNS$)	NPNS	η_{10}	-0,512225	0,785	0,60
Pendidikan					
- Diploma ($I = DP$)	DP	η_{11}	-3,53548	0,109	0,03
- Sarjana ($I = SJ$)	SJ	η_{12}	-1,34419	0,496	0,26

Penghasilan					
-Rendah ($I=RDH$)	RDH	η_{13}	0,040799	0,962	1,04
-Sedang ($I=SDG$)	SDG	η_{14}	21,0639	0,997	1,40578E+09
-Tinggi ($I=TGI$)	TGI	η_{15}	2,81888	0,049	16,76
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	η_{16}	20,0597	0,999	5,15007E+08
-Lebih dari 5 kali ($I=LBDL$)	LBDL	η_{17}	42,1285	0,996	1,97770E+18

Log-Likelihood = -27,708

Test that all slopes are zero: G = 65,979, DF = 17, P-Value = 0,0004

Keterangan: -*nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 8 memperlihatkan bahwa potensi daya tarik objek wisata sentra kuliner dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y6] memberikan hasil P-value=0,0004 (0,04 %). Hasil analisis pada Tabel 10, variabel jenis kelamin menghasilkan nilai duga sebesar ($\eta_4 = -2,04978$) dengan (P-value=0,016) dan (Odd ratio=0,13), yang bermakna bahwa responden laki-laki dalam mengapresiasi potensi daya tarik objek wisata sentra kuliner lebih rendah sebesar 0.13 kali dibandingkan responden perempuan. Hal ini sejalan dengan penelitian [29] yang menyatakan bahwa jenis kelamin wisatawan cenderung menentukan jenis dan pilihan dalam melakukan perjalanan, termasuk menentukan jenis dan pilihan makanan yang mereka santap. Begitu juga dengan pendidikan responden, dari Tabel 11 diperoleh data $\eta_{11} = -3,53548$, P-value=0,109 dan OR=0,03, memberikan makna bahwa responden yang berpendidikan diploma lebih tinggi 0,03 kali dalam mengapresiasi potensi daya tarik objek wisata sentra kuliner dibandingkan dengan responden yang berpendidikan menengah. Tingkat pendidikan mempengaruhi konsumen dalam pemilihan produk yang diinginkannya, hal ini dikarenakan tingkat pendidikan seseorang mempengaruhi cara berfikir dan persepsinya terhadap suatu produk yang dikonsumsi [30].

4.3.7 Potensi Daya Tarik Objek Wisata Sentra Souvenir

Hasil analisis tentang potensi daya tarik wisata sentra souvenir disajikan Tabel 8.

Tabel 8. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Sentra Souvenir Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		λ_0	-18,3315	0,999	
Demografi					
Umur					
- Dewasa Awal ($I=DAW$)	DAW	λ_1	-4,17888	0,023	0,02
-Dewasa Akhir ($I=DAR$)	DAR	λ_2	-21,0723	0,999	0,00
-Lansia Awal ($I=LAW$)	LAW	λ_3	-18,8773	0,999	0,00
Kelamin ($I=Lk$)	KLM	λ_4	-4,13990	0,000	0,02
Alamat					
-Kota dalam Provinsi ($I=KTDP$)	KTDP	λ_5	1,93489	0,050	6,92
-Kabupaten luar Provinsi ($I=KBLP$)	KBLP	λ_6	-22,8835	0,998	0,00
-Kota luar Provinsi ($I=KTLP$)	KTLP	λ_7	-2,95497	0,024	0,05
-Luar Negeri ($I=LN$)	LN	λ_8	-0,648662	0,737	0,52
Pekerjaan					
-Pegawai Negeri Sipil ($I=PNS$)	PNS	λ_9	-21,0560	0,999	0,00
-Non Pegawai Negeri Sipil ($I=NPNS$)	NPNS	λ_{10}	-0,219282	0,871	0,80
Pendidikan					
-Diploma ($I=DP$)	DP	λ_{11}	1,18563	0,531	3,27
-Sarjana ($I=SJ$)	SJ	λ_{12}	0,235723	0,879	1,27
Penghasilan					
-Rendah ($I=RDH$)	RDH	λ_{13}	0,364375	0,711	1,44
-Sedang ($I=SDG$)	SDG	λ_{14}	-0,798226	0,636	0,45
-Tinggi ($I=TGI$)	TGI	λ_{15}	0,389223	0,778	1,48
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	λ_{16}	17,9559	0,999	62830176,65
-Lebih dari 5 kali ($I=LBDL$)	LBDL	λ_{17}	2,19685	0,315	9,00

Log-Likelihood = -27,235

Test that all slopes are zero: G = 71,604, DF = 17, P-Value = 0,0004

Keterangan: -*nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 8 memperlihatkan bahwa potensi daya tarik obyek wisata sentra souvenir dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y7] memberikan hasil P-value = 0,0004 (0,04%). Hasil analisis Tabel 8, berdasarkan variabel umur menghasilkan nilai duga sebesar ($\lambda_1 = -4,17888$) dengan (P-value=0,023) dan (OR=0,02), yang bermakna bahwa responden yang berumur dewasa awal dalam mengapresiasi potensi daya tarik obyek wisata sentra souvenir meningkat sebesar 0,02 kali dibandingkan responden yang berumur remaja akhir. Ini berarti bahwa umur responden mempengaruhi calon wisatawan untuk berkunjung ke obyek wisata sentra souvenir. Berbeda dengan [31] semakin tinggi tingkat umur pengunjung, semakin kecil jumlah pengunjung ke objek wisata.

4.3.8 Potensi Daya Tarik Wisata Kebun Nangkada

Hasil wawancara dengan menggunakan kuisioner terhadap potensi daya tarik wisata kebun nangkada disajikan dalam Tabel 9.

Tabel 9. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Kebun Nangkada Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
Constant		ν_0	-19,8661	0,999	
Demografi					
Umur					
- Dewasa Awal ($I=DAW$)	DAW	ν_1	-3,00924	0,079	0,05
-Dewasa Akhir ($I=DAR$)	DAR	ν_2	-22,1413	0,999	0,00
-Lansia Awal ($I=LAW$)	LAW	ν_3	-20,9950	0,999	0,00
Kelamin ($I=Lk$)	KLM	ν_4	-2,92471	0,000	0,05
Alamat					
-Kota dalam Provinsi ($I=KTDP$)	KTDP	ν_5		0,685	
-Kabupaten luar Provinsi ($I=KBLP$)	KBLP	ν_6	0,302661	0,999	1,35
-Kota luar Provinsi ($I=KTLP$)	KTLP	ν_7	-22,7275	0,002	0,00
-Luar Negeri ($I=LN$)	LN	ν_8	-4,31391	0,664	0,01
Pekerjaan					
-Pegawai Negeri Sipil ($I=PNS$)	PNS	ν_9	-0,808786	0,999	0,45
-Non Pegawai Negeri Sipil ($I=NPNS$)	NPNS	ν_{10}	-22,6333	0,888	0,00
Pendidikan					
-Diploma ($I=DP$)	DP	ν_{11}	0,150047		1,16
-Sarjana ($I=SI$)	SI	ν_{12}	1,61382	0,363	5,02
Penghasilan					
-Rendah ($I=RDH$)	RDH	ν_{13}	-0,517368	0,637	0,60
-Sedang ($I=SDG$)	SDG	ν_{14}			
-Tinggi ($I=TGI$)	TGI	ν_{15}	2,10793	0,025	8,23
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	ν_{16}	1,47222	0,240	4,36
-Lebih dari 5 kali ($I=LBDL$)	LBDL	ν_{17}	2,90913	0,019	18,34
			19,0271	0,999	1,83379E+08
			0,969438	0,566	2,64

Log-Likelihood = -30,834

Test that all slopes are zero: G = 56,826, DF = 17, P-Value = 0,0004

Keterangan: *nyata pada taraf 15%; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Hasil output Minitab 16 pada Tabel 9 memperlihatkan bahwa potensi daya tarik kebun nangkada berbasis demografi, pendidikan dan rutinitas berwisata dapat dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y8] memberikan hasil P-value = 0,0004 (0,04%). Hasil analisis pada Tabel 9, berdasarkan variabel alamat, responden yang berada di kota luar provinsi dalam mengapresiasi potensi daya tarik obyek wisata kebun nangkada akan meningkat menjadi 0,01 kali ($\nu_7 = -4,31391$, P-value=0,002, OR=0,01).

4.3.9 Potensi Daya Tarik Wisata Tarian Kuda Lumping

Hasil wawancara dengan menggunakan kuisioner terhadap potensi daya tarik wisata tarian kuda lumpung disajikan dalam Tabel 10.

Tabel 14. Hasil Optimasi Parameter Model Potensi Potensi Obyek Wisata Tarian Kuda Lumping Berbasis Demografi, Pendidikan dan Rutinitas Berwisata

Predictor	Simbol dalam Model	Simbol	Coef	P	OR
-----------	--------------------	--------	------	---	----

	Model				
Constant		χ_0	-6,58911	0,209	
Demografi					
Umur					
- Dewasa Awal ($I=DAW$)	DAW	χ_1	-0,843338	0,509	0,43
-Dewasa Akhir ($I=DAR$)	DAR	χ_2	16,1621	1,000	10449423,05
-Lansia Awal ($I=LAW$)	LAW	χ_3	17,4466	0,999	37754156,24
Kelamin ($I=Lk$)	KLM	χ_4	0,845677	0,299	2,33
Alamat					
-Kota dalam Provinsi ($I=KTDP$)	KTDP	χ_5	3,39376	0,006	29,78
-Kabupaten luar Provinsi ($I=KBLP$)	KBLP	χ_6	1,80932	0,502	6,11
-Kota luar Provinsi ($I=KTLP$)	KTLP	χ_7	0,900643	0,492	2,46
-Luar Negeri ($I=LN$)	LN	χ_8	4,56620	0,003	96,18
Pekerjaan					
-Pegawai Negeri Sipil ($I=PNS$)	PNS	χ_9	16,4786	0,999	14340124,90
-Non Pegawai Negeri Sipil ($I=NPNS$)	NPNS	χ_{10}	-0,361601	0,771	0,70
Pendidikan					
-Diploma ($I=DP$)	DP	χ_{11}	4,17458	0,120	65,01
-Sarjana ($I=SJ$)	SJ	χ_{12}	1,29358	0,301	3,65
Penghasilan					
-Rendah ($I=RDH$)	RDH	χ_{13}	0,845286	0,426	2,33
-Sedang ($I=SDG$)	SDG	χ_{14}	3,62452	0,012	37,51
-Tinggi ($I=TGI$)	TGI	χ_{15}	5,19302	0,001	180,01
Rutinitas Berwisata					
-Secara berkelompok ($I=KLMP$)	KLMP	χ_{16}	1,47990	0,761	4,39
-Lebih dari 5 kali ($I=LBDL$)	LBDL	χ_{17}	-1,32153	0,264	0,27

Log-Likelihood = -30,834

Test that all slopes are zero: G = 72,511, DF = 17, P-Value = 0,0004

Keterangan: -*nyata pada taraf 15% ; ** nyata pada taraf 5%

P-value = 0,0004 menyatakan bahwa berbeda nyata dengan tingkat kesalahan di bawah 10%

Dari hasil output Minitab 16 pada Tabel 10 memperlihatkan bahwa potensi daya tarik obyek wisata tarian kuda lumping dinilai sangat baik. Penilaian tersebut didasarkan bahwa untuk [Y9] memberikan hasil P-value = 0,0004 (0,04%). Berdasarkan hasil penelitian Tabel 14, data demografi yang terdiri dari jenis kelamin, umur, alamat dan pekerjaan, variabel alamat responden yang berasal dari luar negeri dalam mengapresiasi potensi daya tarik obyek wisata tarian kuda lumping lebih tinggi yaitu 96,18 kali dari responden yang berada di kabupaten dalam provinsi. Ini dapat dilihat dari hasil analisis ($\chi_8=4,56620$, P-value=0,003, OR=96,18).

5. SIMPULAN

Pengembangan ekowisata yang semula terdiri dari 7 obyek ekowisata dapat dikembangkan menjadi 16 obyek ekowisata, aktivitas gajah liar yang selalu masuk kedalam wilayah desa yang berbatasan dengan kawasan Taman Nasional Way Kambas dapat dijadikan obyek ekowisata yang terintegrasi dengan potensi sumber daya wisata pedesaan dan potensi-potensi yang ada di Desa Braja Harjosari dari analisis hasil output Minitab 16 berbasis demografi, pendidikan dan rutinitas wisata dinilai sangat baik.

KEPUSTAKAAN

- [1] Cross, M. (2013). *Pengertian Media Sosial* [diakses 24 Agustus 2017]. Tersedia dari <https://www.pakarkomunikasi.com>
- [2] Danniel. (2004). *Pengantar Ekonomi Pertanian*. Bumi Aksara. Jakarta.
- [3] Widjaja, H.A.W. (2003). *Otonomi Desa*. PT Raja Grafindo Persada. Jakarta.
- [4] Putra, C.K. (2013). *Pemberdayaan Masyarakat Desa dalam Pemberdayaan Masyarakat Desa*. Jurnal Administrasi. Malang.
- [5] Setiadi. (2013). *Konsep dan Praktek Penulisan Riset Keperawatan*. Graha Ilmu. Yogyakarta.
- [6] Rachman. (2015). *Strategi Komunikasi Pemasaran Desa Wisata Melalui Media*. Laporan Studi Pustaka IPB. Bogor.

- [7] Saktiawan. (2010). *Pentingnya Membangun Partisipasi Masyarakat Dalam Pengembangan Desa Wisata* [diakses 12 Desember 2016]. Tersedia dari <http://buletin.betungkerihun.wordpress.com>.
- [8] Wiratno. (2004). *Berkaca Di Cermin Retak "Refleksi Konservasi dan Implikasi Bagi Pengelolaan Taman Nasional"*. FOReST Press, The Gibbon Foundation Indonesia, Departemen Kehutanan, PILINGO Movement. Jakarta.
- [9] TIES (The International Ecotourism Society). (2006). *Fact Sheet: Global Ecotourism* [diakses 2 September 2016]. Tersedia dari <https://www.ecotourism.org>.
- [10] Dowling, R. K. and Page, S.J. (2002). *Ecotourism*. Prentice Hall. London.
- [11] Gumelar, S. S. (2010). *Strategi Pengembangan Dan Pengelolaan Resort And Leisure*. Hand Out Mata Kuliah Concept Resort And Leisure, Universitas Pendidikan Indonesia. Bandung.
- [12] Manurung. (2002). *Ecotourism in Indonesia*. In: Hundloe, T (ed.). *Linking Green Productivity to Ecotourism : Experiences in the Asia-Pacific Region*. Asian Productivity Organization (APO). Japan. 98-103.
- [13] Ritonga, T. A. (2012). *Pemanfaatan Gajah Jinak Dalam Kegiatan Conservation Response Unit (Cru) Di Tangkahan*. Universitas Sumatera Utara. Sumatera Utara.
- [14] Firqan, I. (2012). *Melirik peran dan daya guna taman konservasi Lampung* [diakses 23 November 2012]. Tersedia dari <http://astacala.org/wp/2012/03/melirik-peran-dan-dayaguna-taman-konservasi-gajah-di-lampung.html>
- [15] Ankre. (2006). *Zonning and Opportunity Spectrum Planning In A Discountinous Environment*. Paper Planning for Tourism and Outdoor Recreation in the Lulea Archipelago. Sweden.
- [16] Hamzah, Y, I. (2013). *Potensi Media Sosial Sebagai Sarana Promosi Interaktif bagi Pariwisata Indonesia*. Jurnal Kepariwisata Indonesia Vo. 8 No. 3 2013. Jakarta : Puslitbangjak Badan Pengembangan Sumberdaya Kementerian Parekraf.
- [17] Ankre. (2006). *Zonning and Opportunity Spectrum Planning In A Discountinous Environment*. Paper Planning for Tourism and Outdoor Recreation in the Lulea Archipelago. Sweden.
- [18] Nisrina. (2015). *Bisnis Online, Manfaat Media Sosial Dalam Meraup Uang*. Kobis. Yogyakarta.
- [19] Inskeep, E. (1991). *Tourism Planning, and Integrated and Sustainable Development Approach*. Van Nostrand Reinhold. New york.
- [20] Wiendu. (1993). *Concept, Perspective, and Challenges, makalah bagian dari Laporan Konferensi Internasional mengenai Pariwisata Budaya*. Gadjah Mada University Press. Yogyakarta.
- [21] Aukerman. (2004). *Water Recreation Opportunity Spectrum (WROS) Users' Guidebook*. Bureau of Reclamation. United States Department of the Interior.
- [22] Sujadi. (2003). *Metodologi Penelitian Pendidikan*. Rineka Cipta. Jakarta.
- [23] Departemen Pariwisata, Pos dan Telekomunikasi. 1987. *Surat Keputusan Menteri Pariwisata, Pos Telekomunikasi No. KM94/HK/ 103 MPPT 87*. Kementrian Pariwisata, Pos Telekomunikasi RI.
- [24] Sunaryo B. (2013). *Kebijakan Pembangunan Destinasi Pariwisata: Konsep dan Aplikasinya di Indonesia*. Gava Media. Yogyakarta.

- [25] Yoeti, O. (1997). *Pengantar ilmu pariwisata*. Angkasa. Bandung.
- [26] Pendit, N. S. (2003). *Ilmu Pariwisata Sebuah Pengantar Perdana*. PT. Pradnya Paramita. Jakarta.
- [27] Subangkit L. (2014). *Faktor-Faktor Kepuasan Pengunjung di Pusat Konservasi Gajah Taman Nasional Way Kambas Lampung*. Jurnal Sylva Lestari (2:3).101-110.
- [28] Premono BT & Kunarso A. (2008). *Pengaruh perilaku pengunjung terhadap jumlah kunjungan di taman wisata alam panti kayu Palembang*. Jurnal penelitian hutan dan konservasi alam. 5(5):423-433.
- [29] Harsana W dan Maria TW. (2005). Makalah: *Analisis Pasar Ditinjau dari Persepsi Wisatawan Terhadap Kuliner di Kabupaten Sleman*. Universitas Negeri Yogyakarta. 33 halaman.
- [30] Sungkawa I, Dwi P, dan Eva F. (2015). *Hubungan antara Persepsi dan Prefensi Konsumen dengan Pengambilan Keputusan Pembelian Buah Lokal*. Jurnal Agrijati. (28:1). 79-99.
- [31] Modjanggo, F, Arief, S, Susti. (2015). *Faktor-faktor yang Mempengaruhi Jumlah Pengunjung Ke Objek Ekowisata Pantai Siuri, Desa Toinas Kecamatan Pamona Barat Kabupaten Poso*. Warta Rimba. (3:2). 88-95.

DIAGONALISASI SECARA UNITER MATRIKS HERMITE DAN APLIKASINYA PADA PENGAMANAN PESAN RAHASIA

Abdurrois¹⁾, Dorrah Aziz, dan Aang Nuryaman

Jurusan Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam
Universitas Lampung, Jl. Prof. Soemantri Brodjonegoro No. 1, Bandar Lampung 35145
abdurrois89@gmail.com¹⁾

ABSTRAK

Matriks Hermite adalah matriks yang hasil dari transpos konjugatnya sama dengan matriks itu sendiri. Diagonalisasi matriks Hermite secara uniter merupakan proses untuk mendekomposisikan matriks Hermite menjadi matriks diagonal dimana unsur-unsur dari diagonal utamanya merupakan nilai eigen dari matriks Hermite. Dalam artikel ini dikaji penerapan diagonalisasi matriks Hermite secara uniter sebagai landasan membuat matriks kunci untuk proses enkripsi dan deskripsi pesan rahasia melalui algoritma kriptografi Hill Cipher. Beberapa hasil enkripsi dan deskripsi dengan ukuran matriks berbeda disajikan dalam artikel ini.

Kata kunci: Matriks Hermite, Diagonalisasi, Kriptografi.

1. PENDAHULUAN

Matriks merupakan salah satu cabang dari ilmu aljabar linier yang memiliki peran yang sangat penting di dalam matematika. Pentingnya peranan matriks ini dapat dilihat dari begitu banyaknya penggunaan matriks dalam berbagai bidang antara lain aljabar, statistika, metode numerik, persamaan diferensial dan lain-lain.

Matriks bujur sangkar merupakan salah satu syarat dalam menentukan diagonalisasi dalam sebuah matriks. Diagonalisasi matriks banyak diterapkan dalam berbagai ilmu matematika, misalnya dalam irisan kerucut dan persamaan differensial dimana dalam diagonalisasi yang dilakukan pada matriks dengan unsur real. Penerapan diagonalisasi juga dapat dilakukan pada matriks dengan unsur bilangan kompleks, contohnya matriks *Hermite*. Matriks *Hermite* merupakan matriks bujur sangkar dengan unsur bilangan kompleks yang memenuhi sifat $A = A^*$ dimana A^* adalah matriks transpos konjugat dari A . Pendiagonalan suatu matriks *Hermite* sangatlah diperlukan terutama saat menghitung matriks *Hermite* A^n dimana matriks A^n digunakan sebagai kunci penyandi untuk pengamanan pesan rahasia.

Matriks *Hermite* dapat diaplikasikan untuk proses pengamanan pesan rahasia, hal ini layak diterapkan di era globalisasi karena kerahasiaan adalah suatu hal yang sangat penting di jaman serba modern saat ini. Dalam penelitian ini matriks *Hermite* yang digunakan adalah ordo 5×5 dan 6×6 serta diagonal utamanya adalah bilangan real kemudian memilih matriks A^2 untuk ordo 5×5 dan E^3 untuk ordo 6×6 sebagai kunci penyandinya dimana A dan E adalah matriks *Hermite*. Berdasarkan masalah tersebut, penulis ingin mengembangkan salah satu manfaat dari diagonalisasi pada matriks kompleks khususnya pada matriks *Hermite* dengan mengaplikasikannya pada pengamanan pesan rahasia dengan menggunakan alat bantu software Matlab R2013b.

2. LANDASAN TEORI

2.1 Diagonalisasi Matriks

Sebuah matriks bujur sangkar A dikatakan dapat didiagonalisasikan jika terdapat sebuah matriks P yang dapat dibalik sedemikian rupa sehingga $P^{-1}AP$ adalah sebuah matriks diagonal, matriks P dikatakan mendiagonalisasi A [1].

Prosedur untuk mendiagonalisasikan sebuah matriks :

1. Menentukan n vektor eigen dari A yang bebas linier, misalkan $\mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_n$
2. Membentuk sebuah matriks P dengan $\mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_n$ sebagai vektor-vektor kolomnya.
3. Matriks $P^{-1}AP$ kemudian akan menjadi diagonal dengan $\lambda_1, \lambda_2, \dots, \lambda_n$ sebagai entri-entri diagonalnya secara berurutan, dimana λ_i adalah nilai eigen yang terkait dengan $\mathbf{p}_i$ untuk $i = 1, 2, \dots, n$.

Sebuah matriks bujursangkar A dengan entri-entri kompleks dikatakan secara uniter dapat didiagonalkan apabila terdapat sebuah matriks uniter P sedemikian rupa sehingga $P^{-1}AP (= P^*AP)$ adalah matriks diagonal, dan matriks P dikatakan secara uniter mendiagonalisasi A [2].

2.2 Matriks Kompleks

Jika A adalah sebuah matriks yang memiliki entri-entri bilangan kompleks, maka transpos konjugat matriks A , yang dinotasikan dengan A^* , didefinisikan sebagai

$$A^* = \overline{A}^T \quad (1)$$

dimana $\overline{A}$ adalah sebuah matriks yang entri-entrinya adalah konjugat-konjugat kompleks dari entri-entri yang bersesuaian pada matriks A dan $\overline{A}^T$ adalah transpos dari matriks $\overline{A}$.

Sebuah matriks bujur sangkar A dengan entri-entri bilangan kompleks disebut matriks uniter jika

$$A^{-1} = A^* \quad (2)$$

Sebuah matriks bujur sangkar A dengan entri-entri bilangan kompleks disebut matriks *Hermite* jika

$$A = A^* \quad (3)$$

Matriks *Hermite* merupakan bentuk lain dari matriks simetri pada matriks dengan unsur bilangan riil. Pada matriks dengan elemen riil, matriks simetri didefinisikan dengan $A = A^T$, sama halnya dengan matriks *Hermite* yaitu $A = A^*$, dimana A^* merupakan transpos konjugat dari A .

Sebuah matriks bujur sangkar A dengan entri-entri bilangan kompleks disebut matriks normal jika

$$AA^* = A^*A \quad (4)$$

Setiap matriks *Hermite* merupakan matriks normal karena $AA^* = AA = A^*A$ dan setiap matriks uniter adalah matriks normal karena $AA^* = I = A^*A$ [2].

2.3 Kriptografi

Kriptografi berasal dari bahasa Yunani, terdiri dari dua suku kata yaitu *kripto* dan *graphia*. *Kripto* artinya menyembunyikan, sedangkan *graphia* artinya tulisan. Kriptografi adalah ilmu yang mempelajari teknik-teknik matematika yang berhubungan dengan aspek keamanan informasi, seperti kerahasiaan data, keabsahan data, integritas data, serta autentikasi data. Tetapi tidak semua aspek keamanan informasi dapat diselesaikan dengan kriptografi. Kriptografi dapat pula diartikan sebagai ilmu atau seni untuk menjaga keamanan pesan. Ketika suatu pesan dikirim dari suatu tempat ke tempat lain, isi pesan tersebut mungkin dapat disadap oleh pihak lain yang tidak berhak untuk mengetahui isi pesan tersebut. Untuk menjaga pesan, maka pesan tersebut dapat diubah menjadi suatu kode yang tidak dapat dimengerti oleh pihak lain.

Enskripsi adalah sebuah proses penyandian yang melakukan perubahan sebuah kode (pesan) dari yang bisa dimengerti (*plaintexts*) menjadi sebuah kode yang tidak bisa dimengerti (*cipherteks*). Sedangkan proses kebalikannya untuk mengubah cipherteks menjadi plaintexts disebut dekripsi. Proses enkripsi dan dekripsi memerlukan suatu mekanisme dan kunci tertentu [3].

2.4 Hill Cipher

Pada tahun 1929 Lester S. Hill menciptakan *Hill Cipher*. Teknik kriptografi ini diciptakan dengan maksud untuk dapat menciptakan *cipher* (kode) yang tidak dapat dipecahkan menggunakan teknik analisis frekuensi. *Hill Cipher* tidak mengganti setiap abjad yang sama pada *plaintext* dengan abjad lainnya yang sama pada *ciphertext* karena menggunakan perkalian matriks pada dasar enkripsi dan dekripsinya. *Hill Cipher* yang merupakan *polyalphabetic cipher* dapat dikategorikan sebagai *block cipher*. Karena teks yang akan diproses akan dibagi menjadi blok-blok dengan ukuran tertentu. Setiap karakter dalam satu blok akan saling mempengaruhi karakter lainnya dalam proses enkripsi dan dekripsinya, sehingga karakter yang sama belum tentu dipetakan menjadi karakter yang sama pula [4].

Secara umum dengan menggunakan matriks $K_{m \times m}$ sebagai kunci. Jika elemen pada baris i dan kolom j dari matriks K adalah $k_{i,j}$, maka dapat ditulis $K = k_{i,j}$. Untuk $x = (x_1, x_2, \dots, x_m)$ dan $y = e_k(x) = (y_1, y_2, \dots, y_m)$ sebagai berikut:

$$(y_1, y_2, \dots, y_m) = (x_1, x_2, \dots, x_n) \begin{pmatrix} k_{11} & k_{12} & \dots & k_{1n} \\ k_{21} & k_{22} & \dots & k_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ k_{n1} & k_{n2} & \dots & k_{nn} \end{pmatrix} \quad (5)$$

dengan kata lain, $y = xK$.

Untuk melakukan deskripsi menggunakan matriks invers K^{-1} . Jadi deskripsi dilakukan dengan rumus $x = yK^{-1}$ [5].

3. METODE PENELITIAN

Metode yang dipergunakan dalam penelitian ini adalah studi kepustakaan (literatur). Adapun langkah-langkah yang akan dilakukan penulis dalam menyelesaikan penelitian ini adalah sebagai berikut:

1. Mendiagonalisasi matriks *Hermite* secara uniter.
2. Menghitung matriks A^n , $n \in \mathbb{Z}^+$ dengan menerapkan pendioaganalan matriks *Hermite* untuk dijadikan matriks kunci penyandi pengamanan pesan rahasia.
3. Melakukan enkripsi dan deskripsi dengan menggunakan algoritma kriptografi *Hill Cipher*.

4. HASIL DAN PEMBAHASAN

4.1 Diagonalisasi Matriks *Hermite*

Matriks *Hermite* merupakan salah satu contoh dari matriks kompleks yang dapat di diagonalisasikan secara uniter. Anggota diagonal utama dari matriks *Hermite* adalah bilangan real. Diagonalisasi matriks *Hermite* berarti membentuk matriks yang diagonal dari matriks *Hermite* yang telah diketahui. Matriks *Hermite* dapat didiagonalisasikan secara uniter apabila kita telah mendapatkan matriks yang ortogonal atau pada kompleks dikatakan uniter.

Selanjutnya akan dijelaskan diagonalisasi pada matriks *Hermite* untuk ordo 5×5 dan 6×6 melalui contoh berikut:

Contoh 1. Diagonalisasi Secara Uniter Matriks *Hermite* Ordo 5×5

1. Diberikan matriks *Hermite* A dengan ordo 5×5 .

$$A = \begin{bmatrix} 1 & 0 & 1-i & 0 & 0 \\ 0 & 1 & 3-i & 0 & 0 \\ 1+i & 3+i & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

Yang memenuhi sifat $A = \overline{A}^T = A^*$.

2. Menentukan polinomial karakteristik dari matriks A

$$\det(\lambda I - A) = 0$$

maka

$$\det \begin{pmatrix} 1 & 0 & 1-i & 0 & 0 \\ 0 & 1 & 3-i & 0 & 0 \\ 1+i & 3+i & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} = 0$$

Sehingga polinomial karakteristik A adalah

$$\lambda^5 - 6\lambda^4 + 2\lambda^3 + 20\lambda^2 - 27\lambda + 10 = 0$$

3. Menentukan nilai eigen dan vektor eigen dari matriks A ,

Dengan memfaktorkan polinomial dari matriks sebagai berikut:

$$\lambda^5 - 6\lambda^4 + 2\lambda^3 + 20\lambda^2 - 27\lambda + 10 = 0$$

$$(\lambda + 2)(\lambda - 1)^3(\lambda - 5)^2 = 0$$

Sehingga nilai eigen dari matriks adalah $\lambda = -2$, $\lambda = 1$, dan $\lambda = 5$.

Secara definisi,

$$x = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{bmatrix}$$

merupakan sebuah vektor eigen dari A yang diasosiasikan dengan λ jika dan hanya jika x adalah solusi nontrivial bagi

$$\begin{bmatrix} \lambda - 1 & 0 & -1 + i & 0 & 0 \\ 0 & \lambda - 1 & -3 + i & 0 & 0 \\ -1 - i & -3 - i & \lambda - 2 & 0 & 0 \\ 0 & 0 & 0 & \lambda - 1 & 0 \\ 0 & 0 & 0 & 0 & \lambda - 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (6)$$

Untuk menentukan vektor-vektor eigen yang diasosiasikan dengan λ dengan mensubstitusikan nilai eigen ke persamaan (6).

Untuk $\lambda = -2$ berdasarkan persamaan (6) dengan menggunakan operasi baris elementer maka diperoleh vektor eigen untuk $\lambda = -2$ adalah

$$u_1 = \begin{bmatrix} \frac{-1+i}{3} \\ \frac{3+i}{3} \\ 1 \\ 0 \\ 0 \end{bmatrix}$$

Dengan cara yang sama pada $\lambda = 1$, dan $\lambda = 5$ maka diperoleh

$$\mathbf{u}_2 = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}, \mathbf{u}_3 = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}, \mathbf{u}_4 = \begin{bmatrix} -2+i \\ 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}, \text{ dan } \mathbf{u}_5 = \begin{bmatrix} \frac{1-i}{4} \\ \frac{3-i}{4} \\ 1 \\ 0 \\ 0 \end{bmatrix}$$

4. Menerapkan proses Gram-Schmidt untuk mendapatkan basis ortonormal.

Menentukan kolom pertama $\mathbf{p}_1$ dengan mencari $\|\mathbf{u}_1\|$ terlebih dahulu sebagai berikut:

$$\|\mathbf{u}_1\| = \sqrt{\left|\frac{-1+i}{3}\right|^2 + \left|\frac{3+i}{3}\right|^2 + |1|^2 + |0|^2 + |0|^2} = \sqrt{\frac{7}{3}}$$

$$\text{Sehingga diperoleh } \mathbf{p}_1 = \frac{\mathbf{u}_1}{\|\mathbf{u}_1\|} = \begin{bmatrix} \frac{-1+i}{\sqrt{21}} \\ \frac{-3-i}{\sqrt{21}} \\ \frac{3}{\sqrt{21}} \\ 0 \\ 0 \end{bmatrix}$$

Dengan cara yang sama pada $\mathbf{u}_2, \mathbf{u}_3, \mathbf{u}_4$, dan $\mathbf{u}_5$ diperoleh

$$\mathbf{p}_2 = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}, \mathbf{p}_3 = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}, \mathbf{p}_4 = \begin{bmatrix} \frac{-2+i}{\sqrt{6}} \\ \frac{1}{\sqrt{6}} \\ 0 \\ 0 \\ 0 \end{bmatrix}, \text{ dan } \mathbf{p}_5 = \begin{bmatrix} \frac{1-i}{2\sqrt{7}} \\ \frac{3-i}{2\sqrt{7}} \\ \frac{2}{\sqrt{7}} \\ 0 \\ 0 \end{bmatrix}$$

5. Dengan proses Gram-schmidt maka diperoleh matriks P yang mendiagonalisasikan matriks A sebagai berikut:

$$\overline{P}^T = \begin{bmatrix} \frac{-1-i}{\sqrt{21}} & \frac{-3+i}{\sqrt{21}} & \frac{3}{\sqrt{21}} & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 \\ \frac{-2-i}{\sqrt{6}} & \frac{1}{\sqrt{6}} & 0 & 0 & 0 \\ \frac{1+i}{2\sqrt{7}} & \frac{3+i}{2\sqrt{7}} & \frac{2}{\sqrt{7}} & 0 & 0 \end{bmatrix} = P^{-1}$$

Sehingga

$$P^{-1}AP = \begin{bmatrix} -2 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 5 \end{bmatrix} = D$$

Berdasarkan pembuktian di atas maka jelaslah bahwa matriks A dapat didiagonalkan secara uniter oleh matriks P .

Contoh 2. Diagonalisasi Secara Uniter Matriks Hermite Ordo 6×6

1. Diberikan matriks Hermite A dengan ordo 6×6 .

$$E = \begin{bmatrix} 1 & 0 & 2i & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 & 2+i & 0 \\ -2i & 0 & 1 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & 1 & 0 & 0 \\ 0 & 2-i & 0 & 0 & -2 & 0 \\ 0 & 0 & 0 & -1-i & 0 & 1 \end{bmatrix}$$

Yang memenuhi sifat $E = E^*$.

2. Menentukan polinomial karakteristik dari matriks

$$\det(\lambda I - E) = 0$$

maka

$$\det \left(\begin{bmatrix} \lambda-1 & 0 & 2i & 0 & 0 & 0 \\ 0 & \lambda-2 & 0 & 0 & 2+i & 0 \\ -2i & 0 & \lambda-1 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & \lambda-1 & 0 & 0 \\ 0 & 2-i & 0 & 0 & \lambda+2 & 0 \\ 0 & 0 & -1-i & 0 & 0 & \lambda-1 \end{bmatrix} \right) = 0$$

Sehingga polinomial karakteristik E adalah

$$\lambda^6 - 4\lambda^5 - 19\lambda^4 + 64\lambda^3 + 75\lambda^2 - 252\lambda + 135 = 0.$$

3. Menentukan nilai eigen dan vektor eigen dari matriks E ,

Dengan memfaktorkan polinomial dari matriks sebagai berikut:

$$\lambda^6 - 4\lambda^5 - 19\lambda^4 + 64\lambda^3 + 75\lambda^2 - 252\lambda + 135 = 0$$

$$(\lambda - 5)(\lambda - 3)(\lambda + 3)(\lambda - 1)^2 = 0$$

Sehingga nilai eigen dari matriks adalah $\lambda = 5$, $\lambda = 3$, $\lambda = -3$, dan $\lambda = 1$.

Secara definisi,

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{bmatrix}$$

merupakan sebuah vektor eigen dari E yang diasosiasikan dengan λ jika dan hanya jika $\mathbf{x}$ adalah solusi nontrivial bagi

$$\begin{bmatrix} \lambda-1 & 0 & 2i & 0 & 0 & 0 \\ 0 & \lambda-2 & 0 & 0 & 2+i & 0 \\ -2i & 0 & \lambda-1 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & \lambda-1 & 0 & 0 \\ 0 & 2-i & 0 & 0 & \lambda+2 & 0 \\ 0 & 0 & -1-i & 0 & 0 & \lambda-1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (4.5)$$

Untuk menentukan vektor-vektor eigen yang diasosiasikan dengan λ dengan mensubstitusikan nilai eigen ke persamaan (4.5).

Untuk $\lambda = 5$ berdasarkan persamaan (4.5) maka diperoleh:

$$\begin{bmatrix} 4 & 0 & 2i & 0 & 0 & 0 \\ 0 & 3 & 0 & 0 & 2+i & 0 \\ -2i & 0 & 4 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & 4 & 0 & 0 \\ 0 & 2-i & 0 & 0 & 7 & 0 \\ 0 & 0 & -1-i & 0 & 0 & 4 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}$$

dengan menggunakan operasi baris elementer maka diperoleh vektor eigen

untuk $\lambda = 5$ adalah

$$\mathbf{v}_1 = \begin{bmatrix} -1-i \\ 0 \\ -2+2i \\ 1-2i \\ 0 \\ 1 \end{bmatrix},$$

Dengan cara yang sama pada $\lambda = 3$, $\lambda = -3$, dan $\lambda = 1$ maka diperoleh

$$\mathbf{v}_2 = \begin{bmatrix} 0 \\ 2+i \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}, \quad \mathbf{v}_3 = \begin{bmatrix} 0 \\ \frac{-2-i}{5} \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}, \quad \mathbf{v}_4 = \begin{bmatrix} -1-i \\ 0 \\ 2-2i \\ 1-2i \\ 0 \\ 1 \end{bmatrix}, \quad \mathbf{v}_5 = \begin{bmatrix} \frac{1+i}{2} \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}, \quad \text{dan} \quad \mathbf{v}_6 = \begin{bmatrix} \frac{-1+3i}{2} \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}$$

4. Menerapkan proses Gram-Schmidt untuk mendapatkan basis ortonormal.

Berdasarkan vektor-vektor eigen yang telah diperoleh maka akan ditentukan basis orthogonal dengan menggunakan proses Gram-Schmidt sebagai berikut:

$$\mathbf{w}_1 = \mathbf{u}_1 = \begin{bmatrix} -1-i \\ 0 \\ -2+2i \\ 1-2i \\ 0 \\ 1 \end{bmatrix}$$

$$\mathbf{w}_2 = \mathbf{u}_2 - \frac{\langle \mathbf{u}_2, \mathbf{w}_1 \rangle}{\langle \mathbf{w}_1, \mathbf{w}_1 \rangle} \mathbf{w}_1 = \begin{bmatrix} 0 \\ 2+i \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}$$

$$\mathbf{w}_3 = \mathbf{u}_3 - \frac{\langle \mathbf{u}_3, \mathbf{w}_1 \rangle}{\langle \mathbf{w}_1, \mathbf{w}_1 \rangle} \mathbf{w}_1 - \frac{\langle \mathbf{u}_3, \mathbf{w}_2 \rangle}{\langle \mathbf{w}_2, \mathbf{w}_2 \rangle} \mathbf{w}_2 = \begin{bmatrix} 0 \\ \frac{-2-i}{5} \\ 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}$$

$$\mathbf{w}_4 = \mathbf{u}_4 - \frac{\langle \mathbf{u}_4, \mathbf{w}_1 \rangle}{\langle \mathbf{w}_1, \mathbf{w}_1 \rangle} \mathbf{w}_1 - \frac{\langle \mathbf{u}_4, \mathbf{w}_2 \rangle}{\langle \mathbf{w}_2, \mathbf{w}_2 \rangle} \mathbf{w}_2 - \frac{\langle \mathbf{u}_4, \mathbf{w}_3 \rangle}{\langle \mathbf{w}_3, \mathbf{w}_3 \rangle} \mathbf{w}_3 = \begin{bmatrix} -1-i \\ 0 \\ 2-2i \\ 1-2i \\ 0 \\ 1 \end{bmatrix}$$

$$w_5 = u_5 - \frac{\langle u_5, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 - \frac{\langle u_5, w_2 \rangle}{\langle w_2, w_2 \rangle} w_2 - \frac{\langle u_5, w_3 \rangle}{\langle w_3, w_3 \rangle} w_3 - \frac{\langle u_5, w_4 \rangle}{\langle w_4, w_4 \rangle} w_4 = \begin{bmatrix} 1+i \\ 2 \\ 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}$$

$$w_6 = u_6 - \frac{\langle u_6, w_1 \rangle}{\langle w_1, w_1 \rangle} w_1 - \frac{\langle u_6, w_2 \rangle}{\langle w_2, w_2 \rangle} w_2 - \frac{\langle u_6, w_3 \rangle}{\langle w_3, w_3 \rangle} w_3 - \frac{\langle u_6, w_4 \rangle}{\langle w_4, w_4 \rangle} w_4 - \frac{\langle u_6, w_5 \rangle}{\langle w_5, w_5 \rangle} w_5 = \begin{bmatrix} -1+3i \\ 3 \\ 0 \\ 0 \\ 1 \\ 0 \\ -1-2i \\ 3 \end{bmatrix}$$

- a. Menentukan kolom pertama q_1 dengan mencari $\|w_1\|$ terlebih dahulu sebagai berikut:

$$\|w_1\| = \sqrt{|-1-i|^2 + |0|^2 + |-2+2i|^2 + |1-2i|^2 + |0|^2 + |1|^2} = 4$$

$$\text{Sehingga diperoleh } q_1 = \frac{w_1}{\|w_1\|} = \begin{bmatrix} \frac{-1-i}{4} \\ 0 \\ \frac{-1+i}{4} \\ \frac{2}{4} \\ \frac{1-2i}{4} \\ 0 \\ \frac{1}{4} \\ 0 \end{bmatrix}$$

Dengan cara yang sama pada w_2, w_3, w_4, w_5 dan w_6 maka diperoleh

$$q_2 = \begin{bmatrix} 0 \\ \frac{2+i}{\sqrt{6}} \\ 0 \\ 0 \\ \frac{1}{\sqrt{6}} \\ 0 \end{bmatrix}, q_3 = \begin{bmatrix} 0 \\ \frac{-2-i}{\sqrt{30}} \\ 0 \\ 0 \\ \frac{5}{\sqrt{30}} \\ 0 \end{bmatrix}, q_4 = \begin{bmatrix} \frac{-1-i}{4} \\ 0 \\ \frac{1-i}{4} \\ \frac{2}{4} \\ \frac{1-2i}{4} \\ 0 \\ \frac{1}{4} \\ 0 \end{bmatrix}, q_5 = \frac{w_5}{\|w_5\|} = \begin{bmatrix} \frac{1+i}{\sqrt{6}} \\ 0 \\ 0 \\ 0 \\ 0 \\ \frac{2}{\sqrt{6}} \end{bmatrix}, \text{ dan } q_6 = \frac{w_6}{\|w_6\|} = \begin{bmatrix} \frac{-1+3i}{2\sqrt{6}} \\ \frac{3}{2\sqrt{6}} \\ 0 \\ \frac{3}{2\sqrt{6}} \\ 0 \\ \frac{-1-2i}{2\sqrt{6}} \end{bmatrix}$$

5. Membentuk matriks Q yang kolom-kolomnya adalah vektor-vektor basis yang dibangun dengan proses Gram-schmidt maka diperoleh matriks Q yang mendiagonalisasikan matriks E sebagai berikut:

$$Q = \begin{bmatrix} \frac{-1-i}{4} & 0 & 0 & \frac{-1-i}{4} & \frac{1+i}{\sqrt{6}} & \frac{-1+3i}{2\sqrt{6}} \\ 0 & \frac{2+i}{\sqrt{6}} & \frac{-2-i}{\sqrt{30}} & 0 & 0 & 0 \\ \frac{-1+i}{2} & 0 & 0 & \frac{1-i}{2} & 0 & 0 \\ \frac{1-2i}{4} & 0 & 0 & \frac{1-2i}{4} & 0 & \frac{3}{2\sqrt{6}} \\ 0 & \frac{1}{\sqrt{6}} & \frac{5}{\sqrt{30}} & 0 & 0 & 0 \\ \frac{1}{4} & 0 & 0 & \frac{1}{4} & \frac{2}{\sqrt{6}} & \frac{-1-2i}{2\sqrt{6}} \end{bmatrix}$$

6. Membuktikan Q dengan menunjukkan $Q^{-1}EQ = D$ adalah matriks diagonal, dimana $Q^{-1} = \bar{Q}^T = Q^*$ karena P adalah matriks uniter.

$$\bar{Q}^T = \begin{bmatrix} \frac{-1+i}{4} & 0 & \frac{-1-i}{2} & \frac{1+2i}{4} & 0 & \frac{1}{4} \\ 0 & \frac{2-i}{\sqrt{6}} & 0 & 0 & \frac{1}{\sqrt{6}} & 0 \\ 0 & \frac{-2+i}{\sqrt{30}} & 0 & 0 & \frac{5}{\sqrt{30}} & 0 \\ \frac{-1+i}{4} & 0 & \frac{1+i}{2} & \frac{1+2i}{4} & 0 & \frac{1}{4} \\ \frac{1-i}{\sqrt{6}} & 0 & 0 & 0 & 0 & \frac{2}{\sqrt{6}} \\ \frac{-1-3i}{2\sqrt{6}} & 0 & 0 & \frac{3}{2\sqrt{6}} & 0 & \frac{-1+2i}{2\sqrt{6}} \end{bmatrix} = Q^{-1}$$

Sehingga

$$Q^{-1}EQ = \begin{bmatrix} 5 & 0 & 0 & 0 & 0 & 0 \\ 0 & 3 & 0 & 0 & 0 & 0 \\ 0 & 0 & -3 & 0 & 0 & 0 \\ 0 & 0 & 0 & -3 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} = D$$

Berdasarkan pembuktian di atas maka jelaslah bahwa matriks E dapat didiagonalkan secara uniter oleh matriks Q .

4.2 Menghitung Matriks A^n , dimana $n \in \mathbb{Z}^+$

Pendiagonalan suatu matriks *Hermite* sangatlah diperlukan terutama saat kita menghitung matriks *Hermite* A^n karena dengan proses pendagonalan maka untuk menghitung matriks *Hermite* A^n akan relatif lebih singkat dari pada menghitung secara langsung yang tentunya akan memakan waktu yang sangat lama apalagi jika n merupakan bilangan bulat positif yang cukup besar.

Matriks A^n ini digunakan sebagai kunci penyandi untuk pengamanan pesan rahasia yang akan dilakukan pada langkah selanjutnya. Pada kesempatan ini penulis memilih matriks A^2 untuk ordo 5×5 dan E^3 untuk ordo 6×6 sebagai kunci penyandinya.

Kemudian untuk menghitung matriks A^n kita gunakan hasil dari perhitungan $P^{-1}AP$. Karena P adalah matriks uniter maka $P^* = P^{-1}$, jadi $P^{-1}AP = P^*AP$. Kita misalkan $P^{-1}AP = P^*AP = D$ dimana D adalah suatu matriks diagonal. Kalikan kedua ruas dengan P dari kiri dan P^{-1} dari kanan diperoleh:

$$A^n = PD^nP^{-1}$$

- Untuk A^2 , maka berdasarkan Contoh 1. diperoleh sebagai berikut:

$$A^2 = PD^2P^{-1} = \begin{bmatrix} 3 & 4-2i & 3-3i & 0 & 0 \\ 4+2i & 11 & 9-3i & 0 & 0 \\ 3+3i & 9+3i & 16 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- Untuk E^3 , berdasarkan pada Contoh 2. maka akan diperoleh sebagai berikut:

$$E^3 = QD^3Q^{-1} = \begin{bmatrix} 13 & 0 & 38i & 6-18i & 0 & -6-6i \\ 0 & 18 & 0 & 0 & 18+9i & 0 \\ -38i & 0 & 49 & -57-19i & 0 & -19+19i \\ 6+18i & 0 & -57+19i & 31 & 0 & 6-12i \\ 0 & 18-9i & 0 & 0 & -18 & 0 \\ -6+6i & 0 & -19-19i & 6+12i & 0 & 7 \end{bmatrix}$$

4.3 Pengamanan Pesan Rahasia Menggunakan Diagonalisasi Matrik *Hermite*

Pengamanan pesan rahasia menggunakan diagonalisasi matriks *Hermite* merupakan pengembangan dari *Hill Cipher*. Pada penelitian ini banyaknya karakter yang digunakan sebanyak 97 karakter yang terdiri dari huruf besar A-Z sebanyak 26 karakter, huruf kecil a-z sebanyak 26 karakter, angka 0-9 sebanyak 10 karakter, tanda

baca sebanyak 32 karakter, spasi, delete (sama seperti spasi) dan 1 karakter tambahan. Jika terdapat kekurangan pada vektor plaintext terakhir maka ditambahkan karakter “boneka” yaitu spasi. Di bawah ini disajikan tabel konversi karakter yang digunakan dalam pengamanan pesan rahasia pada penelitian ini.

Tabel 1. Konversi karakter dalam pengamanan pesan rahasia

Sp	!	“	#	\$	%	&	`	(	)	*	+	,	-	.	/
0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	1	2	3	4	5	6	7	8	9	:	;	<	=	>	?
16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
@	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
32	33	34	35	36	37	38	39	40	41	42	43	44	45	46	47
P	Q	R	S	T	U	V	W	X	Y	Z	[	\	]	^	—
48	49	50	51	52	53	54	55	56	57	58	59	60	61	62	63
‘	a	b	c	d	e	f	g	h	i	j	k	l	m	n	o
64	65	66	67	68	69	70	71	72	73	74	75	76	77	78	79
p	q	r	s	t	u	v	w	x	y	z	{		}	~	Del
80	81	82	83	84	85	86	87	88	89	90	91	92	93	94	95
■															
96															

Keterangan:

Sp = Spasi

Del = Delete

Secara rinci langkah-langkah pengamanan pesan rahasia menggunakan diagonalisasi matriks *Hermite* melalui contoh sebagai berikut:

Contoh 3. Pengamanan Pesan Rahasia Menggunakan Diagonalisasi Matriks *Hermite* Ordo

5×5

1. Proses enkripsi pesan rahasia

- Diketahui *plaintext* yang akan dienkripsi adalah

Matematika FMIPA Universitas Lampung 2017/2018

Berdasarkan Contoh 4.1 dengan matriks A^2 sebagai kunci penyandinya yaitu

$$\begin{bmatrix} 1 & 0 & 1-i & 0 & 0 \\ 0 & 1 & 3-i & 0 & 0 \\ 1+i & 3+i & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}^2$$

- b. Menghitung matriks A^2 menggunakan diagonalisasi matriks *Hermite* berdasarkan Contoh 1. diperoleh

$$A^2 = \begin{bmatrix} 3 & 4-2i & 3-3i & 0 & 0 \\ 4+2i & 11 & 9-3i & 0 & 0 \\ 3+3i & 9+3i & 16 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- c. Mentranformasikan matriks $A^2 = [a_{ij}]$ ke dalam matriks real $B = [b_{ij}]$, dimana $b_{ij} = |a_{ij}|^2$ dan melakukan operasi mod 97 pada matriks B diperoleh sebagaimana berikut:

$$B = \begin{bmatrix} |3|^2 & |4-2i|^2 & |3-3i|^2 & 0 & 0 \\ |4+2i|^2 & |11|^2 & |9-3i|^2 & 0 & 0 \\ |3+3i|^2 & |9+3i|^2 & |16|^2 & 0 & 0 \\ 0 & 0 & 0 & |1|^2 & 0 \\ 0 & 0 & 0 & 0 & |1|^2 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- d. Mengelompokkan karakter-karakter *plaintext* yang berurutan ke dalam vektor *plaintext* 5×1 , kemudian konversikan karakter tersebut dengan nilai numeriknya, diperoleh sebagaimana berikut:

$$\begin{bmatrix} M \\ a \\ t \\ e \\ m \end{bmatrix} = \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \end{bmatrix}, \begin{bmatrix} a \\ t \\ i \\ k \\ a \end{bmatrix} = \begin{bmatrix} 65 \\ 84 \\ 73 \\ 75 \\ 65 \end{bmatrix}, \begin{bmatrix} Sp \\ F \\ M \\ I \\ P \end{bmatrix} = \begin{bmatrix} 0 \\ 38 \\ 45 \\ 41 \\ 48 \end{bmatrix}, \begin{bmatrix} A \\ Sp \\ U \\ n \\ i \end{bmatrix} = \begin{bmatrix} 33 \\ 0 \\ 53 \\ 78 \\ 73 \end{bmatrix}, \begin{bmatrix} v \\ e \\ r \\ s \\ i \end{bmatrix} = \begin{bmatrix} 86 \\ 69 \\ 82 \\ 83 \\ 73 \end{bmatrix}, \begin{bmatrix} t \\ a \\ s \\ Sp \\ L \end{bmatrix} = \begin{bmatrix} 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix}, \begin{bmatrix} a \\ m \\ p \\ u \\ n \end{bmatrix} = \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \end{bmatrix},$$

$$\begin{bmatrix} g \\ Sp \\ 2 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 71 \\ 0 \\ 18 \\ 16 \\ 17 \end{bmatrix}, \begin{bmatrix} 7 \\ / \\ 2 \\ 1 \end{bmatrix} = \begin{bmatrix} 23 \\ 15 \\ 18 \\ 16 \\ 17 \end{bmatrix}, \text{ dan } \begin{bmatrix} 8 \\ Sp \\ Sp \\ Sp \\ Sp \end{bmatrix} = \begin{bmatrix} 24 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}.$$

- e. Mengalikan matriks B dengan setiap vektor *plaintext* dan melakukan operasi

modular serta konversikan dengan karakter yang setara, diperoleh sebagaimana berikut:

$$\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 3217 \\ 10020 \\ 11969 \\ 69 \\ 77 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 16 \\ 29 \\ 34 \\ 69 \\ 77 \end{bmatrix} = \begin{bmatrix} 0 \\ B \\ e \\ m \end{bmatrix}$$

$$\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 65 \\ 84 \\ 73 \\ 75 \\ 65 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 3579 \\ 9886 \\ 13256 \\ 75 \\ 65 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 87 \\ 89 \\ 64 \\ 75 \\ 65 \end{bmatrix} = \begin{bmatrix} w \\ y \\ ' \\ k \\ a \end{bmatrix}$$

$$\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 38 \\ 45 \\ 41 \\ 48 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 1570 \\ 4962 \\ 6210 \\ 41 \\ 48 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 18 \\ 15 \\ 2 \\ 41 \\ 48 \end{bmatrix} = \begin{bmatrix} 2 \\ / \\ " \\ I \\ P \end{bmatrix}$$

$$\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 33 \\ 0 \\ 53 \\ 78 \\ 73 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 1251 \\ 5430 \\ 3880 \\ 78 \\ 73 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 87 \\ 95 \\ 0 \\ 78 \\ 73 \end{bmatrix} = \begin{bmatrix} w \\ Del \\ Sp \\ n \\ i \end{bmatrix}$$

$$\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 86 \\ 69 \\ 82 \\ 83 \\ 73 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 3630 \\ 10756 \\ 12842 \\ 83 \\ 73 \end{bmatrix} (mod\ 97) = \begin{bmatrix} 41 \\ 86 \\ 38 \\ 83 \\ 73 \end{bmatrix} = \begin{bmatrix} l \\ v \\ F \\ s \\ i \end{bmatrix}$$

$$\begin{aligned}
 &\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix} \pmod{97} = \begin{bmatrix} 3550 \\ 10710 \\ 12508 \\ 0 \\ 44 \end{bmatrix} \pmod{97} = \begin{bmatrix} 58 \\ 40 \\ 92 \\ 0 \\ 44 \end{bmatrix} = \begin{bmatrix} Z \\ H \\ | \\ \text{Sp} \\ L \end{bmatrix} \\
 &\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \end{bmatrix} \pmod{97} = \begin{bmatrix} 3565 \\ 10348 \\ 13060 \\ 85 \\ 78 \end{bmatrix} \pmod{97} = \begin{bmatrix} 73 \\ 66 \\ 62 \\ 85 \\ 78 \end{bmatrix} = \begin{bmatrix} i \\ b \\ ^ \\ u \\ n \end{bmatrix} \\
 &\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 71 \\ 0 \\ 18 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 963 \\ 3040 \\ 2394 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 90 \\ 33 \\ 66 \\ 16 \\ 17 \end{bmatrix} = \begin{bmatrix} z \\ A \\ b \\ 0 \\ 1 \end{bmatrix} \\
 &\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 23 \\ 15 \\ 18 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 831 \\ 2440 \\ 2880 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 55 \\ 15 \\ 67 \\ 16 \\ 17 \end{bmatrix} = \begin{bmatrix} W \\ / \\ c \\ 0 \\ 1 \end{bmatrix} \\
 &\begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 24 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \pmod{97} = \begin{bmatrix} 216 \\ 480 \\ 432 \\ 0 \\ 0 \end{bmatrix} \pmod{97} = \begin{bmatrix} 22 \\ 92 \\ 44 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 6 \\ | \\ L \\ \text{Sp} \\ \text{Sp} \end{bmatrix}
 \end{aligned}$$

- f. Diperoleh pesan rahasia yang telah dienskripsi dari *plaintext*

Matematika FMIPA Universitas Lampung 2017/2018

menjadi *ciphertext*

0=Bemwy`ka2/"IPw~~Del~~SpniIvFsiZH|SpLib^unzAb01W/c016|LSpSp

atau dengan mengubah ~~Del~~ dan Sp menjadi spasi maka *ciphertext* yang diperoleh adalah

0=Bemwy`ka2/"IPw niIvFsiZH| Lib^unzAb01W/c016|L .

2. Proses deskripsi pesan rahasia

- a. Diketahui *plaintext* yang akan dienskripsi adalah

0=Bemwy`ka2/"IPw~~Del~~SpniIvFsiZH|SpLib^unzAb01W/c016|LSpSp

Berdasarkan Contoh 4.1 dengan matriks A^2 sebagai kunci penyandinya yaitu

$$\begin{bmatrix} 1 & 0 & 1-i & 0 & 0 \\ 0 & 1 & 3-i & 0 & 0 \\ 1+i & 3+i & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- b. Menghitung matriks A^2 menggunakan diagonalisasi matriks *Hermite* berdasarkan Contoh 1. diperoleh

$$A^2 = \begin{bmatrix} 3 & 4-2i & 3-3i & 0 & 0 \\ 4+2i & 11 & 9-3i & 0 & 0 \\ 3+3i & 9+3i & 16 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- c. Mentranformasikan matriks $A^2 = [a_{ij}]$ ke dalam matriks real $B = [b_{ij}]$, dimana $b_{ij} = |a_{ij}|^2$ dan melakukan operasi mod 97 pada matriks B diperoleh sebagaimana berikut:

$$B = \begin{bmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

- d. Mencari invers dari matriks B didefinisikan $B^{-1} = C = [c_{ij}]$ dan melakukan operasi modular pada matriks C .

Matriks B memiliki invers jika dan hanya jika determinannya tidak nol. Namun karena matriks bekerja pada $\mathbb{Z}_{97}$, maka matriks B memiliki invers modulo 97 jika dan hanya jika $\text{FPB}(\det(B), 97)=1$

$$\begin{aligned} \det(B)(\text{mod } 97) &= \det \begin{pmatrix} 9 & 20 & 18 & 0 & 0 \\ 20 & 34 & 90 & 0 & 0 \\ 18 & 90 & 62 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} (\text{mod } 97) \\ &= -27284(\text{mod } 97) = 70 \end{aligned}$$

dan

$$(\det(B)(\text{mod } 97))^{-1} = 70^{-1}(\text{mod } 97) = 79$$

Karena $\text{FPB}(70, 97) = 1$, maka invers modulo 97 dari matriks B diperoleh sebagai berikut:

$$\begin{aligned} C &= \left(\frac{1}{\det(B)} \cdot \text{Adj}(B) \right) (\text{mod } 97) \\ &= \frac{1}{\det(B)(\text{mod } 97)} \cdot \text{Adj}(B)(\text{mod } 97) \\ &= (\det(B)(\text{mod } 97))^{-1} (\text{mod } 97) \cdot \text{Adj}(B)(\text{mod } 97) \\ &= 79 \begin{bmatrix} -6612 & 380 & 1368 & 0 & 0 \\ 380 & 234 & -450 & 0 & 0 \\ 1368 & -450 & -184 & 0 & 0 \\ 0 & 0 & 0 & -27284 & 0 \\ 0 & 0 & 0 & 0 & -27284 \end{bmatrix} (\text{mod } 97) \\ &= \begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \end{aligned}$$

- e. Mengelompokkan karakter-karakter *ciphertext* yang berurutan ke dalam vektor *ciphertext* 5×1 , kemudian konversikan karakter tersebut dengan nilai numeriknya, diperoleh sebagai berikut:

$$\begin{aligned} \begin{bmatrix} 0 \\ = \\ B \\ e \\ m \end{bmatrix} &= \begin{bmatrix} 16 \\ 29 \\ 34 \\ 69 \\ 77 \end{bmatrix}, \begin{bmatrix} w \\ y \\ ' \\ k \\ a \end{bmatrix} = \begin{bmatrix} 87 \\ 89 \\ 64 \\ 75 \\ 65 \end{bmatrix}, \begin{bmatrix} 2 \\ / \\ " \\ 1 \\ p \end{bmatrix} = \begin{bmatrix} 18 \\ 15 \\ 2 \\ 41 \\ 48 \end{bmatrix}, \begin{bmatrix} w \\ Del \\ Sp \\ n \\ i \end{bmatrix} = \begin{bmatrix} 87 \\ 95 \\ 0 \\ 78 \\ 73 \end{bmatrix}, \begin{bmatrix} l \\ v \\ f \\ s \\ i \end{bmatrix} = \begin{bmatrix} 41 \\ 86 \\ 38 \\ 83 \\ 73 \end{bmatrix}, \begin{bmatrix} Z \\ H \\ | \\ Sp \\ L \end{bmatrix} = \begin{bmatrix} 58 \\ 40 \\ 92 \\ 0 \\ 44 \end{bmatrix}, \begin{bmatrix} i \\ b \\ ^ \\ u \\ n \end{bmatrix} = \begin{bmatrix} 73 \\ 66 \\ 62 \\ 85 \\ 78 \end{bmatrix}, \begin{bmatrix} z \\ A \\ b \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 90 \\ 33 \\ 66 \\ 16 \\ 17 \end{bmatrix} \\ \begin{bmatrix} W \\ / \\ c \\ 0 \\ 1 \end{bmatrix} &= \begin{bmatrix} 55 \\ 15 \\ 67 \\ 16 \\ 17 \end{bmatrix}, \text{ dan } \begin{bmatrix} 6 \\ l \\ L \\ Sp \\ Sp \end{bmatrix} = \begin{bmatrix} 22 \\ 92 \\ 44 \\ 0 \\ 0 \end{bmatrix} \end{aligned}$$

- f. Mengalikan matriks C dengan setiap vektor *ciphertext* dan melakukan operasi modular serta konversikan angka tersebut dengan karakter yang setara, diperoleh sebagaimana berikut:

$$\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 16 \\ 29 \\ 34 \\ 69 \\ 77 \end{bmatrix} (\text{mod } 97) = \begin{bmatrix} 3343 \\ 4042 \\ 2121 \\ 69 \\ 77 \end{bmatrix} (\text{mod } 97) = \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \end{bmatrix} = \begin{bmatrix} M \\ a \\ t \\ e \\ m \end{bmatrix}$$

$$\begin{aligned}
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 87 \\ 89 \\ 64 \\ 75 \\ 65 \end{bmatrix} \pmod{97} = \begin{bmatrix} 13257 \\ 12209 \\ 6475 \\ 75 \\ 65 \end{bmatrix} \pmod{97} = \begin{bmatrix} 65 \\ 84 \\ 73 \\ 75 \\ 65 \end{bmatrix} = \begin{bmatrix} a \\ t \\ i \\ k \\ a \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 18 \\ 15 \\ 2 \\ 41 \\ 48 \end{bmatrix} \pmod{97} = \begin{bmatrix} 2425 \\ 1784 \\ 1015 \\ 41 \\ 48 \end{bmatrix} \pmod{97} = \begin{bmatrix} 0 \\ 38 \\ 45 \\ 41 \\ 48 \end{bmatrix} = \begin{bmatrix} Sp \\ F \\ M \\ I \\ P \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 87 \\ 95 \\ 0 \\ 78 \\ 73 \end{bmatrix} \pmod{97} = \begin{bmatrix} 12643 \\ 9409 \\ 5873 \\ 78 \\ 73 \end{bmatrix} \pmod{97} = \begin{bmatrix} 33 \\ 0 \\ 53 \\ 78 \\ 73 \end{bmatrix} = \begin{bmatrix} A \\ Sp \\ U \\ n \\ i \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 41 \\ 86 \\ 38 \\ 83 \\ 73 \end{bmatrix} \pmod{97} = \begin{bmatrix} 8428 \\ 8605 \\ 5320 \\ 83 \\ 73 \end{bmatrix} \pmod{97} = \begin{bmatrix} 86 \\ 69 \\ 82 \\ 83 \\ 73 \end{bmatrix} = \begin{bmatrix} v \\ e \\ r \\ s \\ i \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 58 \\ 40 \\ 92 \\ 0 \\ 44 \end{bmatrix} \pmod{97} = \begin{bmatrix} 8620 \\ 9474 \\ 4060 \\ 0 \\ 44 \end{bmatrix} \pmod{97} = \begin{bmatrix} 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix} = \begin{bmatrix} t \\ a \\ s \\ Sp \\ L \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 73 \\ 66 \\ 62 \\ 85 \\ 78 \end{bmatrix} \pmod{97} = \begin{bmatrix} 10832 \\ 10165 \\ 5124 \\ 85 \\ 78 \end{bmatrix} \pmod{97} = \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \end{bmatrix} = \begin{bmatrix} a \\ m \\ p \\ u \\ n \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 90 \\ 33 \\ 66 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 10935 \\ 9312 \\ 3801 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 71 \\ 0 \\ 18 \\ 16 \\ 17 \end{bmatrix} = \begin{bmatrix} g \\ Sp \\ 2 \\ 0 \\ 1 \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 55 \\ 15 \\ 67 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 6813 \\ 6708 \\ 2443 \\ 16 \\ 17 \end{bmatrix} \pmod{97} = \begin{bmatrix} 23 \\ 15 \\ 18 \\ 16 \\ 17 \end{bmatrix} = \begin{bmatrix} 7 \\ / \\ 2 \\ 0 \\ 1 \end{bmatrix} \\
 &\begin{bmatrix} 94 & 47 & 14 & 0 & 0 \\ 47 & 56 & 49 & 0 & 0 \\ 14 & 49 & 14 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 22 \\ 92 \\ 44 \\ 0 \\ 0 \end{bmatrix} \pmod{97} = \begin{bmatrix} 7008 \\ 8342 \\ 5432 \\ 0 \\ 0 \end{bmatrix} \pmod{97} = \begin{bmatrix} 24 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 8 \\ Sp \\ Sp \\ Sp \\ Sp \end{bmatrix}
 \end{aligned}$$

g. Diperoleh pesan rahasia yang telah dideskripsi dari *ciphertext*

0=Bemwy`ka2/"IPw niIvFsiZH| Lib^unzAb01W/c016|L

menjadi *plaintext*

MatematikaSpFMIPASpUniversitasSpLampungSp2017/2018SpSpSpSp

atau dengan mengubah Sp menjadi spasi maka *plaintext* yang diperoleh adalah

Matematika FMIPA Universitas Lampung 2017/2018.

Contoh 4. Pengamanan Pesan Rahasia Menggunakan Diagonalisasi Matriks Hermite Ordo

6×6

1. Proses enkripsi pesan rahasia

- a. Diketahui *plaintext* yang akan dienkripsi adalah

Matematika FMIPA Universitas Lampung 2017/2018

Berdasarkan Contoh 4.2 dengan matriks E^3 sebagai kunci penyandinya yaitu

$$\begin{bmatrix} 1 & 0 & 2i & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 & 2+i & 0 \\ -2i & 0 & 1 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & 1 & 0 & 0 \\ 0 & 2-i & 0 & 0 & -2 & 0 \\ 0 & 0 & 0 & -1-i & 0 & 1 \end{bmatrix}^3$$

- b. Menghitung matriks E^3 menggunakan diagonalisasi matriks *Hermite* berdasarkan Contoh 2. diperoleh

$$E^3 = \begin{bmatrix} 13 & 0 & 38i & 6-18i & 0 & -6-6i \\ 0 & 18 & 0 & 0 & 18+9i & 0 \\ -38i & 0 & 49 & -57-19i & 0 & -19+19i \\ 6+18i & 0 & -57+19i & 31 & 0 & 6-12i \\ 0 & 18-9i & 0 & 0 & -18 & 0 \\ -6+6i & 0 & -19-19i & 6+12i & 0 & 7 \end{bmatrix}$$

- c. Mentranformasikan matriks $E^3 = [e_{ij}]$ ke dalam matriks real $F = [f_{ij}]$, dimana $f_{ij} = |e_{ij}|^2$ dan melakukan operasi mod 97 pada matriks F diperoleh sebagaimana berikut:

$$F =$$

$$\begin{bmatrix} |13|^2 & 0 & |38i|^2 & |6-18i|^2 & 0 & |-6-6i|^2 \\ 0 & |18|^2 & 0 & 0 & |18+9i|^2 & 0 \\ |-38i|^2 & 0 & |49|^2 & |-57-19i|^2 & 0 & |-19+19i|^2 \\ |6+18i|^2 & 0 & |-57+19i|^2 & |31|^2 & 0 & |6-12i|^2 \\ 0 & |18-9i|^2 & 0 & 0 & |-18|^2 & 0 \\ |-6+6i|^2 & 0 & |-19-19i|^2 & |6+12i|^2 & 0 & |7|^2 \end{bmatrix} \pmod{97}$$

$$= \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix}$$

- d. Mengelompokkan karakter-karakter *plaintext* yang berurutan ke dalam vektor *plaintext* 6×1 , kemudian konversikan karakter tersebut dengan nilai numeriknya, diperoleh sebagaimana berikut:

$$\begin{bmatrix} M \\ a \\ t \\ e \\ m \\ a \end{bmatrix} = \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \\ 65 \end{bmatrix}, \begin{bmatrix} t \\ i \\ k \\ a \\ Sp \\ F \end{bmatrix} = \begin{bmatrix} 84 \\ 73 \\ 75 \\ 65 \\ 0 \\ 38 \end{bmatrix}, \begin{bmatrix} M \\ I \\ P \\ A \\ Sp \\ U \end{bmatrix} = \begin{bmatrix} 45 \\ 41 \\ 48 \\ 33 \\ 0 \\ 53 \end{bmatrix}, \begin{bmatrix} n \\ i \\ v \\ e \\ r \\ s \end{bmatrix} = \begin{bmatrix} 78 \\ 73 \\ 86 \\ 69 \\ 82 \\ 83 \end{bmatrix}, \begin{bmatrix} i \\ t \\ a \\ s \\ Sp \\ L \end{bmatrix} = \begin{bmatrix} 73 \\ 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix}, \begin{bmatrix} a \\ m \\ p \\ u \\ n \\ g \end{bmatrix} = \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \\ 71 \end{bmatrix}, \begin{bmatrix} Sp \\ 2 \\ 0 \\ 1 \\ 7 \\ / \end{bmatrix} = \begin{bmatrix} 0 \\ 18 \\ 16 \\ 17 \\ 23 \\ 15 \end{bmatrix}$$

$$\text{dan } \begin{bmatrix} 2 \\ 0 \\ 1 \\ 8 \\ Sp \\ Sp \end{bmatrix} = \begin{bmatrix} 18 \\ 16 \\ 17 \\ 24 \\ 0 \\ 0 \end{bmatrix}$$

- e. Mengalikan matriks B dengan setiap vektor *plaintext* dan melakukan operasi

modular serta konversikan dengan karakter yang setara, diperoleh sebagaimana berikut:

$$\begin{array}{l}
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \\ 65 \end{bmatrix} \pmod{97} = \begin{bmatrix} 19905 \\ 3454 \\ 14246 \\ 16336 \\ 3446 \\ 15764 \end{bmatrix} \pmod{97} = \begin{bmatrix} 20 \\ 59 \\ 84 \\ 40 \\ 57 \\ 50 \end{bmatrix} = \begin{bmatrix} 4 \\ [\\ t \\ H \\ Y \\ R \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 84 \\ 73 \\ 75 \\ 65 \\ 0 \\ 38 \end{bmatrix} \pmod{97} = \begin{bmatrix} 19719 \\ 2409 \\ 15698 \\ 16245 \\ 1241 \\ 16530 \end{bmatrix} \pmod{97} = \begin{bmatrix} 28 \\ 81 \\ 81 \\ 46 \\ 77 \\ 40 \end{bmatrix} = \begin{bmatrix} < \\ q \\ q \\ N \\ m \\ H \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 45 \\ 41 \\ 48 \\ 33 \\ 0 \\ 53 \end{bmatrix} \pmod{97} = \begin{bmatrix} 13461 \\ 1353 \\ 10346 \\ 11416 \\ 697 \\ 10640 \end{bmatrix} \pmod{97} = \begin{bmatrix} 75 \\ 92 \\ 64 \\ 67 \\ 18 \\ 67 \end{bmatrix} = \begin{bmatrix} k \\ | \\ ' \\ c \\ 2 \\ c \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 78 \\ 73 \\ 86 \\ 69 \\ 82 \\ 83 \end{bmatrix} \pmod{97} = \begin{bmatrix} 23749 \\ 3803 \\ 18004 \\ 20149 \\ 3947 \\ 19108 \end{bmatrix} \pmod{97} = \begin{bmatrix} 81 \\ 20 \\ 59 \\ 70 \\ 67 \\ 96 \end{bmatrix} = \begin{bmatrix} q \\ 4 \\ [\\ f \\ c \\ \blacksquare \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 73 \\ 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix} \pmod{97} = \begin{bmatrix} 19741 \\ 2772 \\ 14658 \\ 17358 \\ 1428 \\ 17096 \end{bmatrix} \pmod{97} = \begin{bmatrix} 50 \\ 56 \\ 11 \\ 92 \\ 70 \\ 24 \end{bmatrix} = \begin{bmatrix} R \\ X \\ + \\ | \\ f \\ 8 \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \\ 71 \end{bmatrix} \pmod{97} = \begin{bmatrix} 22537 \\ 3867 \\ 16268 \\ 19538 \\ 3883 \\ 18654 \end{bmatrix} \pmod{97} = \begin{bmatrix} 33 \\ 84 \\ 69 \\ 41 \\ 3 \\ 30 \end{bmatrix} = \begin{bmatrix} A \\ t \\ e \\ I \\ \# \\ > \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 0 \\ 18 \\ 16 \\ 17 \\ 23 \\ 15 \end{bmatrix} \pmod{97} = \begin{bmatrix} 3629 \\ 985 \\ 2170 \\ 3077 \\ 1065 \\ 2834 \end{bmatrix} \pmod{97} = \begin{bmatrix} 40 \\ 15 \\ 36 \\ 70 \\ 95 \\ 21 \end{bmatrix} = \begin{bmatrix} H \\ / \\ D \\ f \\ \text{Del} \\ 5 \end{bmatrix} \\
 \\
 \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix} \begin{bmatrix} 18 \\ 16 \\ 17 \\ 24 \\ 0 \\ 0 \end{bmatrix} \pmod{97} = \begin{bmatrix} 4414 \\ 528 \\ 3293 \\ 3711 \\ 272 \\ 4019 \end{bmatrix} \pmod{97} = \begin{bmatrix} 49 \\ 43 \\ 92 \\ 25 \\ 78 \\ 42 \end{bmatrix} = \begin{bmatrix} Q \\ K \\ | \\ 9 \\ n \\ J \end{bmatrix}
 \end{array}$$

- f. Diperoleh pesan rahasia yang telah dienskripsi dari *plaintext*

Matematika FMIPA Universitas Lampung 2017/2018

menjadi *ciphertext*

4[tHYR<qqNmHk|`c2cq4[fc■RX+|f8AteI#>H/Df Del 5QK|9nJ

atau dengan mengubah **Del** menjadi spasi maka *ciphertext* yang diperoleh adalah

4[tHYR<qqNmHk|`c2cq4[fc■RX+|f8AteI#>H/Df 5QK|9nJ.

2. Proses deskripsi pesan rahasia

- a. Diketahui *ciphertext* yang akan dienskripsi adalah

4[tHYR<qqNmHk]\`c2cq4[fc■RX+|f8AteI#>H/Df Del 5QK|9nJ

Berdasarkan Contoh 4.1 dengan matriks E^3 sebagai kunci penyandinya yaitu

$$\begin{bmatrix} 1 & 0 & 2i & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 & 2+i & 0 \\ -2i & 0 & 1 & -3-i & 0 & -1+i \\ 0 & 0 & -3+i & 1 & 0 & 0 \\ 0 & 2-i & 0 & 0 & -2 & 0 \\ 0 & 0 & 0 & -1-i & 0 & 1 \end{bmatrix}^3$$

- b. Menghitung matriks E^3 menggunakan diagonalisasi matriks *Hermite* berdasarkan Contoh 2. diperoleh

$$E^3 = \begin{bmatrix} 13 & 0 & 38i & 6-18i & 0 & -6-6i \\ 0 & 18 & 0 & 0 & 18+9i & 0 \\ -38i & 0 & 49 & -57-19i & 0 & -19+19i \\ 6+18i & 0 & -57+19i & 31 & 0 & 6-12i \\ 0 & 18-9i & 0 & 0 & -18 & 0 \\ -6+6i & 0 & -19-19i & 6+12i & 0 & 7 \end{bmatrix}$$

- c. Mentransformasikan matriks $E^3 = [e_{ij}]$ ke dalam matriks real $F = [f_{ij}]$, dimana $f_{ij} = |e_{ij}|^2$ dan melakukan operasi mod 97 pada matriks F diperoleh sebagaimana berikut:

$$F = \begin{bmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{bmatrix}$$

- d. Mencari invers dari matriks F didefinisikan $F^{-1} = G = [g_{ij}]$ dan melakukan operasi modular pada matriks G .

Matriks F memiliki invers jika dan hanya jika determinannya tidak nol. Namun karena matriks bekerja pada $\mathbb{Z}_{97}$, maka matriks F memiliki invers modulo 97 jika dan hanya jika $\text{FPB}(\det(F), 97)=1$

$$\begin{aligned} \det(F)(\text{mod } 97) &= \det \begin{pmatrix} 72 & 0 & 86 & 69 & 0 & 72 \\ 0 & 33 & 0 & 0 & 17 & 0 \\ 86 & 0 & 73 & 21 & 0 & 43 \\ 69 & 0 & 21 & 88 & 0 & 83 \\ 0 & 17 & 0 & 0 & 33 & 0 \\ 72 & 0 & 43 & 83 & 0 & 49 \end{pmatrix} (\text{mod } 97) \\ &= 8123587200(\text{mod } 97) = 63 \end{aligned}$$

dan

$$(\det(F)(\text{mod } 97))^{-1} = 63^{-1}(\text{mod } 97) = 77$$

Karena $\text{FPB}(63, 97) = 1$, maka invers modulo 97 dari matriks F diperoleh sebagai berikut:

$$\begin{aligned} G &= \left(\frac{1}{\det(F)} \cdot \text{Adj}(F) \right) (\text{mod } 97) \\ &= \frac{1}{\det(F)(\text{mod } 97)} \cdot \text{Adj}(F)(\text{mod } 97) \\ &= (\det(F)(\text{mod } 97))^{-1} (\text{mod } 97) \cdot \text{Adj}(F)(\text{mod } 97) \end{aligned}$$

$$= \begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix}$$

- e. Mengelompokkan karakter-karakter *ciphertext* yang berurutan ke dalam vektor *ciphertext* 6×1 , kemudian konversikan karakter tersebut dengan nilai numeriknya, diperoleh sebagai berikut:

$$\begin{bmatrix} 4 \\ [\\ t \\ H \\ Y \\ R \end{bmatrix} = \begin{bmatrix} 20 \\ 59 \\ 84 \\ 40 \\ 57 \\ 50 \end{bmatrix}, \begin{bmatrix} < \\ q \\ N \\ m \\ H \end{bmatrix} = \begin{bmatrix} 28 \\ 81 \\ 81 \\ 46 \\ 77 \\ 40 \end{bmatrix}, \begin{bmatrix} k \\ | \\ ' \\ c \\ 2 \\ c \end{bmatrix} = \begin{bmatrix} 75 \\ 92 \\ 64 \\ 67 \\ 18 \\ 67 \end{bmatrix}, \begin{bmatrix} q \\ 4 \\ [\\ f \\ c \\ \blacksquare \end{bmatrix} = \begin{bmatrix} 81 \\ 20 \\ 59 \\ 70 \\ 67 \\ 96 \end{bmatrix}, \begin{bmatrix} R \\ X \\ + \\ | \\ f \end{bmatrix} = \begin{bmatrix} 50 \\ 56 \\ 11 \\ 92 \\ 70 \\ 24 \end{bmatrix}, \begin{bmatrix} A \\ t \\ e \\ l \\ \# \\ > \end{bmatrix} = \begin{bmatrix} 33 \\ 84 \\ 69 \\ 41 \\ 3 \\ 30 \end{bmatrix}, \begin{bmatrix} H \\ / \\ D \\ f \\ \text{Del} \\ 5 \end{bmatrix} = \begin{bmatrix} 40 \\ 15 \\ 36 \\ 70 \\ 95 \\ 21 \end{bmatrix}$$

$$\text{dan} \begin{bmatrix} Q \\ K \\ | \\ 9 \\ n \\ J \end{bmatrix} = \begin{bmatrix} 49 \\ 43 \\ 92 \\ 25 \\ 78 \\ 42 \end{bmatrix}$$

- f. Mengalikan matriks G dengan setiap vektor *ciphertext* dan melakukan operasi modular serta konversikan angka tersebut dengan karakter yang setara, diperoleh sebagaimana berikut:

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 20 \\ 59 \\ 84 \\ 40 \\ 57 \\ 50 \end{bmatrix} \pmod{97} = \begin{bmatrix} 10618 \\ 7534 \\ 8620 \\ 9866 \\ 7546 \\ 13548 \end{bmatrix} \pmod{97} = \begin{bmatrix} 45 \\ 65 \\ 84 \\ 69 \\ 77 \\ 65 \end{bmatrix} = \begin{bmatrix} M \\ a \\ t \\ e \\ m \\ a \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 28 \\ 81 \\ 81 \\ 46 \\ 77 \\ 40 \end{bmatrix} \pmod{97} = \begin{bmatrix} 11336 \\ 10258 \\ 8320 \\ 10347 \\ 10282 \\ 13036 \end{bmatrix} \pmod{97} = \begin{bmatrix} 84 \\ 73 \\ 75 \\ 65 \\ 0 \\ 38 \end{bmatrix} = \begin{bmatrix} t \\ i \\ k \\ a \\ \text{Sp} \\ F \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 75 \\ 92 \\ 64 \\ 67 \\ 18 \\ 67 \end{bmatrix} \pmod{97} = \begin{bmatrix} 15468 \\ 6928 \\ 13337 \\ 17590 \\ 7372 \\ 15670 \end{bmatrix} \pmod{97} = \begin{bmatrix} 45 \\ 41 \\ 48 \\ 33 \\ 0 \\ 53 \end{bmatrix} = \begin{bmatrix} M \\ l \\ P \\ A \\ \text{Sp} \\ U \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 81 \\ 20 \\ 59 \\ 70 \\ 67 \\ 96 \end{bmatrix} \pmod{97} = \begin{bmatrix} 16180 \\ 5796 \\ 15994 \\ 20148 \\ 5514 \\ 17737 \end{bmatrix} \pmod{97} = \begin{bmatrix} 78 \\ 73 \\ 86 \\ 69 \\ 82 \\ 83 \end{bmatrix} = \begin{bmatrix} n \\ i \\ v \\ e \\ r \\ s \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 50 \\ 56 \\ 11 \\ 92 \\ 70 \\ 24 \end{bmatrix} \pmod{97} = \begin{bmatrix} 12198 \\ 8232 \\ 7146 \\ 13663 \\ 8148 \\ 9356 \end{bmatrix} \pmod{97} = \begin{bmatrix} 73 \\ 84 \\ 65 \\ 83 \\ 0 \\ 44 \end{bmatrix} = \begin{bmatrix} i \\ t \\ a \\ s \\ \text{Sp} \\ L \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 33 \\ 84 \\ 69 \\ 41 \\ 3 \\ 30 \end{bmatrix} \pmod{97} = \begin{bmatrix} 10347 \\ 5412 \\ 7355 \\ 9494 \\ 5898 \\ 11032 \end{bmatrix} \pmod{97} = \begin{bmatrix} 65 \\ 77 \\ 80 \\ 85 \\ 78 \\ 71 \end{bmatrix} = \begin{bmatrix} a \\ m \\ p \\ u \\ n \\ g \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 40 \\ 15 \\ 36 \\ 70 \\ 95 \\ 21 \end{bmatrix} \pmod{97} = \begin{bmatrix} 11155 \\ 7390 \\ 6612 \\ 11269 \\ 6910 \\ 9618 \end{bmatrix} \pmod{97} = \begin{bmatrix} 0 \\ 18 \\ 16 \\ 17 \\ 23 \\ 15 \end{bmatrix} = \begin{bmatrix} \text{Sp} \\ 2 \\ 0 \\ 1 \\ 7 \\ / \end{bmatrix}$$

$$\begin{bmatrix} 64 & 0 & 62 & 87 & 0 & 13 \\ 0 & 62 & 0 & 0 & 68 & 0 \\ 62 & 0 & 30 & 19 & 0 & 82 \\ 87 & 0 & 19 & 82 & 0 & 65 \\ 0 & 68 & 0 & 0 & 62 & 0 \\ 13 & 0 & 82 & 65 & 0 & 76 \end{bmatrix} \begin{bmatrix} 49 \\ 43 \\ 92 \\ 25 \\ 78 \\ 42 \end{bmatrix} \pmod{97} = \begin{bmatrix} 11561 \\ 7970 \\ 9717 \\ 10791 \\ 7760 \\ 12998 \end{bmatrix} \pmod{97} = \begin{bmatrix} 18 \\ 16 \\ 17 \\ 24 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 2 \\ 0 \\ 1 \\ 8 \\ \text{Sp} \\ \text{Sp} \end{bmatrix}$$

g. Diperoleh pesan rahasia yang telah dideskripsi dari *ciphertext*

4[tHYR<qqNmHk]`c2cq4[fc■RX+|f8AteI#>H/Df 5QK|9nJ

menjadi *plaintext*

MatematikaSpFMIPASpUniversitasSpLampungSp2017/2018SpSp

atau dengan mengubah Sp menjadi spasi maka *plaintext* yang diperoleh adalah

Matematika FMIPA Universitas Lampung 2017/2018.

5. SIMPULAN DAN SARAN

Dari uraian pembahasan dapat diambil kesimpulan berikut. Diagonalisasi matriks hermite secara uniter sebagai landasan membuat matriks kunci proses enkripsi dan deskripsi pesan rahasia. Algoritma yang digunakan pada proses pengamanan pesan rahasia ini adalah *Hill Cipher*. Beberapa hasil enkripsi dan deskripsi dengan ukuran matriks berbeda yang telah disajikan menghasilkan bahwa *Hill Cipher* adalah algoritma kriptografi klasik yang sangat kuat dilihat dari segi keamanannya. Semakin besar matriks kunci, semakin sulit untuk dipecahkan oleh orang lain

yang berarti semakin tinggi tingkat kemanannya. Saran dari penulis adalah sebaiknya pesan yang akan dikirim dienkripsi terlebih dahulu menggunakan proses diagonalisasi matriks *Hermite* sehingga pesan yang terkirim hanya dapat dimengerti oleh orang yang berhak menerimanya saja.

KEPUSTAKAAN

- [1] Anton,H. dan Rorres, C. (2004). *Aljabar Linear Elementer,Jilid 1*.Erlangga, Jakarta.
- [2] Anton,H. dan Rorres, C. (2006). *Aljabar Linear Elementer,Jilid 2*. Erlangga, Jakarta.
- [3] Menezes, A.J., Oorschot, P.C.V. dan Vanstone, S.A. (1996). *Handbook of Applied Cryptography*.CRC Press.
- [4] Stinson, D. R. (2006). *Cryptography Theory and Practice*.CRC Press Boca Raton, Florida.
- [5] Ariyus, D. (2008). *Pengantar Ilmu Kriptografi Teori Analisis dan Implementasi*. Andi Offset,Yogyakarta.

PEMBANDINGAN METODE RUNGE-KUTTA ORDE 4 DAN METODE ADAM-BASHFORT MOULTON DALAM PENYELESAIAN MODEL PERTUMBUHAN UANG YANG DIINVESTASIKAN

Intan Puspitasari, Agus Sutrisno, Tiryono Ruby, Muslim Ansori
Jurusan Matematika
Fakultas Matematika dan Ilmu Pengetahuan Alam
Universitas Lampung
Jl. Prof. Sumantri Brojonegoro No.1, Bandar Lampung
Telp. +62721701609, Fax. +62721702767, Website : unila.ac.id

ABSTRAK

Persamaan diferensial sering kali diterapkan pada berbagai model matematika yang menggambarkan masalah dalam kehidupan nyata, salah satunya dalam masalah finansial yaitu investasi. Investasi berupa tabungan bank dapat diaplikasikan menjadi sebuah model matematika, yaitu

$$\frac{dP(t)}{dt} = r \cdot P(t)$$

dimana $P(t)$ merupakan besarnya tabungan pada tahun ke- t (dalam rupiah), r adalah besarnya bunga, dan t adalah waktu ke- t (dalam tahun). Model tersebut dapat diselesaikan dengan dua metode yaitu metode analitik dan metode numerik. Pada umumnya, digunakan ekspansi Taylor untuk menurunkan metode numerik dari suatu model. Akan tetapi, ekspansi Taylor akan membutuhkan turunan tingkat tinggi yang menyebabkan kompleksitas perhitungan bertambah. Berbeda dengan ekspansi Taylor, skema Runge-Kutta dan Adam-Bashfort Moulton adalah alternatif dari metode numerik untuk mendapatkan konvergensi tinggi tanpa memerlukan turunan tingkat tinggi. Oleh karena itu, penelitian ini menggunakan metode Runge-Kutta orde empat dan metode Adam-Bashfort Moulton dalam penyelesaian model pertumbuhan uang yang diinvestasikan. Dari kedua metode tersebut ditentukan metode terbaik dalam mengaproksimasi nilai penyelesaian model tersebut dengan melihat nilai galat dari kedua metode tersebut. Dalam penelitian ini ditentukan 2 contoh kasus dari model pertumbuhan uang yang diinvestasikan. Selanjutnya, dicari solusi numerik berdasarkan dari kedua metode tersebut dengan menggunakan software MATLAB R2013b. Kemudian, bandingkan galat kedua solusi numerik tersebut terhadap solusi analitiknya. Setelah dilakukan tahap-tahap pada metode penelitian, dapat disimpulkan bahwa semakin kecil bunga pertahunnya, maka hasil aproksimasi semakin mendekati hasil eksaknya. Sebaliknya, semakin besar bunga pertahunnya, maka selisih antara hasil aproksimasi dan hasil eksaknya akan semakin besar. Metode Runge-Kutta Orde 4 lebih baik dalam mengaproksimasi suatu nilai pada $x(i)$ yang besar dibandingkan dengan metode Adam-Bashfort Moulton. Sebaliknya, dari kedua contoh kasus tersebut terlihat bahwa metode Adam-Bashfort Moulton lebih baik dalam mengaproksimasi suatu nilai pada $x(i)$ yang kecil dibandingkan metode Runge-Kutta Orde 4.

Kata kunci : Runge-Kutta, Adam-Bashfort Moulton

1. PENDAHULUAN

Persamaan diferensial sering kali diterapkan pada berbagai model matematika yang menggambarkan masalah dalam kehidupan nyata, salah satunya dalam bidang finansial yaitu investasi. Investasi memiliki berbagai jenis, diantaranya yaitu saham, deposito berjangka, emas, tabungan bank, dan masih banyak lagi. Namun, yang akan dibahas pada penelitian ini hanyalah investasi berupa tabungan bank. Investasi berupa tabungan bank ini dapat diaplikasikan menjadi sebuah model matematika.

Model tersebut berbentuk persamaan diferensial, yaitu

$$\frac{dP(t)}{dt} = r \cdot P(t)$$

dimana $P(t)$ merupakan besarnya tabungan pada tahun ke- t (dalam rupiah), r adalah besarnya bunga, dan t adalah tahun ke- t (dalam tahun). Model tersebut dapat diselesaikan dengan dua metode yaitu metode analitik dan metode numerik. Metode analitik disebut juga metode sejati karena memberikan solusi sejati atau solusi yang sesungguhnya, yaitu solusi yang memiliki galat sama dengan nol. Sedangkan, metode numerik adalah satu-satunya metode alternatif yang ada dalam upaya menyelesaikan persoalan-persoalan matematis dengan mengkaji parametrik dari persoalan dari medan yang bersifat sembarang.

Pada umumnya, digunakan ekspansi Taylor untuk menurunkan metode numerik dari suatu model. Akan tetapi, ekspansi Taylor akan membutuhkan turunan tingkat tinggi yang menyebabkan kompleksitas perhitungan bertambah. Berbeda dengan ekspansi Taylor, skema Runge-Kutta adalah alternatif dari metode numerik untuk mendapatkan konvergensi tinggi tanpa memerlukan turunan tingkat tinggi.

Metode Runge-Kutta yang paling mendekati konvergen ialah yang berorde empat. Metode Runge-Kutta Orde 4 merupakan metode langkah tunggal yang memiliki nilai galat terkecil, sedangkan pada metode langkah ganda dapat digunakan metode Adam-Bashfort Moulton untuk mengaproksimasikan penyelesaian model tersebut dengan galat terkecil. Oleh karena itu, penelitian ini akan menggunakan metode Runge-Kutta orde empat dan metode Adam-Bashfort Moulton dalam penyelesaian model pertumbuhan uang yang diinvestasikan. Dari kedua metode tersebut akan ditentukan metode terbaik dalam mengaproksimasi nilai penyelesaian model tersebut dengan melihat nilai galat dari kedua metode tersebut.

2. LANDASAN TEORI

2.1 Model Matematika

Model matematika suatu fenomena adalah suatu ekspresi matematika yang diturunkan dari fenomena tersebut. Ekspresi dapat berupa persamaan, sistem persamaan atau ekspresi-ekspresi matematika yang lain seperti fungsi maupun relasi. Model matematika digunakan untuk menjelaskan karakteristik fenomena yang dimodelkannya, dapat secara kualitatif dan kuantitatif [1].

2.2 Persamaan Diferensial

Persamaan diferensial adalah persamaan yang memuat variable bebas, variable tak bebas dan derivatif-derivatif dari variable tidak bebas terhadap variable bebasnya. Tingkat (orde) persamaan diferensial adalah tingkat tertinggi dari derivatif yang terdapat dalam persamaan diferensial. Derajat suatu persamaan diferensial adalah pangkat tertinggi dari derivatif tertinggi dalam persamaan diferensial [2].

2.3 Persamaan Diferensial Biasa

Persamaan diferensial biasa adalah persamaan yang memuat turunan terhadap fungsi yang memuat satu variabel bebas. Jika x adalah fungsi dari t , maka contoh persamaan diferensial biasa adalah

$$\frac{dx}{dt} = t^2 \cos x \quad (1)$$

Dimana persamaan tersebut memiliki order satu. Order dari persamaan diferensial adalah turunan tertinggi pada fungsi tak diketahui (peubah tak bebas) yang muncul dalam persamaan diferensial [3].

2.4 Persamaan Diferensial Biasa (PDB) Linier

Persamaan diferensial biasa linier memiliki bentuk umum

$$a_n(t) \frac{d^n x}{dt^n} + a_{n-1}(t) \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1(t) \frac{dx}{dt} + a_0(t)x = f(t) \quad (2)$$

Dengan $a_n \neq 0, a_n, a_{n-1}, \dots, a_0$ disebut koefisien persamaan diferensial. Fungsi $f(t)$ disebut input atau unsur nonhomogen. Jika $f(t)$ disebut input, maka solusi dari persamaan diferensial x_t biasanya disebut output. Jika ruas sebelah kanan $f(t)$ bernilai nol untuk semua nilai t dalam interval yang ditinjau, maka persamaan ini dikatakan homogen, sebaliknya dikatakan nonhomogen. Contoh persamaan diferensial biasa linier adalah

$$\frac{dx}{dt} = 2x + 3t \quad (3)$$

Yang merupakan persamaan diferensial biasa linier nonhomogen order satu [4].

2.5 Persamaan Diferensial Biasa (PDB) Nonlinier

Jika persamaan diferensial biasa tidak dapat dinyatakan dalam bentuk umum persamaan diferensial biasa linier, yaitu pada (2.1), maka persamaan diferensial tersebut adalah persamaan diferensial biasa nonlinier. Contoh persamaan diferensial biasa nonlinier

$$\frac{d^2 x}{dt^2} + 3x^2 = \sin t \quad (4)$$

Yang merupakan persamaan diferensial biasa nonlinier nonhomogen order dua [4].

2.6 Persamaan Diferensial sebagai Model Matematika

Banyak sekali fenomena yang jika dibawa ke dalam model matematika bentuknya berupa persamaan diferensial biasa (PDB) maupun persamaan diferensial parsial (PDP). Fenomena yang demikian disebut lump problems yang dapat dimodelkan dengan PDB. Dapat diartikan bahwa lumps problems menjadi masalah-masalah yang tak terdistribusi sebagai lawan dari masalah-masalah terdistribusi [1].

2.7 Metode Numerik

Metode numerik adalah teknik yang digunakan untuk memformulasikan persoalan matematik sehingga dapat dipecahkan dengan operasi perhitungan atau aritmatika biasa (tambah, kurang, kali, dan bagi). Metode numerik disebut juga sebagai alternatif dari metode analitik, yang merupakan metode penyelesaian persoalan matematika dengan rumus-rumus aljabar yang sudah baku atau lazim. Disebut demikian, karena adakalanya persoalan

matematika sulit diselesaikan atau bahkan tidak dapat diselesaikan secara analitik sehingga dapat dikatakan bahwa persoalan matematik tersebut tidak mempunyai solusi analitik. Sehingga sebagai alternatif, persoalan matematik tersebut diselesaikan dengan metode numerik. Perbedaan antara metode analitik dan metode numerik adalah metode analitik hanya dapat digunakan untuk menyelesaikan permasalahan yang sederhana dan menghasilkan solusi yang sebenarnya atau solusi sejati. Sedangkan metode numerik dapat digunakan untuk menyelesaikan permasalahan yang sangat kompleks dan nonlinier. Solusi yang dihasilkan dari penyelesaian secara numerik merupakan solusi hampiran atau pendekatan yang mendekati solusi eksak atau solusi sebenarnya. Hasil penyelesaian yang didapatkan dari metode numerik dan metode analitik memiliki selisih, dimana selisih tersebut dinamakan kesalahan [5].

2.8 Metode Runge-Kutta

Secara umum, Runge-Kutta digunakan dalam penyelesaian masalah yang berhubungan dengan perhitungan numerik. Model umum dari metode Runge-Kutta tersebut yaitu :

$$y_{i+1} = y_i + (a_1k_1 + a_2k_2 + \dots + a_nk_n)h \quad (5)$$

Dengan a_i adalah konstan dan k_i adalah :

$$\begin{aligned} k_1 &= f(x_i, y_i) \\ k_2 &= f(x_i + p_1 \cdot h, y_i + q_{11} \cdot k_1 \cdot h) \\ k_3 &= f(x_i + p_2 \cdot h, y_i + q_{21} \cdot k_1 \cdot h + q_{22} \cdot k_2 \cdot h) \\ k_4 &= f(x_i + p_{n-1} \cdot h, y_i + q_{n-1,1} \cdot k_1 \cdot h + q_{n-1,2} \cdot k_2 \cdot h + \dots + \\ &\quad q_{n-1,n-1} \cdot k_{n-1} \cdot h) \end{aligned} \quad (6)$$

Dengan p_{n-1} dan $q_{n-1,2}$ adalah konstan. Persamaan diatas adalah fungsi utama dari Runge-Kutta dan k_n adalah fungsi evaluasi dari metode Runge-Kutta [6].

2.9 Metode Runge-Kutta Orde 4

Metode Runge-Kutta mempunyai galat pemotongan lokal yang sebanding dengan Δx^5 . Metode yang sangat terkenal untuk mengaproksimasi solusi masalah nilai awal orde pertama adalah metode Runge-Kutta orde ke empat. Prosedur metode Runge-Kutta orde ke empat untuk menyelesaikan masalah nilai awal tersebut sebagai berikut :

Tahap 1. Bagilah interval $x_0 \leq x \leq b$ menjadi p subinterval dengan menggunakan titik-titik yang berspasi sama :

$$\begin{aligned} x_1 &= x_0 + \Delta x \\ x_2 &= x_1 + \Delta x \\ &\vdots \\ x_p &= x_{p-1} + \Delta x = b \end{aligned} \quad (7)$$

Tahap 2. Untuk $i = 1, 2, 3, \dots, p$, dapatkan barisan aproksimasi berikut :

$$y_n = y_{n-1} + \frac{K_1 + 2K_2 + 2K_3 + K_4}{6} \quad (8)$$

dimana

$$K_1 = g(x_{n-1}, y_{n-1}) \Delta x$$

$$\begin{aligned}K_2 &= g\left(x_{n-1} + \frac{\Delta x}{2}, y_{n-1} + \frac{K_1}{2}\right) \Delta x \\K_3 &= g\left(x_{n-1} + \frac{\Delta x}{2}, y_{n-1} + \frac{K_2}{2}\right) \Delta x \\K_4 &= g(x_{n-1} + \Delta x, y_{n-1} + K_3) \Delta x\end{aligned}\quad (9)$$

Tahap 3.

$$\begin{aligned}K_1 &= g(x, y) \Delta x \\K_2 &= g(x + 0,5 \Delta x, y + 0,5 \Delta x K_1) \Delta x \\K_3 &= g(x + 0,5 \Delta x, y + 0,5 \Delta x K_2) \Delta x \\K_4 &= g(x + \Delta x, y + \Delta x K_3) \Delta x \\y &= y + \left(\frac{1}{6}\right) (K_1 + 2K_2 + 2K_3 + K_4) \\x &= x + \Delta x [7].\end{aligned}\quad (10)$$

2.10 Metode Adam-Bashfort-Multon

Metode sebelumnya yaitu metode euler taylor, runge kutta dinamakan metode satu langkah (single-step) karena hanya menggunakan satu titik untuk mencari titik sebelumnya yaitu untuk mencari $(x_1$ dan $y_1)$ memerlukan titik awal $(x_0$ dan $y_0)$. Sebaliknya, metode banyak langkah (multi-step) memerlukan beberapa nilai awal sebelumnya.

Metode Adam merupakan metode multi-step yang didasarkan pada kalkulus :

$$\frac{dy}{dx} = f(x) \quad (11)$$

$$dy = f(x) dx \quad (12)$$

$$(1) \quad \int dy = y_{i+1} - y_i = \int_{x_i}^{x_{i+1}} f(x) dx \rightarrow y_{i+1} = y_i + \int_{x_i}^{x_{i+1}} f(x) dx$$

$$\int_{x_i}^{x_{i+1}} f(x) dx = \frac{\Delta h}{2} [3f_i - f_{i-1}] \quad \text{untuk 2 titik}$$

$$\int_{x_i}^{x_{i+1}} f(x) dx = \frac{\Delta h}{12} [23f_i - 16f_{i-1} + 5f_{i-2}] \quad \text{untuk 3 titik} \quad (13)$$

Untuk meramalkan suatu titik $f(x)$, $y(x)$ diperlukan 4 titik sebelumnya yaitu titik :

$$(x_{i-3}, f_{i-3}), (x_{i-2}, f_{i-2}), (x_{i-1}, f_{i-1}) \text{ dan } (x_i, f_i)$$

Dari titik ini diramalkan :

$$p_{i+1} = y_i + \frac{h}{2} (55f_i - 59f_{i-1} + 37f_{i-2} - 9f_{i-3}) \quad (14)$$

Kemudian dikoreksi menjadi :

$$y_{i+1} = y_i + \frac{h}{24} (9f_{i+1} - 19f_i - 5f_{i-1} + f_{i-3}) [8] \quad (15)$$

2.11 Model Pertumbuhan Uang yang diinvestasikan

Model pertumbuhan uang yang ditabung di bank juga mempunyai model pertumbuhan yang sama, yaitu bahwa laju pertambahan banyaknya uang yang ditabung, $P(t)$, sebanding dengan banyaknya uang yang ditabung pada waktu t . Jadi model matematikanya berbentuk

$$\frac{dP(t)}{dt} = rP(t) \quad , \quad r > 0 \quad (16)$$

Dimana r adalah besarnya bunga pertahun [7].

2.12 Investasi

Investasi adalah penanaman modal untuk satu atau lebih aktiva yang dimiliki dan biasanya berjangka waktu lama dengan harapan mendapatkan keuntungan di masa-masa yang akan datang [9].

2.13 Jenis-Jenis Investasi

Produk-produk investasi yang tersedia di pasaran antara lain:

1. Tabungan di bank

Dengan menyimpan uang di tabungan, maka akan mendapatkan suku bunga tertentu yang besarnya mengikuti kebijakan bank bersangkutan. Produk tabungan biasanya memperbolehkan kita mengambil uang kapanpun yang kita inginkan.

2. Deposito di bank

Produk deposito hampir sama dengan produk tabungan. Bedanya, dalam deposito tidak dapat mengambil uang kapanpun yang diinginkan, kecuali apabila uang tersebut sudah menginap di bank selama jangka waktu tertentu (tersedia pilihan antara satu, tiga, enam, dua belas, sampai dua puluh empat bulan, tetapi ada juga yang harian). Suku bunga deposito biasanya lebih tinggi daripada suku bunga tabungan. Selama deposito kita belum jatuh tempo, uang tersebut tidak akan terpengaruh pada naik turunnya suku bunga di bank.

3. Saham

Saham adalah kepemilikan atas sebuah perusahaan tersebut. Dengan membeli saham, berarti membeli sebagian perusahaan tersebut. Apabila perusahaan tersebut mengalami keuntungan, maka pemegang saham biasanya akan mendapatkan sebagian keuntungan yang disebut deviden. Saham juga bisa dijual kepada pihak lain, baik dengan harga yang lebih tinggi yang selisih harganya disebut capital gain maupun lebih rendah daripada kita membelinya yang selisih harganya disebut capital loss. Jadi, keuntungan yang bisa didapat dari saham ada dua yaitu deviden dan capital gain.

4. Properti

Investasi dalam properti berarti investasi dalam bentuk tanah atau rumah [10].

3. METODE PENELITIAN

Adapun langkah-langkah yang digunakan dalam penelitian ini adalah sebagai berikut :

1. Menentukan contoh kasus dari model pertumbuhan uang yang diinvestasikan dengan melakukan simulasi pada beberapa parameter yang mempengaruhi model tersebut.
2. Mencari solusi numerik berdasarkan metode Runge-Kutta Orde 4 dengan menggunakan *software* MATLAB berdasarkan nilai-nilai parameter hasil simulasi kasus dengan cara sebagai berikut :
 - a. Mendeklarasikan parameter-parameter dari contoh kasus tersebut kedalam *software* Matlab R2013b.
 - b. Membuat program untuk solusi numerik berdasarkan metode Runge-Kutta orde empat.
3. Mencari solusi numerik berdasarkan metode Adam-Bashfort Moulton dengan menggunakan *software* MATLAB berdasarkan nilai-nilai parameter hasil simulasi kasus dengan cara sebagai berikut :
 - a. Mendeklarasikan parameter-parameter dari contoh kasus tersebut kedalam *software* Matlab R2013b.
 - b. Membuat program untuk solusi numerik berdasarkan metode Adam-Bashfort Moulton.

4. Mencari nilai galat kedua metode tersebut terhadap solusi analitik dari contoh kasus.
5. Menentukan metode terbaik dalam mengaproksimasi solusi model pertumbuhan uang yang diinvestasikan.

4. HASIL DAN PEMBAHASAN

4.1 Contoh Kasus 1 (Bunga per annum/pertahun)

Jika bunga sebesar 10%pa dan besar tabungan diawal sebesar Rp.10.000.000, hitunglah nilai aproksimasi besarnya uang setiap tahun hingga tahun ke-20 menggunakan metode Runge-Kutta orde 4!

Penyelesaian :

Dari contoh kasus tersebut diketahui :

$$P(0) = y(0) = 10.000.000$$

$$t(0) = x(0) = 0$$

$$\Delta x = 1$$

$$r = 10\%pa = 0,1 \text{ (per-tahun)}$$

$$n=20$$

Contoh kasus 1 tersebut diselesaikan dengan menggunakan *software* MATLAB, sehingga didapatkan nilai aproksimasi sebagai berikut :

Tabel 1. Hasil Aproksimasi dengan menggunakan Metode Runge-Kutta Orde 4
dan Metode Adam-Bashfort Moulton (Contoh Kasus 1)

$X(t)$ (dalam tahun)	$Y(t)$ analitik (dalam rupiah)	$Y(t)$ aproksimasi (Runge-Kutta Orde 4) (dalam rupiah)	$Y(t)$ aproksimasi (Adam-Bashfort Moulton) (dalam rupiah)
1	11.051.708,33333333	11.051.709,18075648	11.051.708,33333333
2	12.214.025,70850695	12.214.027,58160170	12.214.025,70850695
3	13.498.584,97062538	13.498.588,07576003	13.498.584,97062538
4	14.918.242,40080686	14.918.246,97641270	14.918.245,40355310
5	16.487.206,38596838	16.487.212,70700128	16.487.213,08338743
6	18.221.179,62091933	18.221.188,00390509	18.221.190,73458307
7	20.137.516,26596777	20.137.527,07470477	20.137.532,65361586
8	22.255.395,63292315	22.255.409,28492468	22.255.418,28127850
9	24.596.014,13780071	24.596.031,11156950	24.596.044,18244477
10	27.182.797,44135165	27.182.818,28459046	27.182.836,18752231
11	30.041.634,90058981	30.041.660,23946433	30.041.683,84626175
12	33.201.138,67778059	33.201.169,22736548	33.201.199,53976755
13	36.692.930,10013834	36.692.966,67619245	36.693.004,84355603
14	40.551.956,13621164	40.551.999,66844675	40.552.047,00771641
15	44.816.839,15635380	44.816.890,70338065	44.816.948,72162644

16	49.530.263,47779349	49.530.324,24395115	49.530.394,66379336
17	54.739.402,56297260	54.739.473,91727201	54.739.558,70555078
18	60.496.391,14668923	60.496.474,64412947	60.496.576,04422073
19	66.858.847,01724582	66.858.944,42279270	66.859.064,99102014
20	73.890.447,67375541	73.890.560,98930651	73.890.703,63595378

Berdasarkan tabel 1, maka dapat ditentukan galat dari kedua metode tersebut. Galat tersebut berguna dalam menentukan keakuratan dari kedua metode dalam mengaproksimasi nilai eksak dari contoh kasus 1 pada model pertumbuhan uang yang diinvestasikan. Nilai galat tersebut dapat dilihat pada tabel berikut :

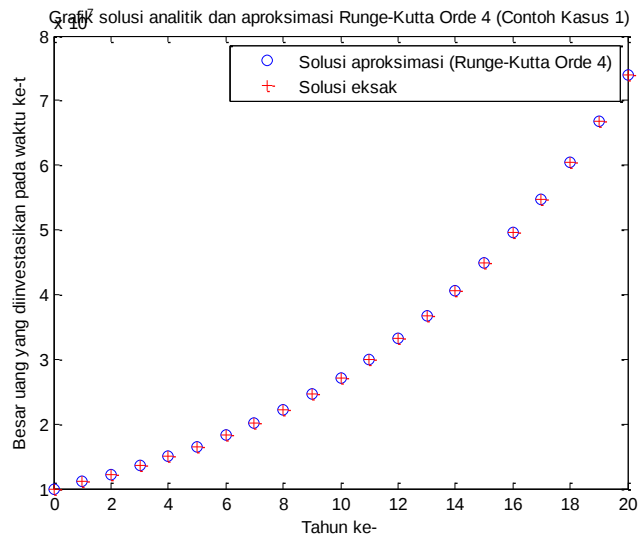
Tabel 2. Nilai Galat dari Metode Runge-Kutta Orde 4 dan Metode Adam-Bashfort

Moulton (Contoh Kasus 1)

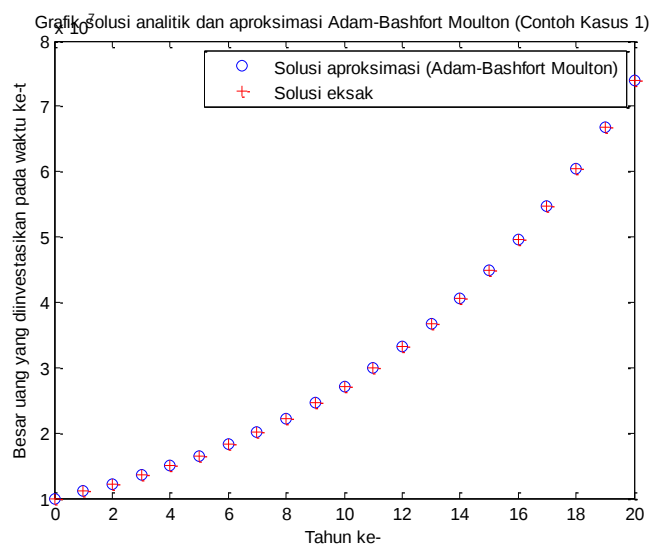
X(i) (dalam tahun)	Nilai Galat Metode Runge-Kutta Orde 4 (dalam rupiah)	Nilai Galat Metode Adam-Bashfort Moulton (dalam rupiah)
1	1,0996599583333	0
2	2,4185317048295	0
3	3,9894225751018	0
4	5,8495087278527	3,0027462411672
5	8,0408811117998	6,6974190510809
6	10,6111569979211	11,1136637404561
7	13,6141647396738	16,3876480869949
8	17,1107103019759	22,6483553498983
9	21,1694350770220	30,0446440614760
10	25,8677755925148	38,7461706623435
11	31,2930369279308	48,9456719420850
12	37,5435930006340	60,8619869574904
13	44,7302283811933	74,7434176951647
14	52,9776379629593	90,8715047687292
15	62,4261026635142	109,5652726367116
16	73,2333613957802	131,1859998703003
17	85,5767018373876	156,1425781846046
18	99,6552950739348	184,8975315019488
19	115,6928020234213	217,9737743213773
20	133,9402826968873	255,9621983766556

Dari tabel tersebut dapat dilihat bahwa nilai aproksimasi dengan kedua metode tersebut pada contoh kasus 1 sudah mendekati nilai analitiknya. Untuk nilai $x(i) \leq 5$, nilai aproksimasi dengan metode Adam-Bashfort Moulton lebih mendekati solusi eksaknya, hal itu ditandai dengan lebih kecilnya galat pada metode Adam-

Bashfort Moulton dibandingkan metode Runge-Kutta Orde 4. Namun, pada nilai $x(i) > 5$ dapat dilihat bahwa nilai galat metode Runge-Kutta Orde 4 lebih kecil dibandingkan metode Adam-Bashfort Moulton, sehingga metode Runge-Kutta Orde 4 lebih baik dibandingkan metode Adam-Bashfort Moulton dalam mengaproksimasikan solusi dari contoh kasus 1 pada $x(i)$ yang semakin besar. Dari *software* MATLAB juga, dapat dilihat grafik dari hasil aproksimasi dengan metode Runge-Kutta Orde 4 dan hasil analitik dari contoh kasus 1, yaitu sebagai berikut :



Gambar 1. Grafik Perbandingan Nilai Aproksimasi dengan Metode Runge-Kutta dan Solusi Analitik pada Contoh Kasus 1.



Gambar 2. Grafik Perbandingan Nilai Aproksimasi dengan Metode Adam-Bashfort Moulton dan Solusi Analitik pada Contoh Kasus 1

Dari kedua grafik tersebut terlihat bahwa kedua metode tersebut layak digunakan dalam mengaproksimasikan nilai eksak dari kasus 1 model pertumbuhan uang yang diinvestasikan, karena pada setiap titik dalam grafik tersebut hampir berhimpitan.

4.2 Contoh Kasus 2 (Bunga harian)

Jika bunga perhari sebesar 0.025% dan besar tabungan diawal sebesar Rp.10.000.000, hitunglah nilai aproksimasi besarnya uang setiap tahun hingga tahun ke-20 menggunakan metode Runge-Kutta orde 4!

Penyelesaian :

Dari contoh kasus tersebut diketahui :

$$P(0) = y(0) = 10.000.000$$

$$t(0) = x(0) = 0$$

$$\Delta x = 1$$

$$r = 0,025\% \text{ perhari} = 0,025\% \times 360 = 9\% \text{ (per-tahun)} = 0,09 \text{ (per-tahun)}$$

Contoh kasus 2 tersebut diselesaikan dengan menggunakan *software* MATLAB, sehingga didapatkan nilai aproksimasi sebagai berikut :

Tabel 3. Hasil Aproksimasi dengan menggunakan Metode Runge-Kutta Orde 4
dan Metode Adam-Bashfort Moulton (Contoh Kasus 2).

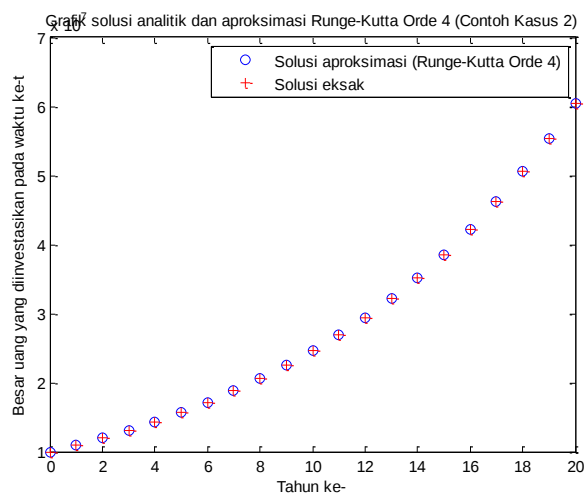
$X(i)$ (dalam tahun)	$Y(i)$ analitik (dalam rupiah)	$Y(i)$ aproksimasi (Runge-Kutta Orde 4) (dalam rupiah)	$Y(i)$ aproksimasi (Adam-Bashfort Moulton) (dalam rupiah)
1	10.941.742,83705210	10.941.742,33750000	10.941.742,33750000
2	11.972.173,63121810	11.972.172,53802400	11.972.172,53802400
3	13.099.644,50733248	13.099.642,71311520	13.099.642,71311520
4	14.333.294,14560340	14.333.291,52802159	14.333.293,33653435
5	15.683.121,85490169	15.683.118,27478840	15.683.122,26461198
6	17.160.068,62184859	17.160.063,92112722	17.160.070,47542374
7	18.776.105,79264343	18.776.099,79200039	18.776.109,35985808
8	20.544.332,10643888	20.544.324,60272556	20.544.337,69375382
9	22.479.079,86676471	22.479.070,63009851	22.479.087,82300668
10	24.596.0311,1156949	24.596.019,88210017	24.596.041,83337015
11	26.912.344,72349263	26.912.331,20779673	26.912.358,66117915
12	29.446.795,51065524	29.446.779,37771719	29.446.813,17484586
13	32.219.926,38528500	32.219.907,26201901	32.219.948,35425114
14	35.254.214,87365382	35.254.192,33991571	35.254.241,80153393
15	38.574.255,30696974	38.574.228,89000238	38.574.287,93293014
16	42.206.958,16996552	42.206.927,33821547	42.206.997,32841390
17	46.181.768,22299781	46.181.732,37923384	46.181.814,85496623
18	50.530.903,16563867	50.530.861,63929575	50.530.958,33146601
19	55.289.614,77624004	55.289.566,81490369	55.289.679,66969735
20	60.496.474,64412945	60.496.419,40406668	60.496.550,60814787

Berdasarkan tabel 3, maka dapat ditentukan galat dari kedua metode tersebut. Galat tersebut berguna dalam menentukan keakuratan dari kedua metode dalam mengaproksimasi nilai eksak dari contoh kasus 2 pada model pertumbuhan uang yang diinvestasikan. Nilai galat tersebut dapat dilihat pada tabel berikut :

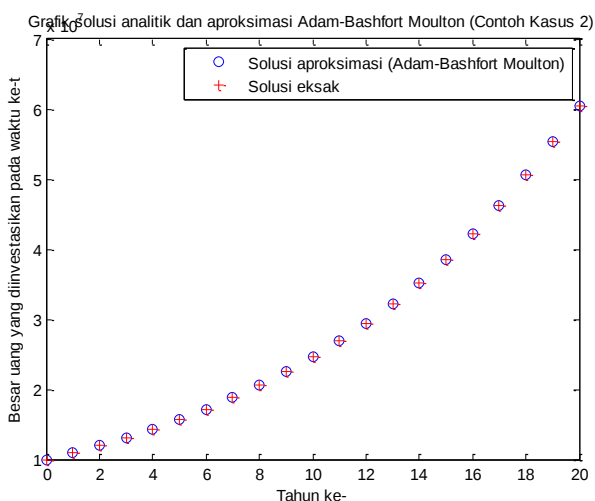
Tabel 4. Nilai Galat dari Metode Runge-Kutta Orde 4 dan Metode Adam-Bashfort
 Moulton (Contoh Kasus 2)

X(i) (dalam tahun)	Nilai Galat Metode Runge-Kutta Orde 4 (dalam rupiah)	Nilai Galat Metode Adam-Bashfort Moulton (dalam rupiah)
1	0.499552099034190	0.499552099034190
2	1.093194099143148	1.093194099143148
3	1.794217279180884	1.794217279180884
4	2.617581808939576	0.809069050475955
5	3.580113289877772	0.409710289910436
6	4.700721371918917	1.853575147688389
7	6.00643040984869	3.567214649170637
8	7.503713317215443	5.587314940989018
9	9.236666198819876	7.956241969019175
10	11.229469321668148	10.721800658851862
11	13.515695899724960	13.937686521559954
12	16.132938049733639	17.664190620183945
13	19.123265992850065	21.968966137617826
14	22.533738113939762	26.927880108356476
15	26.416967362165451	32.625960402190685
16	30.831750050187111	39.158448383212090
17	35.843763969838619	46.631968423724174
18	41.526342920958996	55.165827341377735
19	47.961336351931095	64.893457308411598
20	55.240062765777111	75.964018419384956

Dari tabel tersebut dapat dilihat bahwa nilai aproksimasi dengan kedua metode tersebut pada contoh kasus 2 sudah mendekati nilai analitiknya. Untuk nilai $x(i) \leq 3$, nilai aproksimasi dari kedua metode tersebut bernilai sama, hal itu ditandai dengan sama besarnya nilai galat dari kedua metode tersebut. Pada nilai $x(i)=4$ hingga $x(i)=10$, nilai aproksimasi dari metode Adam-Bashfort lebih baik dibandingkan aproksimasi dari metode Runge-Kutta Orde 4 karena galat yang dimiliki metode Adam-Bashfort Moulton lebih rendah dibandingkan dengan metode Runge-Kutta Orde 4 pada nilai $x(i)$ tersebut. Namun, pada nilai $x(i) > 10$ dapat dilihat bahwa nilai galat metode Runge-Kutta Orde 4 lebih kecil dibandingkan metode Adam-Bashfort Moulton, sehingga metode Runge-Kutta Orde 4 lebih baik dibandingkan metode Adam-Bashfort Moulton dalam mengaproksimasi solusi dari contoh kasus 2 pada $x(i)$ yang semakin besar. Dari *software* MATLAB juga, dapat dilihat grafik dari hasil aproksimasi dengan metode Runge-Kutta Orde 4 dan hasil analitik dari contoh kasus 1, yaitu sebagai berikut :



Gambar 3. Grafik Perbandingan Nilai Aproksimasi dengan Metode Runge-Kutta dan Solusi Analitik pada Contoh Kasus 2.



Gambar 4. Grafik Perbandingan Nilai Aproksimasi dengan Metode Adam-Bashfort Moulton dan Solusi Analitik pada Contoh Kasus 2.

Dari kedua grafik tersebut terlihat bahwa kedua metode tersebut layak digunakan dalam mengaproksimasi nilai eksak dari kasus 2 model pertumbuhan uang yang diinvestasikan, karena pada setiap titik dalam grafik tersebut hampir berimpitan.

4. KESIMPULAN

Dari hasil dan pembahasan penelitian yang telah dilakukan, maka dapat disimpulkan bahwa metode Runge-Kutta Orde 4 dan Adam-Bashfort Moulton dapat diterapkan pada model pertumbuhan uang yang diinvestasikan tersebut. Semakin kecil bunga pertahunnya, maka hasil aproksimasi semakin mendekati hasil eksaknya. Hal ini ditunjukkan pada contoh kasus 2 yang hasil aproksimasinya lebih mendekati nilai eksaknya dibandingkan dengan contoh kasus 1. Sebaliknya, semakin besar bunga pertahunnya, maka selisih antara hasil aproksimasi dan

hasil eksaknya akan semakin besar. Hal ini ditunjukkan pada contoh kasus 1 yang memiliki selisih yang besar antara hasil aproksimasi dan hasil eksaknya.

Metode Runge-Kutta Orde 4 lebih baik dalam mengaproksimasikan suatu nilai pada $x(i)$ yang besar dibandingkan dengan metode Adam-Bashfort Moulton, hal itu terlihat pada kedua contoh kasus. Sebaliknya, dari kedua contoh kasus tersebut terlihat bahwa metode Adam-Bashfort Moulton lebih baik dalam mengaproksimasikan suatu nilai pada $x(i)$ yang kecil dibandingkan metode Runge-Kutta Orde 4.

KEPUSPTAKAAN

- [1] Cahyono, Edi. 2013. *Pemodelan Matematika*. Graha Ilmu, Yogyakarta.
- [2] Wardiman. 1981. *Persamaan Diferensial*. Citra Offset, Yogyakarta.
- [3] Campbell. Haberman. 2008. *Introduction to Differential Equations with Dynamical Systems*. Princeton University Press, New Jersey.
- [4] Hidayat. 2006. *Persamaan Diferensial Parsial*. UPT Penerbitan Universitas Jember, Jember.
- [5] Triatmodjo. 2012. *Metode Numerik Dilengkapi dengan Program Komputer*. Beta Offset, Yogyakarta.
- [6] Muhammad, Singgih Tahwin. dkk. 2015. Pengkajian metode extended runge kutta dan penerapannya pada persamaan diferensial biasa. *Jurnal Sains dan Seni ITS Vol. 4, No.1, (2015) 2337-3520 (2301-928X Print)*.
- [7] Kartono. 2011. *Persamaan Diferensial*. C.V Andi Offset, Yogyakarta.
- [8] Salusu, Abraham. 2008. *Metode Numerik : Dilengkapi dengan Animasi Matematika dan Panduan Singkat Maple*. Graha Ilmu, Yogyakarta.
- [9] Sunariyah, 2006, *Pengantar Pengetahuan Pasar Modal*, Edisi ke Lima, UPP AMP YKPN, Yogyakarta.
- [10] Senduk. 2004. *Seri Perencana Keuangan Keluarga : Mencari Penghasilan Tambahan*. Alex Media Komputoindo, Jakarta.

MENYELESAIKAN PERSAMAAN DIFERENSIAL LINEAR ORDE- n NON HOMOGEN DENGAN FUNGSI GREEN

Fathurrohman Al Ayubi¹⁾, Dorrah Aziz, dan Muslim Ansori
Jurusan Matematika, Fakultas Matematika dan Ilmu Pengetahuan Alam
Universitas Lampung, Jl. Prof. Soemantri Bodjonegoro No. 1, Bandar Lampung 35145
alfathur26@gmail.com¹⁾

ABSTRAK

Dalam penelitian ini akan disajikan bagaimana menyelesaikan persamaan diferensial linear orde- n non homogen dengan fungsi Green melalui transformasi Laplace. Solusi umum dari persamaan diferensial linear orde- n non homogen terdiri dari solusi homogen dan solusi non homogen. Solusi non homogen sering juga disebut solusi partikular. Selanjutnya dari solusi partikular ini dapat diselesaikan dengan fungsi Green melalui transformasi Laplace. Berdasarkan hasil penelitian ini, didapat bahwa persamaan diferensial linear orde- n non homogen dapat diselesaikan dengan fungsi Green melalui transformasi Laplace. Solusi umum yang diperoleh yaitu :

$$y(x) = y_h(x) + \int_0^x f(t) \cdot w(x-t) dt$$

Kata Kunci: Persamaan Diferensial Linear Orde- n Non Homogen, Fungsi Green, Transformasi Laplace

1. PENDAHULUAN

Latar Belakang

Matematika adalah salah satu ilmu pengetahuan yang mengalami perkembangan seiring dengan perkembangan ilmu pengetahuan lainnya. Matematika mempunyai peranan penting untuk ilmu pengetahuan lain seperti kimia, biologi, fisika, ekonomi, dan lain-lain. Salah satu ilmu matematika yang mempunyai peranan penting dengan ilmu pengetahuan lainnya adalah persamaan diferensial.

Persamaan diferensial merupakan persamaan yang berkaitan dengan turunan suatu fungsi atau memuat suku-suku dari fungsi tersebut dan turunannya. Menurut peubah bebasnya, persamaan diferensial dibagi menjadi dua, yaitu persamaan diferensial biasa (satu variabel) dan persamaan diferensial parsial (dua atau lebih variabel). Persamaan diferensial biasa dapat dibagi menurut kelinieran, orde, dan koefisiennya [1]. Persamaan diferensial yang akan dibahas dalam penelitian ini adalah persamaan diferensial linear orde- n non homogen dengan koefisien konstan.

Persamaan diferensial linear orde- n non homogen dengan koefisien konstan sering kali diselesaikan dengan beberapa metode penyelesaian, antara lain: metode koefisien tak tentu, metode invers operator, dan lain-lain. Selain metode-metode tersebut masih ada cara lain untuk menyelesaikan persamaan diferensial linear orde- n non homogen dengan koefisien konstan, yaitu dengan metode fungsi Green [2].

Metode fungsi Green adalah metode penyelesaian yang dalam proses menemukan penyelesaian suatu persamaan diferensial linear orde- n non homogen dengan koefisien konstan, terlebih dahulu ditentukan nilai fungsi Green dari suatu persamaan diferensial tersebut. Nilai fungsi Green dapat ditemukan dengan menggunakan transformasi Fourier, transformasi Laplace, dan variasi parameter [2].

Berdasarkan latar belakang di atas, pada penelitian ini digunakan metode fungsi Green dalam penyelesaian suatu persamaan diferensial linear orde- n non homogen dengan koefisien konstan melalui transformasi Laplace.

Rumusan Masalah

Rumusan masalah dalam penelitian ini adalah bagaimana cara menyelesaikan persamaan diferensial linear orde- n non homogen melalui transformasi Laplace?

Tujuan

Tujuan dari penelitian ini adalah mengetahui cara menyelesaikan persamaan diferensial linear orde- n non homogen dengan koefisien konstan menggunakan metode fungsi Green melalui transformasi Laplace.

2. LANDASAN TEORI

Definisi 1. Persamaan Diferensial Linear Orde- n Non Homogen

Bentuk umum dari persamaan diferensial linear orde- n non homogen yaitu :

$$a_n(x)y^n + a_{n-1}(x)y^{n-1} + \dots + a_1(x)y' + a_0(x)y = g(x) \quad (1)$$

dimana $a_i(x)$ ($i = 0, 1, \dots, n-1, n$) konstan, $g(x) \neq 0$, dan semua nilai dari

$y, y', \dots, y^{n-1}$, dan y^n dengan selang $x > 0$ dan $\frac{f(x)}{a_n(x)} = \phi(x)$ [1].

Definisi 2. Fungsi Green

Fungsi $G(x, t)$ dari persamaan

$$a_n(x) + a_{n-1}(x)y^{(n-1)} + \dots + a_1(x)y' + a_0(x)y = f(x) \quad (2)$$

dikatakan fungsi Green untuk masalah nilai awal persamaan diferensial di atas jika memenuhi kondisi sebagai berikut :

1. $G(x, t)$ terdefinisi pada daerah $R = I \times I$ dari semua titik (x, t) dengan x dan t terletak pada selang I .
2. $G(x, t)$, $\frac{\partial G}{\partial x}$, $\frac{\partial^2 G}{\partial x^2}$, $\dots$, $\frac{\partial^n G}{\partial x^n}$ merupakan fungsi kontinu pada $R = I \times I$.
3. Untuk setiap x_0 dalam selang I dan fungsi $f \in C(I)$, fungsi $y_p(x) = \int_{x_0}^x G(x, t)f(t)dt$ adalah solusi persamaan diferensial (1) yang memenuhi kondisi awal $y_p(x_0) = y_p'(x_0) = y_p''(x_0) = \dots = y_p^{(n-1)}(x_0) = 0$ [3]

Definisi 3. Transformasi Laplace

Misalkan $f(t)$ suatu fungsi dari t terdefinisi untuk $t > 0$. Kemudian $\int_0^\infty e^{-st} f(t) dt$, jika ada dinamakan transformasi suatu fungsi dari s , katakan $F(s)$. Fungsi $F(s)$ ini dinamakan transformasi Laplace dari $f(t)$ dan dinotasikan oleh

$$\mathcal{L}\{f(t)\} = F(s) = \int_0^\infty e^{-st} f(t) dt \quad (3)$$

Operasi yang baru ditunjukkan yang menghasilkan $F(s)$ dari suatu fungsi $f(t)$ yang diberikan, disebut transformasi Laplace [3].

Definisi 4. Transformasi Laplace Invers

Jika $\mathcal{L}\{f(t)\} = F(s)$ maka $f(t)$ dinamakan transformasi Laplace Invers dari $F(s)$ dan dinotasikan dengan $f(t) = \mathcal{L}^{-1}\{F(s)\}$. Kemudian untuk mencari $\mathcal{L}^{-1}\{F(s)\}$ kita harus mencari suatu fungsi dari t yang transformasi Laplacanya adalah $F(s)$ [4].

Definisi 5. Konvolusi [1]

Konvolusi dari dua fungsi $f(x)$ dan $g(x)$ adalah

$$f(x) * g(x) = \int_0^x f(t)g(x-t)dt \quad (4)$$

3. METODOLOGI PENELITIAN

Metode Penelitian

Penelitian ini bersifat studi literatur dengan mengkaji jurnal-jurnal dan buku-buku teks yang berkaitan dengan bidang yang teliti. Adapun langkah-langkah yang dilakukan dalam penelitian ini adalah:

1. Transformasi Laplace pada kedua sisi dari persamaan diferensial linear orde- n non homogen, sehingga diperoleh $Y(s)$.
2. Mengambil transformasi Laplace invers untuk memperoleh $\mathcal{L}^{-1}\{Y(s)\}$.
3. Dengan menggunakan teorema Konvolusi, diperoleh solusi umum fungsi Green persamaan diferensial linear orde- n non homogen, akan ditunjukkan bahwa fungsi Green $G(x, t)$ untuk persamaan diferensial linear orde- n non homogen [1].

4. HASIL DAN PEMBAHASAN

Diberikan persamaan diferensial linear orde- n non homogen dengan koefisien konstan:

$$a_n(x)y^n + a_{n-1}(x)y^{n-1} + \dots + a_1(x)y' + a_0(x)y = g(x) \quad (5)$$

Solusi umum untuk persamaan (5) adalah :

$$y(x) = y_h(x) + y_p(x) \quad (6)$$

Dengan $y_h(x)$ merupakan solusi umum persamaan diferensial homogenya dan $y_p(x)$ merupakan solusi khususnya atau transformasi Laplace [4].

Untuk memperoleh solusi khusus $y_p(x)$ persamaan diferensial linear orde- n non homogen maka dilakukan transformasi Laplace pada kedua sisi persamaan (5):

$$a_n(x)\mathcal{L}(y^n) + a_{n-1}(x)\mathcal{L}(y^{(n-1)}) + \dots + a_1(x)\mathcal{L}(y') + a_0(x)\mathcal{L}(y) = \mathcal{L}(g(x)) \quad (7)$$

di mana

$$\mathcal{L}(y^n(x)) = s^n Y(s) - s^{(n-1)}y(0) - s^{(n-2)}y'(0) - \dots - sy^{(n-2)}(0) - y^{(n-1)}(0) \quad (8)$$

$$\mathcal{L}(y^{(n-1)}(x)) = s^{(n-1)}Y(s) - s^{(n-2)}y(0) - \dots - sy^{(n-2)}(0) - y^{(n-1)}(0) \quad (9)$$

$\vdots$

$$\mathcal{L}(y'(x)) = sY(s) - y(0) \quad (10)$$

$$\mathcal{L}(y(x)) = Y(s). \quad (11)$$

Misalkan diberikan kondisi awal untuk $y(x)$ pada $x = 0$ yaitu:

$$y(0) = c_0, y'(0) = c_1, \dots, y^{(n-1)}(0) = c_{n-1}$$

Substitusikan kondisi awal ke persamaan (8), (9), dan (10) untuk menyederhanakan persamaan tersebut.

$$\mathcal{L}(y^n(x)) = s^n Y(s) - c_0 s^{(n-1)} - c_1 s^{(n-2)} - \dots - c_{n-2} s - c_{n-1} \quad (12)$$

$$\mathcal{L}(y^{(n-1)}(x)) = s^{(n-1)}Y(s) - c_0 s^{(n-2)} - \dots - c_{n-2} s - c_{n-1} \quad (13)$$

$\vdots$

$$\mathcal{L}(y'(x)) = sY(s) - c_0 \quad (14)$$

$$\mathcal{L}(y(x)) = Y(s) \quad (15)$$

Jika $c_0, c_1, \dots, c_n$ adalah konstanta-konstanta yang diketahui adalah nol maka persamaan (12), (13), dan (14) menjadi

$$\mathcal{L}(y^n(x)) = s^n \cdot Y(s) \quad (16)$$

$$\mathcal{L}(y^{(n-1)}(x)) = s^{(n-1)} \cdot Y(s) \quad (17)$$

$\vdots$

$$\mathcal{L}(y'(x)) = s \cdot Y(s) \quad (18)$$

$$\mathcal{L}(y(x)) = Y(s). \quad (19)$$

Substitusikan persamaan (16), (17), (18), dan (19) ke persamaan (7).

Sesuai dengan definisi 3 bahwa $\mathcal{L}(f(x)) = F(s)$ maka persamaan (7) menjadi

$$a_n(x)s^n Y(s) + a_{n-1}(x)s^{(n-1)}Y(s) + \dots + a_1(x)sY(s) + a_0(x)Y(s) = F(s)$$

$$[a_n(x)s^n + a_{n-1}(x)s^{(n-1)} + \dots + a_1(x)s + a_0(x)]Y(s) = F(s)$$

$$Y(s) = F(s) \frac{1}{a_n(x)s^n + a_{n-1}(x)s^{(n-1)} + \dots + a_1(x)s + a_0(x)} \quad (20)$$

Misalkan untuk $\frac{1}{a_n(x)s^n + a_{n-1}(x)s^{(n-1)} + \dots + a_1(x)s + a_0(x)} = G(s)$ maka

$$Y(s) = F(s) \cdot G(s) \quad (21)$$

Diketahui $Y(s) = \mathcal{L}^{-1}(y(x))$.

Berdasarkan definisi transformasi Laplace bahwa $Y(s)$ dan $F(s)$ adalah transformasi Laplace dari $y(x)$ dan $f(x)$. Sudah pasti $F(s)$ diketahui dan persamaan (21) kebalikan dari $y(x)$.

Karena persamaan (5) untuk $x > 0$ maka solusi khusus dari persamaan (5) dapat dicari menggunakan Teorema Konvolusi sehingga didapatkan:

Diketahui :

$$F(s) = \int_0^\infty e^{-su} f(u) du \text{ dan } G(s) = \int_0^\infty e^{-sv} g(v) dv$$

Maka

$$F(s) \cdot G(s) = \int_0^\infty e^{-su} f(u) du \cdot \int_0^\infty e^{-sv} g(v) dv$$

$$F(s) \cdot G(s) = \int_0^\infty \int_0^\infty e^{-s(u+v)} f(u) g(v) du dv$$

Misal : $u + v = t ; v = t - u$

$$F(s) \cdot G(s) = \int_{t=0}^\infty \int_{u=0}^t e^{-st} f(u) g(t-u) du dt$$

$$F(s) \cdot G(s) = \int_{t=0}^\infty e^{-st} \left[\int_{u=0}^t f(u) g(t-u) du \right] dt$$

$$F(s) \cdot G(s) = \mathcal{L} \left[\int_0^t f(u) g(t-u) du \right]$$

$$F(s) \cdot G(s) = \int_0^t f(u) g(t-u) du$$

Jadi, solusi partikuler untuk persamaan diferensial orde- n non homogen dengan koefisien konstan adalah:

$$y_p(x) = F(x) \cdot W(x) = \int_0^x f(t) \cdot w(x-t) dt \quad (22)$$

Jadi solusi umum dari persamaan (5) adalah

$$y(x) = y_h(x) + y_p(x)$$

$$y(x) = y_h(x) + F(x) \cdot G(x)$$

$$y(x) = y_h(x) + \int_0^x f(t) \cdot w(x-t) dt \quad (23)$$

5. SIMPULAN

Kesimpulan

Berdasarkan hasil dan pembahasan maka diperoleh kesimpulan bahwa persamaan diferensial non homogen orde- n dengan koefisien konstanta dapat diselesaikan dengan fungsi Green dengan menggunakan metode transformasi Laplace. Solusi umum dari persamaan diferensial non homogen orde- n dengan koefisien konstan adalah:

$$y(x) = y_h(x) + \int_0^x f(t) \cdot w(x-t) dt$$

Saran

Pada penelitian ini, fungsi Green yang dibahas dalam penelitian ini hanyalah fungsi Green dari persamaan diferensial non homogen orde- n dengan koefisien konstan, di mana untuk mendapatkan fungsi Green menggunakan metode transformasi Laplace. Untuk itu, diperlukan pengembangan lebih lanjut tentang fungsi Green dan cara mendapatkan fungsi Green. Misalnya mencari fungsi Green dari persamaan diferensial parsial dengan metode transformasi Fourier.

DAFTAR PUSTAKA

- [1] Alwi, Wahidah, dkk. 2015. Fungsi Green Yang Dikonstruksikan Pada Persamaan Diferensial Linear Tak Homogen Orde- n . *Jurnal MSA*. Vol. 3, No.1, Januari-Juni 2015. UINAM, Makassar.
- [2] Bronson, R. dan Costa, G. 2003. *Persamaan Diferensial*. Erlangga, Jakarta.
- [3] Kartono. 2002. *Penuntun Belajar Persamaan Diferensial*. Andi Offset, Yogyakarta.
- [4] Djauhari, Eddy. 2015. Membangun Fungsi Green Dari Persamaan Linear Non Homogen Tingkat- n . *Jurnal Matematika Integratif*. Vol. 11, No.2, Oktober 2015. Universitas Padjadjaran, Bandung.

Penyelesaian Kata Ambigu Pada Proses POS Tagging Menggunakan Algoritma Hidden Markov Model (HMM)

Agus Mulyanto¹, Yeni Agus Nurhuda², Nova Wiyanto³

Fakultas Teknik dan Ilmu Komputer Prodi Informatika

Universitas Teknokrat Indonesia - Lampung

Jl. Z.A. Pagar ALam No.9-11 Kedaton Bandar Lampung

agus.mulyanto@teknokrat.ac.id, agus.nurhuda@teknokrat.ac.id, novawiyanto@gmail.com

ABSTRAK

Part of Speech Tagging (POS Tagging) adalah proses memberi label pada setiap kata dalam kalimat dengan POS atau *tag* yang sesuai dengan kelas kata seperti kata kerja, kata keterangan, kata sifat, dan lainnya. Pelabelan kata dapat dilakukan dengan dua cara yakni berbasis aturan (*rule-based*) dan probabilitas (*probability-based*). Penelitian ini melakukan POS tagging berbasis probabilitas dengan menggunakan *Hidden Markov Model* (HMM) yang digunakan pada aplikasi SMIQA (*Swamedikasi Interactive Question Answering System*) yaitu sebuah aplikasi untuk tanya jawab otomatis untuk berkonsultasi terkait penyakit ringan, gejala, dan obat-obatan yang dapat digunakan. Masalah yang dihadapi dalam proses POS Tagging adalah ditemukannya kata ambigu atau kata yang memiliki dua atau lebih makna dalam sebuah dialog yang menyebabkan sulitnya menemukan tag atau kelas kata yang tepat pada sebuah kata dalam sebuah kalimat. Proses POS Tagging dalam penelitian ini dilakukan dalam enam tahapan yaitu proses parsing, tokenisasi menggunakan *Multi-word Expression (MWE) Tokenizer*, pengenalan entitas (*Name Entity Recognizer*), dan penyelesaian kata ambigu dengan metode HMM. Korpus atau dataset yang digunakan sebagai data training dalam penelitian ini adalah data set yang berkaitan dengan penyakit ringan yang telah memiliki label pada setiap katanya. Penelitian ini menggunakan *corpus* bahasa Indonesia yang dibuat oleh Universitas Indonesia yang berjumlah 903.180 kata yang di tambah dengan sekitar 2.200 kata dibidang kesehatan. Pengujian performa penggunaan *Hidden Markov Model* pada proses penyelesaian kata ambigu dilakukan dengan menggunakan kalimat yang mengandung kata mabigu sebanyak 40 kalimat. Hasil yang didapat untuk akurasi *Precision* rata-rata 91,472 %, sedangkan pada pengujian *recall* rata-rata hasil POS Tagging yang didapat 84,877 % dengan nilai standar deviasi 15,435.

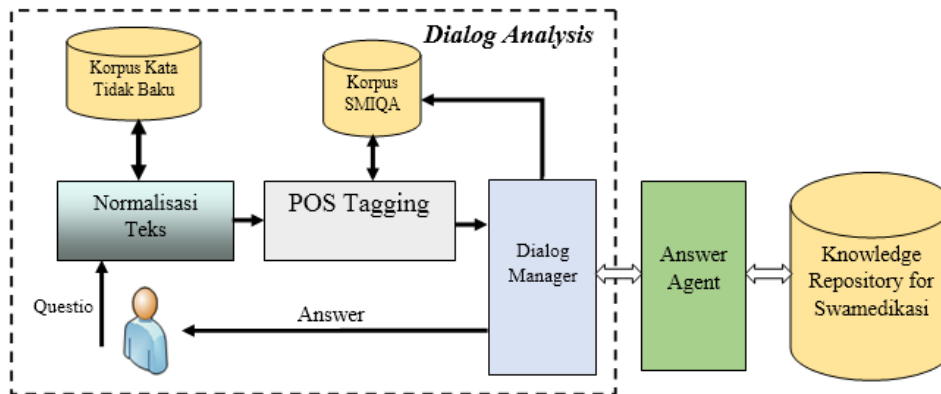
Kata Kunci: *Part of Speech Tagging*, Swamedikasi, *Interactive Question Answering System (IQA)*, *Hidden Markov Model*.

1. PENDAHULUAN

Sistem tanya jawab (*Interactive Question Answering* atau IQA) merupakan sebuah sistem berbasis komputer yang mampu memberikan jawaban terhadap pertanyaan yang diajukan oleh user atau pengguna pada domain yang luas (*open domain*) atau pada domain tertentu atau khusus (*closed domain*) secara otomatis dan interaktif. Sistem tanya jawab otomatis telah banyak diteliti sebelumnya, salah satunya adalah sistem tanya jawab interaktif untuk menangani pelayanan konsumen berupa interaksi antara perusahaan dengan konsumen. Sistem tersebut dirancang untuk memiliki kemampuan berkomunikasi secara aktif dengan pengguna. Hasilnya terjalinlah interaksi yang baik antara perusahaan dan konsumen sehingga konsumen puas terhadap layanan yang diberikan oleh perusahaan [16]. Selain itu Pada tahun 2012, Ayu Purwarianti dan Zulfikar Hakim melakukan penelitian yang sama yaitu membuat IQA untuk *Custumner Service Representative*, dan kemudian Fakhriyanto dan Takdir tahun 2015 membuat sebuah sistem penjawab otomatis untuk permintaan data publikasi di Badan Pusat Statistik (BPS).

Swamedikasi merupakan upaya pengobatan yang dilakukan secara mandiri tanpa resep dokter yang banyak dilakukan oleh masyarakat untuk mengatasi keluhan penyakit ringan dengan menggunakan obat-obatan yang dapat dibeli secara bebas, sehingga masyarakat memerlukan pedoman yang terpadu dan dapat dipercaya agar tidak terjadi kesalahan pengobatan (*medication error*)[5]. Beberapa penyakit yang tergolong kedalam penyakit ringan diantaranya: batuk, flu, demam, nyeri, sakit maag, kecacingan, diare, biang keringat, jerawat, panu, ketombe, kudis, kutil, luka bakar, luka iris dan serut [5]. Pada penelitian ini dibangun sebuah sistem tanya jawab

interaktif (*Interactive Question Answering*) berupa aplikasi berbasis android yang diberi nama SMIQA. Aplikasi ini diharapkan dapat digunakan oleh masyarakat untuk berdialog dan berkonsultasi dalam melakukan swamedikasi penyakit ringan seperti batuk, flu, dan lain-lain. Arsitektur SMIQA terdiri atas beberapa 5 komponen utama yaitu komponen untuk normalisasi kata, *Part-Of-Speech Tagger*, *Dialogue manager*, *Answer Agent*, dan *Knowledge Repository* seperti pada gambar 1.



Gambar 1. Arsitektur SMIQA

Fokus pembahasan pada paper ini adalah pada komponen *Part-Of-Speech Tagger* yaitu bagian yang bertugas untuk memberikan label atau tag berdasarkan kelas kata secara otomatis pada suatu kata dalam kalimat [2] yang nantinya akan digunakan oleh dialog manajer untuk menganalisis dialog. Pelabelan kata dapat dilakukan dengan dua cara yakni berbasis aturan (*rule-based*) dan probabilitas (*probability-based*) dari sebuah model yang dibangun. *Rule-Based* tagging dilakukan dengan cara *top-down*, yaitu melakukan konsultasi dengan ahli linguistik untuk mendefinisikan aturan-aturan yang biasa digunakan manusia. Sedangkan *probability-based tagging* dilakukan dengan cara *bottom-up*, yaitu menggunakan korpus sebagai data training untuk menentukan secara probalistik tag yang terbaik untuk sebuah kata dalam sebuah konteks [3].

Terdapat beberapa metode berbasis probalistik yang dapat digunakan dalam proses *POS Tagging* dalam text berbahasa Indonesia, di antaranya dengan menggunakan *Hidden Markov Model* [6] [7], *Maximun Entropy model* [8], *Conditional Random Field* (CRF) [9] *Maximun Entropy model* dan CRF [10]. *Hidden Markov Model* memiliki kelebihan tersendiri diantara metode probalistik yang lainnya, *Hidden Markov Model* cocok digunakan pada data yang bersifat temporal sequence yaitu rangkaian logika yang keluarannya dipengaruhi oleh hasil sebelumnya [7]. *Part-Of-Speech* tagging yang diterapkan pada SMIQA menggunakan metode *Hidden Markov Model*.

2. LANDASAN TEORI

A. Hidden Markov Model

Hidden Markov Model (HMM) adalah sebuah model statistik dari sebuah sistem yang melakukan perhitungan probabilitas dari suatu kejadian yang tidak dapat diamati berdasarkan kejadian yang dapat diamati [2]. Perhitungan probabilitas dilakukan dengan melihat kejadian-kejadian lain yang dapat diamati secara langsung. *Hidden Markov Model* memiliki 2 macam bagian yaitu *observed state* dan *hidden state*. *Observed state*

merupakan bagian yang dapat diamati secara langsung dan *hidden state* merupakan bagian yang tidak dapat diamat [6].

Dalam proses *POS Tagging* ini, data yang akan diobservasi adalah kumpulan kata atau kalimat dalam sebuah dialog antara user dengan aplikasi SMIQA yang kemudian akan ditentukan kelas kata atau *tag* apa yang tepat. Proses *POS Tagging* dimulai dengan mempertimbangkan semua urutan tag yang mungkin untuk kalimat tersebut. Probabilitas pelabelan kata menggunakan bigram yaitu probabilitas suatu kemunculan *tag* hanya bergantung dari *tag* sebelumnya.

$$P(t_i|w_i) = P(t_i) \times P(t_i | t_{i-1}) \times P(w_i|t_i)$$

Dimana persamaan tersebut dapat diturunkan menjadi

$$P(t_i|w_i) = \frac{P(t_i)}{n} \times \frac{P(t_i | t_{i-1})}{P(t_{i-1})} \times \frac{P(w_i|t_i)}{P(t_i)} \quad (1)$$

Keterangan :

t_i = kelas kata atau *tag* dari w_i yang ada di *corpus*

w_i = kata yang dicari kelas katanya

t_{i-1} = kelas kata atau *tag* sebelum kelas kata dari w_i yang ada di *corpus* sebanyak 1

n = jumlah semua kata didalam *corpus*

P = probabilitas atau peluang

B. Kelas Kata

Kata-kata dikategorikan ke dalam kelas-kelas tertentu dan menjadi posisi penting dalam deskripsi dan studi gramatika [14]. Label atau *tag* yang diberikan ke suatu kata dalam suatu kalimat menunjukkan kelas kata (*word class*) dari kata yang bersangkutan, dalam konteks kalimat tersebut. Kelas kata ini juga disebut sebagai *part of speech*. Kumpulan atau koleksi label atau *tag part of speech* atau kelas kata disebut sebagai *tagset*. Dalam proses *POS tagging* bahasa Indonesia *tagset* yang digunakan seperti ditunjukkan dalam Tabel 1 [10].

Tabel 2.1 Kelas Kata

No	Tag	Description	Example
1.	(	Opening parenthesis	{ [
2.	)	Closing parenthesis	}]
3.	,	Comma	,
4.	.	Sentence terminator	.?!
5.	:	Colon or ellipsis	::
6.	--	Dash	--

7.	“	Opening quotation mark	““
8.	”	Closing quotation mark	””
9.	\$	Dollar	\$
10.	Rp	Rupiah	Rp
11.	SYM	Symbols	%&"""). *+,.<=>@A[fj] U.S U.S.S.R *
12.	NNC	Countable commonnouns	Buku, rumah,karyawan
13.	NNU	Uncountable commonnouns	Air, gula, nasi, hujan
14.	NNG	Genitive Commonnouns	Idealnya, komposisinya, fungsinya, jayanya
15.	NNP	Proper nouns	Jakarta, Soekarno Hatta, Australia, BCA
16.	PRP	Personal pronouns	Saya,aku,dia,kami
17.	PRN	Numberpronouns	Kedua-duanya ,ketiga-tiganya
18.	PRL	Locativepronouns	Sini,situ,sana
19.	WP	WH-pronouns	Apa, siapa, mengapa, bagaimana, berapa
20.	VBT	TransitiveVerbs	Makan, tidur, menyanyi
21.	VBI	Intransitive Verbs	Bermain, terdiam, berputar-putar
22.	MD	Modal or Auxiliaries verbs	Sudah, boleh, harus, mesti, perlu
23.	JJ	Adjectives	Mahal, kaya, besar, malas
24.	CDP	Primary cardinal numerals	Satu, juta, milyar
25.	CDO	Ordinal cardinalnumerals	Pertama, kedua ,ketiga
26.	CDI	Irregular cardinal numerals	Beberapa, segala, semua
27.	CDC	Collective cardinal numerals	Bertiga, bertujuh, berempat
28.	NEG	Negations	Bukan, tidak, belum, jangan
29.	IN	Prepositions	Di, ke, dari, pada, dengan
30.	CC	Coordinateconjunction	Dan, atau, tetapi

31.	SC	Subordinate conjunction	Yang, ketika, setelah
32.	RB	Adverbs	Sekarang ,nanti, sementara ,sebab, sehingga
33.	UH	Interjections	Wah, aduh, astaga, oh, hai
34.	DT	Determiners	Para, ini, masing-masing, itu
35.	WDT	WH-determiners	Apa, siapa, barangsiapa
36.	RP	Particles	Kan, kah, lah ,pun
37.	NN	Common nouns	Buku, rumah, karyawan, air, gula, rumahnya, kambingmu
38.	VB	Verbs	Makan, tidur, menyanyi, bermain, terdiam, berputar-putar
39.	FW	Foreignwords	Absurd, deadline, list, science

C. Kata Ambigu

Menurut Kamus Besar Bahasa Indonesia (KBBI) kata ambigu adalah kata yang memiliki makna lebih dari satu sehingga menimbulkan keraguan, kekaburan, ketidakjelasan, dan sebagainya. Terdapat tiga jenis ambiguitas diantaranya : ambiguitas fonetik, gramatikal, leksikal.

- ambiguitas fonetik adalah kata ambigu yang terjadi karena kesamaan bunyi-bunyian yang diucapkan, contoh : “ani datang kerumah memberi tahu” , kalimat tersebut dapat ditafsirkan sebagai tahu jajanan yang terbuat dari kedelai, atau tahu yang berarti,memberi informasi.
- ambiguitas grametikal adalah kata ambigu yang terbentuk karena adanya pembentukan tata bahasa dalam kata, frasa, dan kalimat. contohnya: " Orang tua", frasa ini mengandung arti ayah dan ibu, atau orang yang telah berumur.
- ambiguitas leksikal adalah kata yang terjadi karena faktor kata itu sendiri. hal ini dikarenakan pada dasarnya suatu kata dapat memiliki lebih dari satu arti tergantung penempatan katanya. Contohnya Budi bisa mengeluarkan bisa ular, kata bisa yang Pertama bermakna, “bisa” dapat diartikan sebagai bisa atau dapat sedangkan dalam kata kedua bisa diartikan bisa sebagai bisa ular [12].

3. METODE PENELITIAN

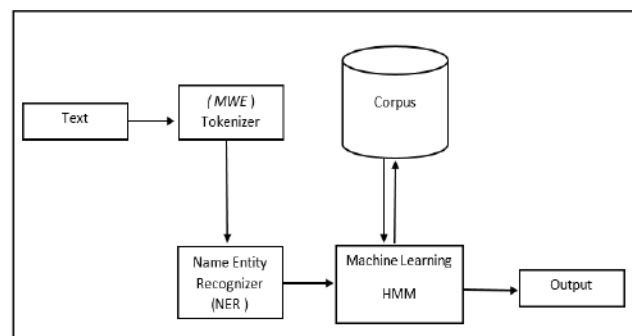
A. Bahan dan Materi Penelitian

Penelitian ini membutuhkan 3 macam data. Data yang pertama adalah data kelas kata atau *tagset*, data kedua adalah data *training* berupa *tagged corpus* atau *corpus* yang sudah diberi kelas katanya dan data ketiga berupa kalimat pertanyaan yang diberikan oleh *user*. *Corpus* yang digunakan sebagai data training pada penelitian ini merupakan korpus dari hasil penelitian yang telah dilakukan oleh Universitas Indonesia sebagai salah satu wakil dari Indonesia dalam proyek Pan Localization (PANL10N, <http://www.panl10.net>) yang berisikan kurang lebih 903.180 kata yang telah di berikan label kelas katanya. Pada awalnya korpus ini masih tersimpan dalam file dengan format teks (*.txt) yang berisikan kata-kata dalam susunan kalimat dan sudah diberikan pelabelan kelas

katanya. Pada penelitian ini korpus tersebut diubah formatnya menjadi format AIML dengan tujuan untuk mempercepat akses dan dapat diakses pada aplikasi mobile berbasis android. Selain itu data korpus juga ditambahkan sekitar 2.200 kata bidang swamedikasi penyakit ringan yang telah tagging dalam bentuk kalimat yang sering digunakan pada dialog. Kalimat seputar swamedikasi sebagian diambil dari Buku Pedoman Penggunaan Obat bebas dan Bebas Terbatas tahun 2007, hasil wawancara dengan dokter, apoteker, tenaga medis, dan sebagian yang lain dari website www.alodokter.com. Sedangkan acuan kelas kata atau *tagset* yang digunakan yaitu *tagset* oleh penelitian [10]. Kelas kata atau *tagset* yang digunakan untuk penelitian ini dapat dilihat pada tabel 2.1.

B. Rancangan Arsitektur Sistem

Pada penelitian yang dilakukan, peneliti berfokus pada *POS Tagging* dan corpus yang digunakan pada aplikasi SMIQA. Probabilitas yang tertinggi terhadap kata tersebut akan dipilih menjadi kelas kata atau *tag* dari kata tersebut. Berikut arsitektur *POS Tagging* yang dibangun dapat dilihat pada gambar.



Gambar 2. Arsitektur POS tagging

1. Inputan yang diberikan berupa text kata dari sebuah dialog yang telah dinormalisasi.
2. *Multi-word Expression (MWE) Tokenizer*

MWE Tokenizer yaitu proses pemotongan kata yang terdiri dari satu atau lebih kata namun memiliki arti yang sama [12]. Pada tahapan ini kalimat yang diinputkan oleh pengguna SMIQA akan di pecah atau di split berdasarkan tanda “spasi” dengan mempertimbangkan tanda baca seperti tanda tanya, tanda seru sebagai pembatas dialog dan dengan memperhatikan apakah kata tersebut termasuk kata frase atau kata bukan frase [12].

Aturan yang dibuat untuk proses pemotongan kata, dimana aturan tersebut adalah sistem akan mengecek setiap *token* pada *corpus*, jika kata tersebut terdapat pada *corpus* maka sistem akan memeriksa kata selanjutnya sampai kata selanjutnya yang tidak cocok dalam membentuk kata frase dan jika kata cocok dalam membentuk kata frase maka sistem akan memberikan token kedalam satu *token*. Jika *token* tidak terdapat pada *corpus*, maka sistem akan memberikan *token* terpisah berdasarkan spasi. Sehingga akan didapatkan token (kata, angka, atau tanda baca) yang terdiri atas *multi-word expression*. *Multi-word expression* adalah token yang terdiri atas lebih dari satu kata akan tetapi memiliki satu makna, seperti “rumah sakit”. Berikut contoh hasil proses *MWE Tokenizer*.

Kalimat : **Apa efek samping dari paracetamol ?**

Hasil Tokenisasi berdasarkan spasi menghasilkan kata : “**apa**”, “**efek**”, “**samping**”, “**dari**”, “**paracetamol**”.

Pada kalimat tersebut terdapat satu frase yaitu “**efek samping**” yang terdiri atas dua kata dengan satu makna.

3. *Name Entity Recognizer*

Setelah proses tokenisasi, maka tahap selanjutnya adalah tahap pengenalan entitas atau yang sering disebut sebagai benda atau objek, apakah dua kata dalam satu makna seperti nama tempat dan yang lainnya.

4. *Machine Learning*

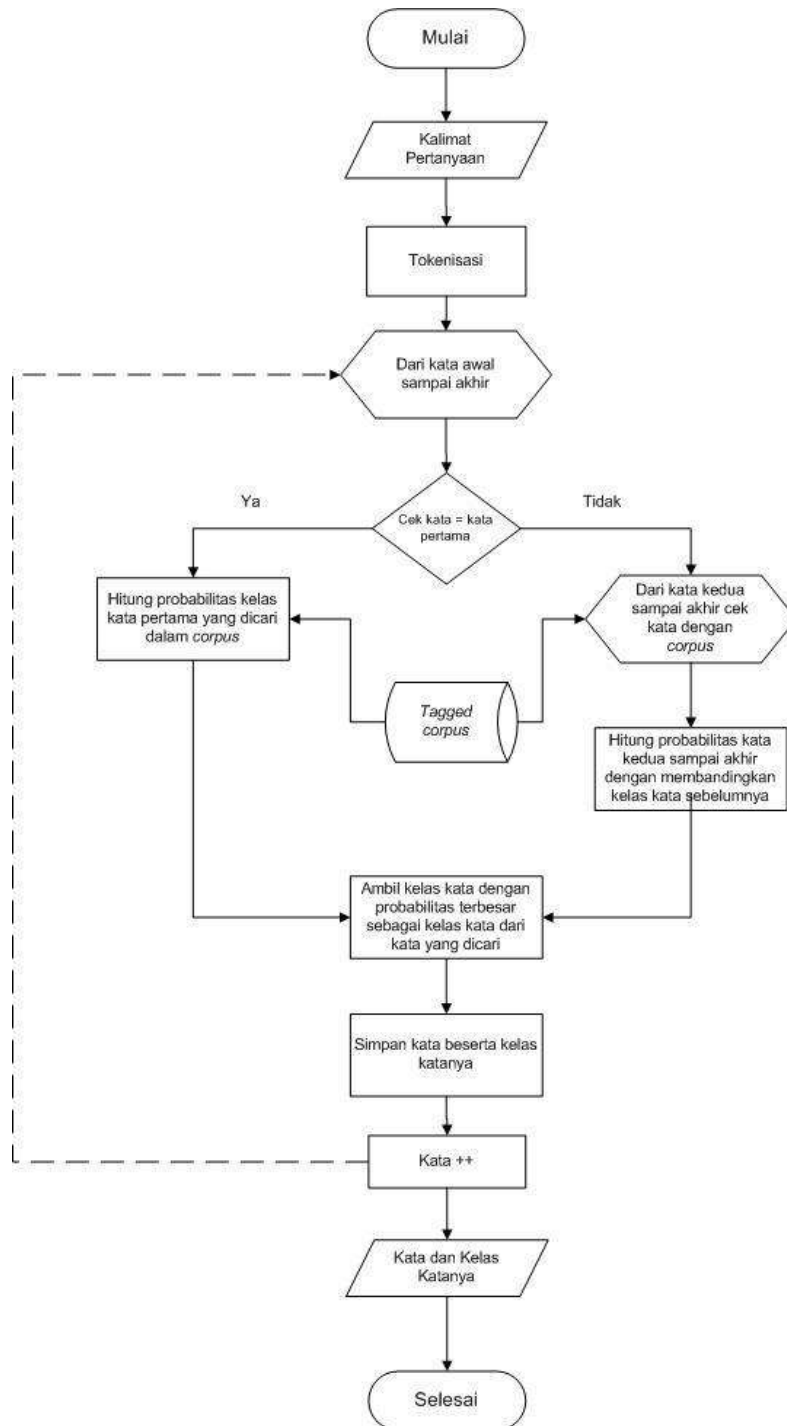
Pada tahapan ini akan dilakukan proses pengelompokkan kelas kata dari setiap token. Mekanisme pengelompokan dilakukan dengan cara, token akan di kelompokkan pada kelompok kata khusus, dengan cara mencari pada kamus kata (korpus SMIQA) yang sudah di tagging. Hasil dari proses ini memungkinkan terjadinya satu token/kata memiliki lebih dari satu tag sehingga token memiliki label ambigu.

Contoh : Mengapa Bisa terjadi Cacingan ?

Hasil : Mengapa Bisa Terjadi cacingan
wp md, nn vbi nn

Selanjutnya *Hidden Markov Model* akan menangani permasalahan token yang memiliki lebih dari satu tag atau disebut ambigu.

Menurut [14] berikut adalah gambaran kerja proses *Hidden Markov Model* :



Gambar 3. Flow chart HMM

Perhitungan probabilitas diawali dengan menghitung probabilitas kata pertama, kelas kata kedua dan seterusnya juga akan dihitung dengan melihat kelas kata sebelumnya [14]. Untuk kata yang tidak ada didalam *corpus*, *output* yang dihasilkan berupa *null*. Dari hasil perhitungan probabilitas tersebut untuk menentukan kelas kata dari kata yang dicari, maka dipilih kelas kata dengan perhitungan nilai probabilitas yang terbesar. Pada kalimat “Mengapa bisa terjadi cacingan” untuk mencari *tag* apa yang tepat di setiap kata dilakukan perhitungan probabilitas. Pada *corpus* untuk kata mengapa, bisa, terjadi dan cacingan. Dari

sebuah kalimat tersebut ada yang memiliki *tag* lebih dari satu *tag*, mengapa memiliki *tag* wp, bisa memiliki *tag* md dan nn, terjadi memiliki *tag* vb dan kata cacingan memiliki *tag* nn. Untuk itu Setiap kata perlu diuji untuk mendapatkan *tag* dengan nilai perhitungan tertinggi supaya *tag* yang didapat tepat, terhadap kata tersebut dengan memperhatikan 1 *tag* sebelum kata yang dicari, pengujian *tag* sebanyak 21 *tagset* dengan melakukan perhitungan probabilitas dan memilih nilai probabilitas tertinggi untuk dijadikan *tag* masing-masing kata dari kalimat uji tersebut, sehingga dihasilkan kata beserta *tag* atau kelas katanya masing-masing seperti “mengapa/wp bisa/md terjadi/vbi cacingan/nn”.

5. Output

Setelah semua proses telah dilakukan tahap akhir adalah keluaran atau output berupa kata yang telah memiliki *tag* masing masing. Berikut contoh hasil output nantinya :

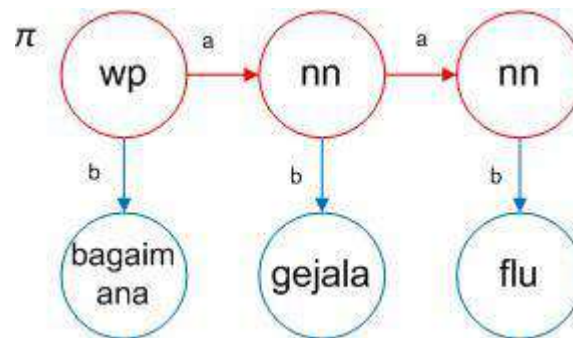
Kalimat : Mengapa Bisa Terjadi Kecacingan
 wp md vbi nn

4. HASIL DAN PEMBAHASAN

A. Pembahasan

Pada penelitian ini perhitungan probabilitas pelabelan kata menggunakan asumsi bigram yaitu probabilitas suatu kemunculan *tag* hanya bergantung dari *tag* sebelumnya, dengan menggunakan teori *Bayessian Interpretation*. Dengan menggunakan persamaan tersebut dapat diperoleh hasil perhitungan probabilitas untuk setiap kata yang dicari kelas katanya. Proses dimulai dengan menghitung probabilitas kelas kata yang ada di dalam *corpus* menggunakan metode HMM. Setiap kata akan dicari probabilitasnya terhadap kelas kata sebelumnya yang ada didalam *corpus*. Perhitungan probabilitas diawali dengan menghitung probabilitas kata pertama, kelas kata kedua dan seterusnya juga akan dihitung dengan melihat kelas kata sebelumnya. Untuk kata yang tidak ada didalam *corpus*, *output* yang dihasilkan berupa *null*. Dari hasil perhitungan probabilitas tersebut untuk menentukan kelas kata dari kata yang dicari, maka dipilih kelas kata dengan perhitungan nilai probabilitas yang terbesar.

Diberikan contoh berupa kalimat yaitu “**bagaimana gejala flu**” untuk mencari *tag* apa yang tepat di setiap kata. Pada *corpus* untuk kata bagaimana, gejala dan flu masing-masing hanya memiliki 1 *tag*, bagaimana memiliki *tag* wp, gejala memiliki *tag* nn dan kata flu memiliki *tag* nn. Setiap kata perlu diuji untuk mendapatkan *tag* apa yang tepat terhadap kata tersebut dengan memperhatikan 1 *tag* sebelum kata yang dicari, pengujian *tag* sebanyak 21 *tagset* dengan melakukan perhitungan probabilitas dan memilih nilai probabilitas tertinggi untuk dijadikan *tag* masing-masing kata dari kalimat uji tersebut. Berikut adalah rantai *markov* dari kalimat uji pada gambar berikut:



Gambar 3.6 Rantai Markov Kalimat Uji

Adapun langkah perhitungan probabilitas dari kalimat uji adalah sebagai berikut :

1. Bagaimana

Kata pertama yaitu “bagaimana”, kata “bagaimana” akan diuji untuk mendapatkan semua nilai probabilitas terhadap kata “bagaimana” tersebut. *Tagset* yang dilakukan dalam pengujian yaitu 21 *tag* yang terdiri dari ",", ".", ":", "nn", "nnc", "nnu", "nng", "nnp", "wp", "vbt", "vbi", "md", "jj", "neg", "in", "cc", "sc", "rb", "dt", "cdp" dan "cdi". Pada *corpus* kata “bagaimana” hanya memiliki *tag* wp sehingga probabilitas terhadap *tag* lain adalah 0 dan probabilitas terbesar terdapat pada *tag* “wp” dengan hasil 0.00428571 dari perhitungan *tag* wp jika *tag* sebelumnya adalah . (titik).

2. Gejala

Kata kedua yaitu “gejala” kata “gejala” juga akan diuji untuk mendapatkan semua nilai probabilitas terhadap kata “gejala” tersebut. *Tagset* yang dilakukan dalam pengujian yaitu 21 *tag* yang terdiri dari ",", ".", ":", "nn", "nnc", "nnu", "nng", "nnp", "wp", "vbt", "vbi", "md", "jj", "neg", "in", "cc", "sc", "rb", "dt", "cdp" dan "cdi". Pada *corpus* kata “gejala” hanya memiliki *tag* nn sehingga probabilitas terhadap *tag* lain adalah 0 dan probabilitas terbesar terdapat pada *tag* nn dengan hasil 0.00450593 dari perhitungan *tag* nn jika *tag* sebelumnya adalah nn.

3. Flu

Kata ketiga yaitu “flu”, kata “flu” akan diuji untuk mendapatkan semua nilai probabilitas terhadap kata “flu” tersebut. *Tagset* yang dilakukan dalam pengujian yaitu 21 *tag* yang terdiri dari ",", ".", ":", "nn", "nnc", "nnu", "nng", "nnp", "wp", "vbt", "vbi", "md", "jj", "neg", "in", "cc", "sc", "rb", "dt", "cdp" dan "cdi". Pada *corpus* kata “flu” hanya memiliki *tag* nn sehingga probabilitas terhadap *tag* lain adalah 0 dan probabilitas terbesar terdapat pada *tag* nn dengan hasil 0.00285799 dari perhitungan *tag* wp jika *tag* sebelumnya adalah in.

B. Hasil Pengujian

Pengujian dilakukan dengan menentukan tingkat akurasi yang didapatkan oleh sistem dengan pengujian recal dan precision [15].

1. Recall adalah proporsi jumlah dokumen yang dapat ditemukan-kembali oleh sebuah proses pencarian di sistem POS Tagging. Rumusnya: Jumlah dokumen relevan yang ditemukan / Jumlah semua dokumen relevan di dalam koleksi.

2. precision adalah proporsi jumlah dokumen yang ditemukan dan dianggap relevan untuk kebutuhan si pencari informasi. Rumusnya: Jumlah dokumen relevan yang ditemukan / Jumlah semua dokumen yang ditemukan.

Tabel 4.1. Hasil pengujian recal dan Precision

Pengujian	Jumlah Kalimat	Jumlah Kata	Total Benar	Total Salah	Akurasi	
					<i>precision</i>	<i>Recall</i>
1	40	71	56	15	78,873	65,116
2	40	252	234	18	92,857	86,667
3	40	249	235	14	94,377	89,354
4	40	210	202	8	96,190	92,661
5	40	243	231	12	95,061	90,588
Rata-rata					91,472	84,877
Standar deviasi					15,435	

Pada tabel 4.1. menunjukan lima hasil pengujian *precision* dan *recall*, masing-masing pengujian berjumlah 40 kalimat. Kalimat uji didapatkan langsung dari kuisioner secara acak yang di bagikan kepada beberapa *respondent*. Pengujian dilakukan dengan dua cara yaitu pengujian *precision* dan *recall*. Hasil tertinggi pengujian *precision* sebesar 96,190 % sedangkan pada pengujian *recall* sebesar 92,661% dengan standar deviasi sebesar 3,651. Dengan nilai standar deviasi tersebut dapat di simpulkan bahwa data uji yang di gunakan persebaran datanya merata. Artinya, lima data uji sudah dapat mewakili pengujian karena nilai akurasinya tidak berbeda jauh dari nilai rata-rata akurasi kedua pengujian.

5. KESIMPULAN

Kesimpulan dari penelitian ini adalah pelebelen kelas kata terhadap kata dalam bidang swamedikasi menggunakan *metode Hidden Markov Model* memiliki hasil tingkat akurasi yang tinggi, yaitu hasil tertinggi dengan pengujian *precision* 96,190 % dan dengan pengujian *recall* hasil tertinggi 92,661% .

KEPUSTAKAAN

- [1] Alwi , H, Dardjowidjojo , S , Lapoliwa , H , Moeliono , A M. 2003 .Tata Bahasa Baku Bahasa Indonesia.Balai Pustaka. Jakarta,Indonesia.
- [2] Jurafsky, D. 2000, *Speech and Language Processing An Introduction to Natural Language Processing, Computational Linguistics, and Speech recognition*, Prentice-Hall Inc, New Jersey.
- [3] Manurung, R. 2016, *Tutorial: Pengenalan terhadap POS tagging dan probabilistic Parsing*, Workshop Nasional INACL, Jakarta.
- [4] Kelly D., Kantor, P. B., Morse, E. L., Scholetz, J. dan Sun, Y., 2009, *Questionnaires for eliciting evaluation data from users of interactive question answering systems*. Natural Language Engineering, 15 , pp 119-141 doi:10.1017/S1351324908004932.

- [5] Direktur Bina Penggunaan Obat Rasion., 2008. *Materi Pelatihan Peningkatan Pengetahuan Dan Keterampilan Memilih Obat Bagi Tenaga Kesehatan*, Departemen Kesehatan RI.
- [6] Wibisono, Y. 2008. *Penggunaan Hidden Markov Model untuk Kompresi Kalimat*. Thesis, Program Magister Informatika, Institut Teknologi Bandung.
- [7] Wicaksono, A.,F dan Purwanti A. 2010, *HMM Based POS Tagger from Bahasa Indonesia*, *Fourth International MALINDO Workshop*, Institut Teknologi Bandung, Bandung.
- [8] Ratnaparkhi, A . 1996, *A Maximum Entropy Model for Part-Of-Speech Tagging*, Universitas Of Pennsylvania,
- [9] Silverberg, M, Ruokolainen, T, linden, K, & Kurimo, M. 2014, *Part-Of-Speech Tagging Using Conditional Random Field : Exploiting Sub-Label Dependencies for Improved Accuracy*, University of Helsinki.
- [10] Pisceldo, F. Adriani, M, & Manurung, R. 2009. *Probabilistic Part Of Speech tagging for Bahasa Indonesia*, *Faculty of Computer Science*, Universitas Indonesia, Depok
- [11] Anh Tuan,L., Yasuhide, Y., Ohkuma, T. 2016, *Sentiment Analysis for Low Resource Language: A study on Informal Indonesia Tweet*, Nanyang Technological University Singapore, Singapore
- [12] Rashel F., Luthfi, A, Dinakaramani, A & manurung, R. 2014 .*Building an indonesia rule based part of speech tagger*, Universitas Indonesia, Depok.
- [13] Christiani, V., M. Jeanny, P and Endah, P. 2012. *Implementasi Brill Tagger Untuk Memberikan POS-Tagging Pada Dokumen Bahasa Indonesia*. vol. 01, no. 03, Universitas Tarumanegara.
- [14] Widhiyanti, K & Harjoko, A. 2012. *POS Tagging Bahasa Indonesia Dengan HMM dan Rule Based*, Vol 8, No 2, Universitas Gajah Mada, Yogyakarta.
- [15] Tague-Sutcliffe, J.M., “Some Perspective on the Evaluation of Information Retrieval System”, *Journal of the American Society for Information Science*, 47(1),1996 : 1-3.
- [16] Chong-Shyong Ong, Min-Yuh Day, Wen-Lian Hsu, 2009, *The measurement of user satisfaction with question answering systems*, *Journal of Information & Management*, pp. 397–403, doi: 10.1016/j.im.2009.07.004

SISTEM TEMU KEMBALI CITRA DAUN TUMBUHAN MENGUNAKAN METODE EIGENFACE

Supiyanto dan Samuel A. Mandowen

Program Studi Sistem Informasi, Jurusan Matematika
Fakultas Matematika dan Ilmu Pengetahuan Alam
Universitas Cenderawasih
Email : supi6976@gmail.com

ABSTRAK

Sistem temu kembali citra daun merupakan suatu aplikasi yang dapat digunakan menemukan kembali citra daun dengan cara melakukan perbandingan antara citra uji dengan citra latih yang ada di database berdasarkan informasi yang ada pada citra tersebut. Metode yang digunakan untuk menemukan kembali citra daun berdasarkan kemiripan warna, bentuk, dan tekstur pada penelitian ini adalah metode eigenface. Data citra yang digunakan pada penelitian ini berupa citra daun sebanyak 50 citra berasal dari 10 tumbuhan yang berbeda. Warna citra berformat grayscale, type citra berformat bitmap (bmp) dengan ukuran 100 x 100 pixel. Sedangkan untuk menentukan kemiripan antara citra uji dengan citra latih yang ada di database menggunakan formula *euclidean distance*. Hasil pengujian, sistem dapat mengenali citra uji dengan benar dengan tingkat akurasi mencapai 100% untuk citra uji sama dengan citra latih, 93% untuk citra uji dan citra latih berasal dari satu helai daun yang sama tetapi berbeda secara posisi dan kondisi dan akurasi dibawah 70% untuk citra uji berbeda dengan citra latih walau berasal dari pohon atau tanaman yang sama.

Kata Kunci : Citra, daun, *Euclidean*, *pixel*, *eigenface*, *grayscale*.

I. Pendahuluan

Daun tumbuhan merupakan salah satu biometrik dari tumbuhan. Hal ini dikarenakan daun pada setiap jenis tumbuhan memiliki bentuk dan tulang daun yang berbeda-beda. Struktur bentuk tulang daun merupakan salah satu fitur unik yang dimiliki daun yang memiliki peranan penting dalam klasifikasi jenis tumbuhan melalui serangkaian proses pengolahan citra. [1]

Dewasa ini, pengenalan citra daun tumbuhan secara otomatis masih menjadi perhatian banyak peneliti. Hal ini dikarenakan dengan semakin berkembangnya teknologi yang digunakan pada masa sekarang ini. Pengenalan citra daun tumbuhan memiliki peranan penting dalam pengenalan jenis-jenis tumbuhan. Dengan sistem pengenalan citra daun yang masih manual, akan mengakibatkan kurang efisiennya tenaga dan ruang penyimpanan hasil. Kurang efisien pada tenaga yaitu, dibutuhkan tenaga ahli atau yang faham mengenai bentuk-bentuk tulang daun, sedangkan kurang efisiennya ruang penyimpanan hasil yaitu, dengan proses yang masih manual, maka kita membutuhkan tempat penyimpanan data yang besar untuk menyimpan hasil identifikasi yang telah dilakukan. [2]

Pada prosesnya, menemukan kembali citra daun tidak mudah dilakukan. Selain ditentukan oleh metode yang akan digunakan, pemilihan fitur tekstur harus tepat sehingga dapat membedakan antara suatu citra dengan citra yang lain. Secara umum fitur tekstur dapat diperoleh dari penerapan operator lokal, analisis statistik dan perhitungan transformasi asal (*domain*). Tetapi, penggunaan banyak parameter dalam identifikasi dapat menurunkan akurasinya, sehingga sangat penting untuk memilih fitur yang tepat [3].

Metode *Principal Component Analysis* (PCA) atau dikenal dengan metode *eigenface* merupakan teknik statistik yang banyak digunakan untuk proses pengenalan pola pada suatu objek. Pola yang dihasilkan adalah citra ciri dengan dimensi data yang lebih kecil karena metode *eigenface* pada dasarnya adalah bertujuan untuk menyederhanakan variabel yang diamati dengan cara mereduksi dimensinya tetapi dengan tidak kehilangan banyak informasi penting yang terkandung di dalamnya [4].

Penelitian berhubungan dengan metode *eigenface* diantaranya [5], melakukan penelitian tentang pengenalan citra wajah orang Papua dengan metode *eigenface* yang mampu mengenali dengan benar sampai 96%.

Sesuai dengan latar belakang masalah di atas, maka peneliti tertarik membuat penelitian adalah bagaimana membuat suatu sistem temu kembali citra daun tumbuhan menggunakan metoda *eigenface* dan ingin mengetahui seberapa akurat metode *eigenface* dapat menemukan kembali citra daun dengan benar.

II. METODE PENELITIAN

Sistem yang dibangun adalah sistem temu kembali citra daun yang dilakukan dengan mengenali pola dari citra daun tersebut. Sistem temu kembali citra daun didasari pada pengenalan pola dengan metode *eigenface*. Proses sistem temu kembali citra daun dilakukan dengan membandingkan citra query dengan citra latih yang ada di database.

Secara garis besar, penelitian ini terdiri dari beberapa tahap yaitu:

A. Studi Literatur

Penulisan ini dimulai dengan studi literature atau studi pustaka. Studi pustaka bertujuan untuk mencari teori-teori yang mendukung mengenai teori yang digunakan dalam penelitian ini seperti teori tentang pengolahan citra, metode *eigenface* atau PCA (*Principal Component Analysis*), Jarak Euclidean, dan lain sebagainya

B. Pengumpulan Data

Pada tahap ini dilakukan pengumpulan data citra daun tanaman yang akan digunakan sebagai data penelitian. Selanjutnya data citra ini akan di bagi atas 2 bagian yakni data training dan data testing.

C. Proses pengolahan citra

Setelah pengambilan data citra, tahap selanjutnya adalah tahap pengolahan citra sebelum data tersebut digunakan untuk penelitian. Berikut ini proses pengolahan citra.

1. Tahap *pre-processing*

Proses yang dilakukan tahap ini adalah proses *cropping*. Tujuannya adalah mengambil bagian citra yang diperlukan yaitu bagian tulang daunnya saja dengan ukuran citra menjadi 100 pixel × 100 pixel. Proses *cropping* dilakukan dengan menggunakan program aplikasi photoshop.

2. Tahap pengubahan format warna citra

Proses yang dilakukan pada tahap ini adalah pengubahan citra warna menjadi citra keabu-abuan (*greyscale*) dengan type *Bitmap* (BMP). Proses ini juga dilakukan dengan menggunakan program aplikasi photoshop.

D. Analisa Masalah

Tahap ini juga berisikan analisa dari permasalahan yang akan diangkat dalam penelitian ini. Pada penelitian ini, permasalahan yang akan diteliti yaitu bagaimana membangun sistem temu kembali citra daun tumbuhan menggunakan metode eigenface.

E. Merancangan Sistem

Pada tahap ini, akan dirancang untuk mengembangkan program aplikasi perangkat lunak sistem temu kembali citra daun tumbuhan menggunakan metode eigenface. Program aplikasi perangkat lunak ini nantinya akan digunakan untuk menentukan apakah citra yang akan diujikan dapat dikenali dengan benar atau tidak. Secara garis besar untuk membangun program aplikasi perangkat lunak ini dilakukan dengan dua tahapan proses yaitu tahap pelatihan dan tahap pengenalan.

F. Menganalisis Algoritma

Tahap analisa ini bertujuan untuk menganalisa algoritma-algoritma yang berhubungan dengan penelitian ini. Adapun algoritma - algoritma pada sistem temu kembali citra daun tumbuhan menggunakan metode eigenface adalah sebagai berikut [6]:

Misalkan diasumsikan :

- Terdapat M samples image dengan ukuran tinggi = ukuran lebar = N .
- Banyaknya citra di dalam database adalah M
- Banyaknya orang di dalam database adalah h , $\{\vec{a}, \vec{b}, \vec{c}, \dots, \vec{\square}\}$

a. Algoritma untuk pelatihan

Tahap algoritma pelatihan dilakukan guna mengekstraksi ciri-ciri dari beberapa citra daun pelatihan yang terdapat pada database. Adapun beberapa langkah dari algoritma pelatihan yakni :

1. Konversi image $I_{1,2,3,\dots,M}$ dengan ukuran $N * N$ ke dalam bentuk vektor kolom.

$$\begin{array}{cccc}
 \begin{array}{c} \text{Image 1} \\ \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 2} \\ \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 3} \\ \begin{pmatrix} c_1 \\ c_2 \\ \vdots \\ c_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 4} \\ \begin{pmatrix} d_1 \\ d_2 \\ \vdots \\ d_{n^2} \end{pmatrix} \end{array} \\
 \begin{array}{c} \text{Image 5} \\ \begin{pmatrix} e_1 \\ e_2 \\ \vdots \\ e_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 6} \\ \begin{pmatrix} f_1 \\ f_2 \\ \vdots \\ f_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 7} \\ \begin{pmatrix} g_1 \\ g_2 \\ \vdots \\ g_{n^2} \end{pmatrix} \end{array} &
 \begin{array}{c} \text{Image 8} \\ \begin{pmatrix} h_1 \\ h_2 \\ \vdots \\ h_{n^2} \end{pmatrix} \end{array}
 \end{array}$$

$\vec{a} = [a_1, a_2, \dots, a_{N^2}]$ begitu juga $\{\vec{b}, \vec{c}, \dots, \vec{h}\}$.

2. Menghitung rata-rata semua citra m :

$$\vec{m} = \frac{\vec{a} + \vec{b} + \dots + \vec{h}}{M},$$

$$\vec{A} = [\vec{a} - \vec{m}, \vec{b} - \vec{m}, \dots, \vec{h} - \vec{m}] = [\vec{a}_m, \vec{b}_m, \dots, \vec{h}_m] \quad m = 1, 2, 3, \dots, M$$

3. Menghitung kovariance (cov) matrik dan matrik yang lain (L) :

$$Cov = \vec{A}\vec{A}^T \quad L = \vec{A}^T \vec{A}$$

4. Menentukan nilai eigen dan vektor eigen sebanyak sample yang ada dari matrik **cov** atau **L**.
5. Membentuk matrik **V** yang merupakan matrik eigen vektor dari **L**.
6. Kalikan matrik vektor $\vec{A}$ dengan matrik eigen vektor **V** yang hasilnya merupakan representasi dari variasi daun (**U**).
7. Tentukan ruang daun dari tiap-tiap citra sebanyak sample yang digunakan dengan menggunakan rumus :

$$\vec{\Omega} = U^T (\vec{A})_M, \quad m = 1, 2, \dots, M$$

8. Hitung nilai thresholdnya / toleransi θ

$$\theta = \frac{1}{2} \max \|\Omega_j - \Omega_i\| \quad \text{for } i = 1 \dots M$$

b. Algoritma untuk Pengenalan / pengujian

langkah-langkah proses pengujian/pengenalan, secara matematis dapat dituliskan sebagai berikut :

1. Konversi citra uji / masukan $r_{1, 2, \dots}$ dengan ukuran $N \times N$ ke dalam bentuk vektor dengan ukuran panjang N^2 .



$$= \begin{pmatrix} r_1 \\ r_2 \\ \vdots \\ r_{N^2} \end{pmatrix}$$

2. Hitung Matriks Beda $\vec{r}_m$ dengan mengurangkan dengan nilai rata-rata :

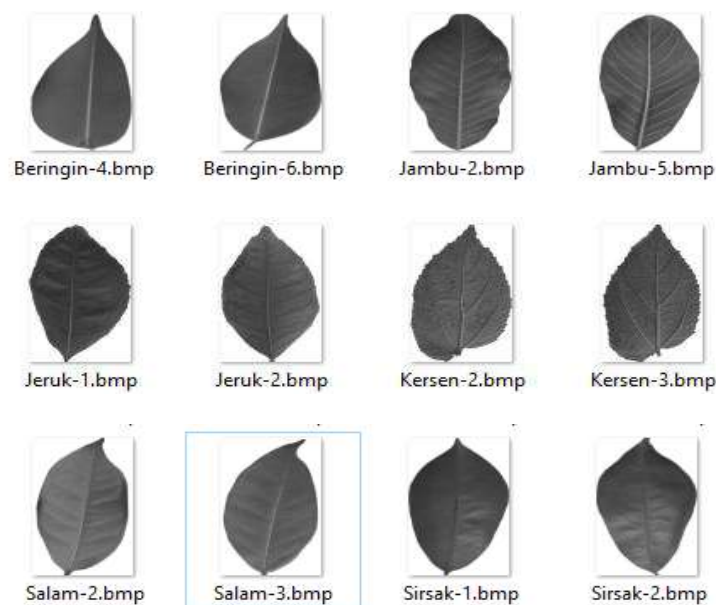
$$\vec{r}_m = \begin{pmatrix} r_1 - m_1 \\ r_2 - m_2 \\ \vdots \\ r_{N^2} - m_{N^2} \end{pmatrix}$$

3. Menghitung bobot pada citra uji dengan cara mengalikan matrik vektor eigen transpose $\vec{V}^T$ dengan matrik $\vec{r}_m$ yang hasilnya matriks (U_{inp}) ; $U_{inp} = \vec{V}^T \vec{r}_m$
4. Menghitung jarak perbedaan antara citra *testing* dengan citra daun *training* dengan menggunakan *euclidian minimum*.
$$\varepsilon_i = \sqrt{\|U - U_{inp}\|^2}; \quad i = 1 \dots M$$
5. Hasil indentifikasinya adalah citra latih yang memiliki jarak terkecil dengan citra uji yang ditampilkan oleh sistem.

III. HASIL DAN PEMBAHASAN

A. Data Citra

Sebagaimana disebutkan di atas bahwa data yang digunakan dalam penelitian ini adalah data primer, yaitu data daun tanaman yang dijadikan sampel penelitian, dipotret oleh peneliti. Daun tanaman diambil dari 10 jenis tanaman yng berbeda, dengan setiap tanaman diwakili minimal 5 daun. Di bawah ini beberapa contoh sampel citra daun yang akan digunakan dalam penelitian.

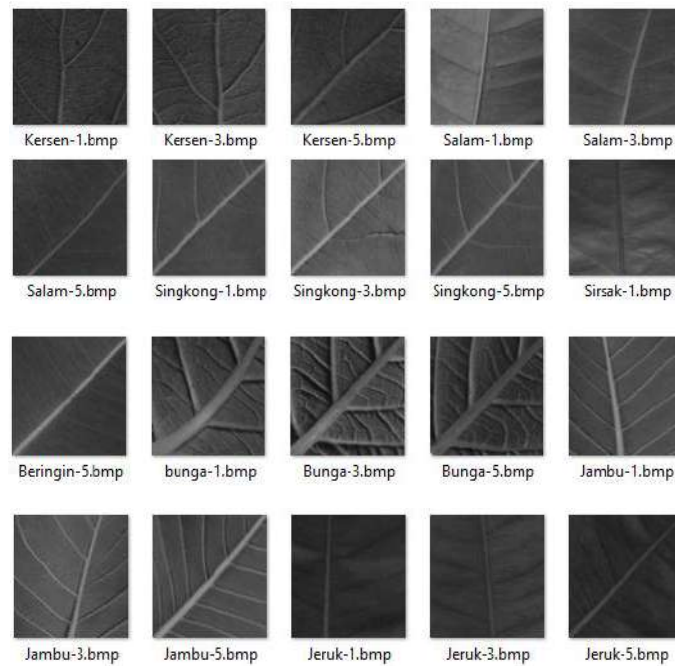


Gambar 2. Contoh sampel data citra daun hasil pemotretan

a. Data Citra Latih (*training*)

Data training adalah sekumpulan data penelitian yang dipakai untuk melatih program, untuk mengetahui karakteristik, model atau pola-pola yang ada pada data tersebut. Berikut ini beberapa contoh citra training yang dipergunakan dalam pelatihan seperti berikut ini :

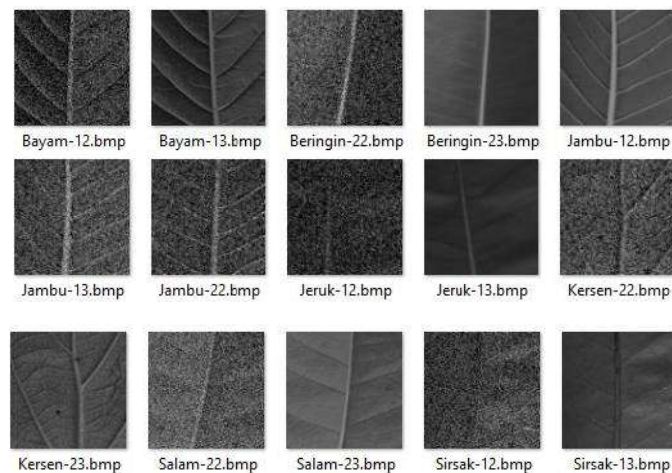




Gambar 3. beberapa contoh citra latih

b. Data Citra Uji (*Testing*)

Tujuan penggunaan data citra uji, untuk mengetahui seberapa akurat sistem dapat mengidentifikasi citra uji dengan benar. Di bawah ini beberapa contoh citra uji yang telah peneliti ujikan pada sistem.



Gambar 4. beberapa contoh citra uji

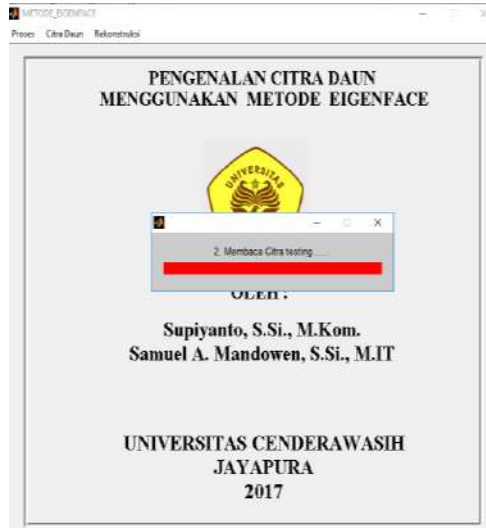
B. Pengujian Sistem

a. Proses Training

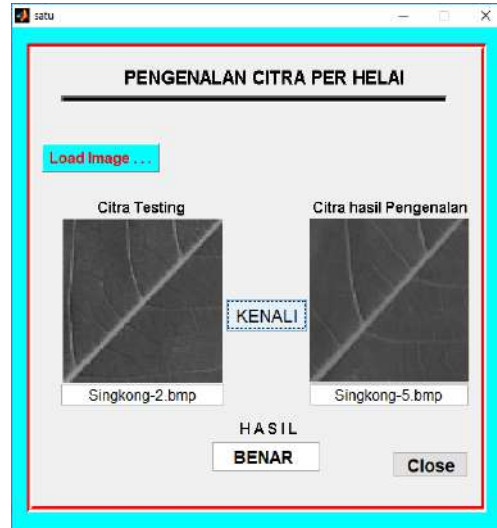
Tahap ini bertujuan untuk menghasilkan nilai bobot dari setiap citra training yang ada. Jika proses training belum dilakukan maka proses pengenalan tidak bisa dilakukan. Beberapa proses yang dilakukan dalam proses training dapat dilihat dalam gambar 5.

b. Proses Pengenalan

Gambar 6, merupakan salah satu contoh proses hasil pengenalan citra uji yang dikenali dengan benar oleh sistem temu kembali citra daun tumbuhan menggunakan metode eigenface.



Gambar 5. beberapa proses yang terjadi ketika pada proses training.



Gambar 6. Contoh pross pengenalan yang dilakukan oleh sistem

c. Hasil Pengujian Program

1. Citra latih sama dengan citra uji

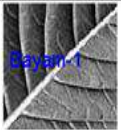
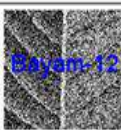
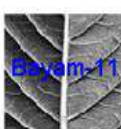
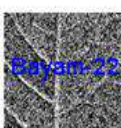
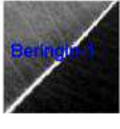
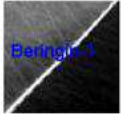
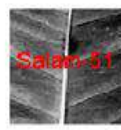
Maksud dari citra latih sama dengan citra uji yaitu citra uji yang digunakan pada proses pengenalan merupakan citra latih itu sendiri. Hasil pengujian menunjukkan bahwa sistem ini berhasil mengenali citra uji dengan benar mencapai 100%. Contoh hasil pengujian terlampir.

2. Citra uji dan citra testing, tapi mash satu daun

Pada pengujian ini, citra uji tidak sama dengan citra latih, tetapi citra uji dan citra latih berasal dari satu helai daun namun berbeda dilihat dari posisi dan kondisi citra;. Hasil pengujian menunjukkan bahwa sistem berhasil mengenali citra uji dengan benar mencapai 93%

3. Citra uji dan citra testing, beda daun

Maksudnya, citra uji tidak sama dengan citra latih tetapi pada pengujian ini citra uji dan citra latih berasal satu pohon tapi beda daun. Hasil pengujian menunjukkan bahwa sistem berhasil mengenali citra uji dengan benar hanya mencapai 60%.

Citra Masukan		Citra Uji ...	Citra Masukan		Citra Uji ...
	dikenai sebagai			dikenai sebagai	
	dikenai sebagai			dikenai sebagai	
	dikenai sebagai			dikenai sebagai	
	dikenai sebagai			dikenai sebagai	
	dikenai sebagai			dikenai sebagai	
	dikenai sebagai			dikenai sebagai	

Gambar 7.Contoh hasil pengujian program

Keterangan : Tulisan warna biru : Citra uji dikenali dengan benar

Tulisan warna merah : Citra uji dikenali dengan salah.

IV. Kesimpulan dan saran

A. Kesimpulan

Dari penelitian ini maka dapat diambil beberapa kesimpulan :

1. Sistem temu kembali citra daun tumbuhan menggunakan metode eigenface yang dibangun ini mampu mengenali citra daun berdasarkan pencocokan kemiripan.
2. Pengujian sistem untuk citra uji yang sama dengan citra training sistem dapat mengenali citra daun dengan benar dengan prosentasi yang diperoleh adalah 100%.
3. Pengujian sistem untuk citra uji yang berbeda dengan citra training tapi masih berasal dari daun yang sama, sistem mengenali citra uji dengan benar dengan prosentasi 93%, sedangkan untuk citra uji dan citra latih yang betul-betul berbeda, sistem mengenali citra uji dengan benar dengan prosentasi dibawah 60%
4. Sebagian besar citra yang gagal dalam pengenalan ini disebabkan faktor tekstur.
5. Pengenalan terhadap citra yang tidak mempunyai wakil dalam pelatihan, akan menyebabkan sistem mengenali sebagai individu yang salah.

B. Saran

Untuk mendapatkan hasil yang lebih baik, ada beberapa hal yang dapat penulis kemukakan sebagai saran yaitu :

1. Diperlukan lebih banyak data citra daun dalam basis data dengan berbagai variasi posisi, dan kondisi, sehingga pengenalan terhadap citra tidak dikenal dapat diminimalkan.
2. Pengambilan keputusan pada proses pengenalan yang digunakan pada penelitian ini hanya memperhitungkan tingkat kemiripan suatu piksel berdasarkan jarak (*euclidean distance*) antar vektor karakteristik daun, sehingga masih perlu diperbaiki dengan memperhatikan juga pola dan tekstur dari sampel (data training) yang dijadikan acuan dalam proses pengenalan.
3. Sistem ini diharapkan dapat dikembangkan lebih lanjut dengan menggabungkan beberapa metode yang dapat mengatasi kegagalan dari system ini karena perbedaan tekstur dan posisi.

KEPUSTAKAAN

- [1] Tjitrosoepomo, G. 2005. *Morfologi Tumbuhan*, Jogjakarta: Gajah Mada University
- [2] Sari dkk., 2014. *Seleksi Fitur Menggunakan Ekstraksi Fitur Bentuk, Warna, Dan Tekstur Dalam Sistem Temu Kembali Citra Daun*, JUTI (Jurnal Ilmiah Teknologi Informasi) by Department of Informatics.
- [3] Arymurthy, A.M. 2010. *Feature/Fitur/Ciri/Object Descriptor*, Faculty of Computer Science : Universitas Indonesia.
- [4] Abdi, H., Williams, L.J. 2010. *Principal Component Analysis*, Wiley Interdisciplinary Reviews: Computational Statistics. 2(4): 433-459.
- [5] Supiyanto dkk., 2009. *Aplikasi Program Delphi dalam Pengenalan Citra Wajah Manusia dengan Menggunakan Metode Eigenface*, Penelitian Hibah Bersaing.
- [6] M. Turk, A. Pentland, 1991, "Eigenfaces for Recognition", *J. Cognitive Neuroscience*, vol. 3, no.1.

EFEKTIVITAS MODEL *PROBLEM SOLVING* DALAM MENINGKATKAN KEMAMPUAN BERFIKIR LANCAR MAHASISWA PADA MATERI pH LARUTAN

Ratu Betta Rudibyani
Dosen Prodi Pendidikan Kimia, FKIP Unila

ABSTRAK

Penelitian ini bertujuan untuk mendeskripsikan efektivitas model Problem Solving pada materi pH larutan dalam meningkatkan kemampuan berpikir lancar. Metode penelitian ini adalah kuasi eksperimen tipe non equivalent control group desain. Kelas sampel ditentukan dengan teknik purposive sampling. Sampel penelitian adalah mahasiswa kelas A dan kelas B Prodi Pendidikan Fisika, tahun 2017/2018. Efektivitas model Problem Solving dilihat berdasarkan perbedaan rata-rata n-Gain yang diperoleh berdasarkan uji t. Rata-rata n-Gain pada kelas A yaitu 0,68 dan kelas B yaitu 0,26. Kesimpulan penelitian ini adalah pembelajaran dengan menggunakan model Problem Solving efektif dalam meningkatkan kemampuan berpikir lancar mahasiswa pada materi pH larutan.

Kata kunci : Efektivitas, problem solving, pH larutan

1. PENDAHULUAN

Kemajuan suatu negara tidak terlepas dari kualitas sumber daya manusia yang dimiliki. Sumber daya manusia yang berkualitas diperoleh melalui proses pendidikan. Tujuan pendidikan yang tercantum dalam Undang-Undang Nomor 20 Tahun 2003 tentang Sistem Pendidikan Nasional adalah untuk mewujudkan peserta didik yang dapat mengembangkan potensi dirinya melalui suasana belajar dan proses pembelajaran yang direncanakan. Peran pendidik dalam pengembangan potensi diri peserta didik sangatlah penting, sebab lemahnya kemampuan guru/ dosen dalam pengendalian kelas dan tidak sesuainya model pembelajaran yang digunakan dapat membuat suasana belajar dan proses pembelajaran menjadi tidak efektif sehingga peserta didik kurang tertarik mengikuti pembelajaran dan tidak dapat menyelesaikan masalah dan mengembangkan potensi yang dimilikinya.

Fakta di lapangan menunjukkan bahwa pembelajaran IPA kimia di Program Studi Pendidikan Fisika, FKIP Unila, belum menekankan pada pemberian pengalaman langsung kepada mahasiswa. Proses pembelajaran masih berpusat pada dosen dengan metode ceramah, tanya jawab dan pemberian tugas kepada mahasiswa. Pada materi-materi tertentu seperti pH larutan, belum mengajak mahasiswa untuk berpikir tingkat tinggi, khususnya berpikir lancar. Hal tersebut yang menyebabkan mahasiswa belum mampu untuk membangun konsep sendiri.

Berdasarkan permasalahan tersebut, perlu dilakukan perubahan proses pembelajaran IPA kimia agar mahasiswa dapat mencapai kompetensi yang diinginkan. Perubahan tersebut antara lain dengan menerapkan model pembelajaran yang dapat melatih keterampilan berpikir mahasiswa sehingga diharapkan dapat merefleksikan pengetahuan yang dimiliki untuk memecahkan masalah dalam menentukan pH larutan. Salah satu model yang dapat diterapkan yaitu model *problem solving*.

Model *problem solving* merupakan model pembelajaran berbasis masalah dimana mahasiswa dituntut untuk menemukan sendiri pemecahan masalahnya. Model *problem solving* dapat memberikan kesempatan kepada mahasiswa untuk bereksplorasi mengumpulkan dan menganalisis data secara lengkap untuk memecahkan masalah yang dihadapi sehingga proses pembelajaran menjadi lebih aktif. Langkah-langkah model *problem solving* antara lain ada masalah yang jelas untuk dipecahkan, mencari data untuk memecahkan masalah tersebut, menetapkan jawaban sementara dari masalah tersebut, menguji kebenaran jawaban sementara tersebut dan menarik kesimpulan [1]. Berdasarkan langkah-langkah pembelajaran tersebut maka kemampuan mahasiswa untuk berpikir kritis, analisis, sistematis dan logis dalam memecahkan masalah secara ilmiah akan sering dilatihkan.

Berpikir kritis adalah pemikiran yang masuk akal dan reflektif yang berfokus untuk memutuskan apa yang mesti dipercaya atau dilakukan [2]. Menurut Ennis terdapat 12 indikator keterampilan berpikir kritis, salah satunya kemampuan memutuskan sumber yang dapat dipercaya. Kemampuan ini dapat dilatihkan kepada mahasiswa selama proses pembelajaran dengan menggunakan model *problem solving*. Langkah-langkah model *problem solving* yang digunakan untuk melatih kemampuan memutuskan sumber yang dapat dipercaya adalah pada tahap 2 mencari data untuk memecahkan masalah dan tahap 4 menguji kebenaran jawaban sementara. Pada tahap 2, mahasiswa diberi kesempatan mencari informasi dari banyak sumber, lalu dilatih mempertimbangkan beberapa sumber yang diperoleh kemudian memutuskan sumber mana yang dapat dipercayai. Kemudian pada tahap 4, mahasiswa diminta dapat menganalisis informasi yang diperoleh pada tahap 2 sehingga dapat memutuskan kebenaran hipotesisnya dan mengetahui kualitas keakuratan sumber informasi yang diperoleh sebelumnya.

Beberapa peneliti telah mengkaji tentang penerapan model pembelajaran *problem solving* dalam meningkatkan keterampilan berpikir kritis siswa. Model pembelajaran *problem solving* dapat meningkatkan keterampilan berpikir kritis dan penguasaan konsep sistem koloid [3]. Pembelajaran dengan model *problem solving* lebih efektif untuk meningkatkan keterampilan berpikir kritis siswa dari pada dengan pembelajaran konvensional pada materi kesetimbangan kimia [4].

Materi pH larutan merupakan materi yang memerlukan proses berfikir kritis karena banyak perhitungan matematis yang harus dikuasai mahasiswa untuk menyelesaikan soal pH larutan.

Berdasarkan uraian di atas, maka artikel ini akan mendeskripsikan: Efektivitas Model *Problem Solving* dalam Meningkatkan Kemampuan Berfikir Lancar Mahasiswa pada Materi pH Larutan.

2. METODE PENELITIAN

Populasi dalam penelitian ini adalah mahasiswa semester ganjil, Program Studi Pendidikan Fisika, FKIP Unila tahun akademik 2017/2018. Kemampuan akademik mahasiswa pada tiap kelas adalah heterogen, sehingga

proporsi jumlah mahasiswa yang memiliki kemampuan akademik yang tinggi, sedang maupun kurang dalam tiap kelasnya hampir sama. Sampel dalam penelitian ini adalah mahasiswa kelas A dan B. Pengambilan sampel dilakukan dengan teknik *purposive sampling*, yaitu teknik pengambilan sampel yang didasarkan pada suatu pertimbangan tertentu berdasarkan ciri atau sifat-sifat populasi yang sudah diketahui sebelumnya. Jenis data yang digunakan dalam penelitian ini adalah data yang bersifat kuantitatif yaitu data hasil tes sebelum pembelajaran diterapkan (pretes) dan hasil tes setelah pembelajaran diterapkan (postes) mahasiswa, serta data yang bersifat kualitatif yaitu data aktivitas belajar mahasiswa.

2.1 Variabel Penelitian

Variabel bebas dalam penelitian ini adalah penggunaan model *problem solving* dan tanpa penggunaan model *problem solving*. Variabel terikat dalam penelitian ini adalah kemampuan berfikir lancar. Variabel kontrol dalam penelitian ini adalah materi yang diberikan yaitu pH larutan.

2.2 Data Penelitian

Jenis data yang digunakan dalam penelitian ini adalah data yang bersifat kuantitatif yaitu data hasil tes sebelum pembelajaran diterapkan (pretes) dan hasil tes setelah pembelajaran diterapkan (postes) siswa, serta data yang bersifat kualitatif yaitu data kinerja guru dan aktivitas belajar siswa.

2.3 Rancangan Penelitian

Penelitian ini menggunakan metode quasi eksperimen dengan desain *Non-Equivalent Control Group Design*. Langkah-langkah yang menunjukkan suatu urutan kegiatan penelitian yaitu:

Tabel 1. Rancangan Penelitian

Kelas	Pretes	Perlakuan	Postes
Kelas A	O ₁	X ₁	O ₂
Kelas B	O ₁	-	O ₂

Keterangan :

X₁ : Pembelajaran kimia menggunakan model pembelajaran *problem solving*

X₂ : Pembelajaran kimia tanpa menggunakan model *problem solving*

O₁ : Kelas A dan kelas B diberi pretes

O₂ : Kelas A dan kelas B diberi postes [5]

A. Instrumen Penelitian

Instrumen yang digunakan pada penelitian ini antara lain silabus, rencana pelaksanaan pembelajaran (RPP), kisi-kisi soal, instrumen tes, rubrik penilaian instrumen tes. Adapun instrumen tes yang digunakan berupa soal pretes dan postes. Soal pretes yang digunakan adalah soal uraian yang mengukur kemampuan berfikir

mahasiswa pada materi asam basa dan soal postes yang digunakan adalah soal uraian yang mengukur kemampuan berfikir mahasiswa pada materi pH Larutan. Dalam pelaksanaannya, kelas A dan kelas B diberikan soal pretes dan postes pada waktu dan soal yang sama.

Pada penelitian ini menggunakan validitas isi yang dilakukan dengan *judgment*. Validitas isi dengan cara *judgment* memerlukan ketelitian dan keahlian penilai, maka dalam hal ini validitas isi dilakukan oleh ahli yaitu dosen kimia, FKIP Unila. Pengujian dilakukan dengan menelaah kisi-kisi, terutama kesesuaian antara tujuan penelitian, tujuan pengukuran, indikator dan butir-butir pertanyaan-annya. Bila ternyata unsur-unsur itu terdapat kesesuaian, maka dapat dinilai bahwa instrumen dianggap valid untuk digunakan dalam mengumpulkan data sesuai kepentingan.

Hipotesis pada penelitian ini adalah rata-rata *n-Gain* mahasiswa dalam berfikir lancar pada materi pH larutan yang diterapkan model *problem solving* lebih tinggi daripada rata-rata *n-Gain* mahasiswa dengan tanpa menggunakan model *problem solving*.

B. Teknik Analisis Data dan Pengujian Hipotesis

Teknik analisis data

Data yang diolah dalam penelitian ini adalah data yang diperoleh dari hasil pretes dan postes untuk mengukur kemampuan mahasiswa dalam berfikir lancar yang diberikan kepada kelas A dan kelas B.

Nilai pretes atau postes dirumuskan sebagai berikut :

$$\text{Nilai} = \frac{\text{Jumlah skor yang diperoleh siswa}}{\text{Jumlah skor maksimal}} \times 100$$

Suatu pembelajaran dikatakan efektif apabila adanya perbedaan yang signifikan secara statistik terhadap hasil belajar mahasiswa di kelas A dan kelas B yang ditunjukkan dengan peningkatan nilai pretes-postes mahasiswa kelas A lebih tinggi dibandingkan peningkatan nilai pretes-postes siswa di kelas B [6].

Untuk mengetahui efektivitas kemampuan mahasiswa dalam berfikir lancar pada materi pH larutan antara kelas yang diterapkan model *problem solving* dan tanpa diterapkan model *problem solving*, maka dilakukan analisis nilai gain ternormalisasi. Perhitungan gain ternormalisasi (*n-Gain*) bertujuan untuk mengetahui peningkatan nilai pretes dan postes kedua kelas. Besarnya peningkatan dihitung dengan rumus *n-Gain*, yaitu :

$$n\text{-Gain} = \frac{(\text{nilai postes-nilai pretes})}{(\text{nilai maksimum -nilai pretes})} \quad [7]$$

Uji normalitas data dilakukan untuk mengetahui apakah kedua kelompok sampel berasal dari populasi yang berdistribusi normal atau tidak. Rumusan hipotesis untuk uji normalitas adalah :

H_0 : sampel berasal dari populasi berdistribusi normal

H_1 : sampel berasal dari populasi berdistribusi tidak normal

Uji ini biasanya menggunakan uji *Chi-Kuadrat* :

$$\chi^2_{hitung} = \sum_{i=1}^n \frac{(O_i - E_i)^2}{E_i}$$

[8]

dengan kriteria uji : terima H_0 jika $\chi^2_{hitung} < \chi^2_{tabel}$ dengan taraf signifikan 5%

Keterangan :

χ^2 : nilai *Chi-Kuadrat*

O_i : frekuensi pengamatan

E_i : frekuensi yang diharapkan

n : banyaknya kelas interval

a. Uji homogenitas dua varians

Uji homogenitas dua varians digunakan untuk mengetahui apakah dua kelompok sampel mempunyai varians yang homogen atau tidak.

H_0 : kedua kelas penelitian mempunyai varians yang homogen

H_1 : kedua kelas penelitian mempunyai varians yang tidak homogen

Rumus statistik untuk uji homogenitas (F) :

$$F_{hitung} = \frac{S_1^2}{S_2^2}$$

$$S^2 = \frac{n \sum fix_i^2 - (\sum fix_i)^2}{n(n-1)}$$

Keterangan :

S_1^2 = varians terbesar

S_2^2 = varians terkecil

Kriteria uji : terima H_0 jika $F_{hitung} < F_{tabel}$, dengan taraf nyata 5%

b. Uji Persamaan dua rata-rata

Uji kesamaan dua rata-rata digunakan untuk menentukan apakah pada awalnya kedua kelas penelitian memiliki kemampuan memfokuskan pertanyaan yang berbeda secara signifikan atau tidak. Hipotesis dirumuskan dalam bentuk pasangan hipotesis nol (H_0) dan hipotesis alternatif (H_1).

Rumusan hipotesis:

$H_0 : \mu_{1x} = \mu_{2x}$: rata-rata nilai pretes kemampuan awal mahasiswa dalam memfokuskan pertanyaan pada kelas A sama dengan rata-rata nilai pretes kemampuan awal mahasiswa dalam memfokuskan pertanyaan pada kelas B.

$H_1 : \mu_{1x} \neq \mu_{2x}$: rata-rata nilai pretes kemampuan awal mahasiswa dalam memfokuskan pertanyaan pada kelas A tidak sama dengan rata-rata nilai pretes kemampuan awal mahasiswa dalam memfokuskan pertanyaan pada kelas B.

Keterangan :

μ_1 : rata-rata nilai pretes (x) pada kelas A

μ_2 : rata-rata nilai pretes (x) pada kelas B

x : kemampuan siswa dalam memfokuskan pertanyaan

Jika data yang diperoleh terdistribusi normal dan homogen, maka pengujian menggunakan uji statistik parametrik, yaitu menggunakan uji-t [8]:

$$t_{hitung} = \frac{\bar{X}_1 - \bar{X}_2}{s \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \quad (4)$$

dan

$$s^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$$

Keterangan :

$\bar{X}_1$ = rata-rata nilai pretes kemampuan mahasiswa dalam berfikir lancar pada kelas A

$\bar{X}_2$ = rata-rata nilai pretes kemampuan mahasiswa dalam berfikir lancar pada kelas B

s^2 = varians gabungan

n_1 = jumlah siswa pada kelas A

n_2 = jumlah siswa pada kelas B

s_1^2 = varians kelas A

s_2^2 = varians kelas B

Dengan kriteria pengujian: terima H_0 jika $-t_{tabel} < t_{hitung} < t_{tabel}$ dengan derajat kebebasan $d(k) = n_1 + n_2 - 2$ dan tolak H_0 untuk harga t lainnya. Dengan menentukan taraf signifikan $\alpha = 5\%$ peluang $(1 - \alpha)$.

c. Uji Perbedaan Dua Rata-Rata

Uji perbedaan dua rata-rata digunakan untuk menentukan seberapa efektif perlakuan terhadap sampel dengan melihat *n-Gain* antara pembelajaran pada kelas A dan B. Hipotesis dirumuskan dalam bentuk pasangan hipotesis nol (H_0) dan hipotesis alternatif (H_1).

Rumusan hipotesis yang digunakan adalah sebagai berikut:

$H_0 : \mu_1 \leq \mu_2$: Rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas eksperimen (yang diterapkan model *problem solving*) lebih tinggi daripada rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas kontrol (yang tidak diterapkan model *problem solving*)

$H_1 : \mu_1 > \mu_2$: Rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas eksperimen (yang diterapkan model *problem solving*) lebih rendah daripada rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas kontrol (yang tidak diterapkan model *problem solving*)

Keterangan:

μ_1 : rata-rata *n-Gain* (x) pada kelas eksperimen

μ_2 : rata-rata *n-Gain* (x) pada kelas kontrol

x : kemampuan siswa dalam memfokuskan pertanyaan

Jika data yang diperoleh terdistribusi normal dan homogen, maka pengujian menggunakan uji statistik parametrik, yaitu menggunakan uji-t [8]:

$$t_{hitung} = \frac{\bar{X}_1 - \bar{X}_2}{s \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \quad (4)$$

dan

$$s^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$$

Keterangan :

$\bar{X}_1$ = rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas A

$\bar{X}_2$ = rata-rata *n-Gain* kemampuan siswa dalam memfokuskan pertanyaan pada kelas B

s^2 = varians gabungan

n_1 = jumlah siswa pada kelas A

n_2 = jumlah siswa pada kelas B

s_1^2 = varians kelas A

s_2^2 = varians kelas B

Dengan kriteria pengujian: terima H_0 jika $t > t_{1-1/2\alpha}$ dengan derajat kebebasan $d(k) = n_1 + n_2 - 2$ dan tolak H_0 untuk harga t lainnya. Dengan menentukan taraf signifikan $\alpha = 5\%$ peluang $(1 - \alpha)$.

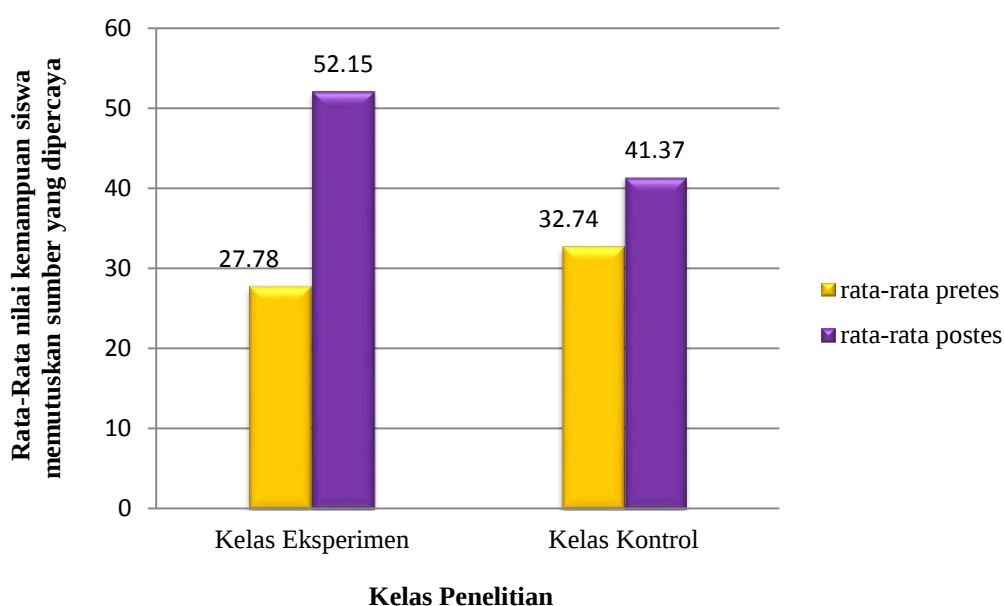
3. HASIL DAN PEMBAHASAN

Rata-rata nilai pretes, nilai postes dan *n-Gain* kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya ditampilkan pada Tabel 1.

Tabel 1. Rata – rata nilai pretes, postes dan *n-Gain* kemampuan mahasiswa di kelas A dan B

Kelas	Rata- Rata		
	Pretes	Postes	<i>n-Gain</i>
A	27,78	52,15	0,35
B	32,74	41,37	0,10

Perbedaan rata-rata nilai pretes dan rata-rata nilai postes kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya antara kelas A dan kelas B ditampilkan pada Gambar 1.



Gambar 1. Rata-rata nilai pretes dan postes kemampuan mahasiswa di kelas A dan kelas B.

Gambar tersebut memperlihatkan bahwa pada kelas A rata-rata nilai kemampuan awal (pretes) mahasiswa dalam memutuskan sumber yang dapat dipercaya sebesar 27,78 lebih kecil dari rata-rata kemampuan awal mahasiswa di kelas B yang sebesar 32,74. Setelah diberi perlakuan pembelajaran dengan model *problem solving* dan diuji rata-rata nilai kemampuan akhir (postes) mahasiswa dalam memutuskan sumber yang dapat dipercaya pada kelas A sebesar 52,15 lebih besar dibandingkan rata-rata nilai kemampuan akhir mahasiswa di kelas kontrol yang pembelajarannya tanpa model *problem solving* yaitu hanya 41,37. Jika dilihat dari Gambar 1 peningkatan pada kelas A lebih tinggi dari pada kelas B. Hal ini menunjukkan bahwa kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya pada akhir pembelajaran di kelas A lebih baik dibandingkan di kelas B.

Berdasarkan data penelitian dan analisisnya menunjukkan bahwa rata-rata nilai *n-Gain* kemampuan siswa dalam memutuskan sumber yang dapat dipercaya pada materi pH larutan yang menggunakan model *problem solving*

lebih tinggi dibandingkan tanpa menggunakan model *problem solving*. Mengapa hal tersebut terjadi untuk mengetahuinya dilakukan pengkajian sesuai dengan fakta yang terjadi pada langkah-langkah model *problem solving* yang digunakan dalam penelitian ini sesuai dengan teori dari [1].

Tahap 1. Mengorientasikan siswa pada masalah. Pada kelas A, mahasiswa dikelompokkan sesuai secara heterogen. Pada tahap ini, dosen mengajukan fenomena sehingga memunculkan masalah dan dapat membangkitkan rasa ingin tahu mahasiswa dalam rangka memotivasi mahasiswa untuk terlibat dalam pemecahan masalah tersebut.

Pada LKS 1, mahasiswa diberikan fakta-fakta tentang sifat-sifat larutan asam, basa dan garam. Pada LKS pertama ini mahasiswa masih sangat mengalami kesulitan dalam menemukan permasalahan dari fakta yang diberikan, untuk itu dosen membantu dengan memberikan pertanyaan-pertanyaan yang memancing mahasiswa untuk menemukan permasalahan yang berhubungan dengan sifat larutan asam, basa dan garam. Pada LKS 2, mahasiswa sudah mulai mampu menemukan sendiri permasalahan dari fakta yang diberikan bahwa ada yang menyebabkan larutan garam dapat bersifat asam, basa atau netral. Permasalahan ini memang harus tumbuh dari mahasiswa sesuai dengan taraf kemampuannya. Mahasiswa akan mengalami kebingungan dan mempunyai rasa keingintahuan yang tinggi terhadap fakta baru yang mengarah pada berkembangnya daya nalar tingkat tinggi mahasiswa yang diawali dengan kata-kata seperti mengapa dan bagaimana. Munculnya pertanyaan-pertanyaan tersebut sekaligus merupakan indikator kesiapan mahasiswa untuk menempuh tahap-tahap berikutnya. Hal ini sesuai dengan yang diungkapkan [9] bahwa perumusan pertanyaan-pertanyaan dan penciptaan masalah-masalah merupakan bagian yang paling penting dan paling kreatif dari sains. Pada LKS 3, mahasiswa mengalami kesulitan kembali untuk menemukan permasalahan dari fakta yang diberikan, hal ini disebabkan siswa kurang memahami isi dari orientasi masalah pada LKS 3. dosen harus membantu dengan menjelaskan maksud dari orientasi masalah tersebut dan mengarahkan mahasiswa untuk menemukan masalah dari fakta yang diberikan.

Tahap 2. Mencari data atau keterangan yang dapat digunakan untuk memecahkan masalah. Mahasiswa mencari data misalnya, dengan jalan membaca buku, mencari di internet, bertanya dan lain-lain. Pada tahap ini setelah mahasiswa merumuskan masalah, dosen mengarahkan mahasiswa agar mendapatkan sumber informasi yang sesuai dan sebanyak-banyaknya untuk mendapatkan penjelasan dari permasalahan yang mereka ajukan. Pada tahap ini mahasiswa dilatihkan untuk dapat memilah-milih sumber-sumber yang akan dijadikan informasi untuk pemecahan masalahnya.

Pada LKS 1, mahasiswa masih terlihat bingung mencari informasi yang sesuai dan hanya terpaku oleh satu sumber, namun setelah dosen menjelaskan media-media yang dapat digunakan untuk mencari informasi tersebut, terlihat beberapa kelompok menggunakan beberapa sumber untuk dijadikan sebagai informasi. Informasi yang mahasiswa dapatkan dari banyak sumber seperti dari buku dan internet dapat meyakinkan informasi yang mereka dapatkan benar atau tidak. Pada LKS 2, mahasiswa sudah terlihat mulai aktif dan dapat bekerjasama mencari

informasi yang diinginkan. Ternyata terdapat kelompok yang mendapatkan informasi yang berbeda dari sumber-sumber yang mereka dapatkan, namun setelah diberi pengarahan oleh dosen mereka menentukan sendiri sumber yang digunakan untuk tahap selanjutnya. Fakta dilapangan tersebut membuktikan bahwa pada tahap kedua ini dapat melatih kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya. Perbedaan informasi yang diperoleh tersebut ternyata membuat mahasiswa lebih merasa ingin tahu untuk memecahkan permasalahannya dan membuktikan informasi mana yang benar. Hal ini sesuai yang diungkapkan [9] bahwa ada kalanya mahasiswa membandingkan hal-hal yang salah dan berpikir dengan caranya sendiri agar mahasiswa menjadi pemikir-pemikir yang diharapkan. Pada LKS 3, mahasiswa telah mampu mendapatkan informasi yang lebih detail untuk memecahkan permasalahannya.

Pada tahap 2 ini, setiap mahasiswa dalam bekerjasama di kelompoknya lebih mudah teramati, mahasiswa bekerjasama mencari data untuk memecahkan masalahnya, mahasiswa yang kerjasamanya kurang, terlihat lebih santai dan sibuk mengobrol.

Tahap 3. Merumuskan hipotesis masalah. Pada tahap ini, dosen meminta mahasiswa untuk memberikan hipotesis awal terhadap jawaban atas permasalahan yang dikemukakan. Mahasiswa kembali berdiskusi dan bekerja sama dalam kelompok serta menggunakan informasi yang telah didapatkan untuk menjawab pertanyaan dan menetapkan hipotesis dari permasalahan tersebut. Mahasiswa merumuskan hipotesis yang artinya merumuskan kemungkinan-kemungkinan jawaban atas masalah tersebut yang masih perlu diuji kebenarannya.

Pada tahap 3 ini, afektif siswa dalam keberanian mengemukakan pendapat/ alasan dapat dikembangkan. Perkembangannya terlihat dari semakin berani dan percaya diri. Hal ini terlihat saat mahasiswa disuruh untuk menyebutkan hipotesisnya di depan kelas dan juga banyak mahasiswa yang mulai berani mengemukakan pendapatnya.

Tahap 4. Menguji kebenaran jawaban sementara. Pada tahap ini, mahasiswa melakukan kegiatan-kegiatan untuk menguji hipotesis dari permasalahannya sesuai dengan langkah pengujian setiap LKS nya yang berisi pertanyaan untuk mengarahkan mahasiswa dalam menemukan pemecahan masalah sehingga akhirnya dapat memutuskan kebenaran hipotesisnya dan mengetahui kualitas keakuratan sumber informasi yang diperoleh sebelumnya.

Pada tahap 4 ini, mahasiswa di setiap kelompok terlihat semakin aktif melakukan diskusi, dan rasa ingin tahu meningkat ketika melakukan identifikasi gambar mikroskopis serta keberanian mengemukakan pendapatnya juga meningkat. Hal ini terlihat saat mahasiswa mengungkapkan hasil identifikasinya. Pada tahap ini diamati bahwa mahasiswa telah berhasil dibimbing untuk menggali pengetahuan mereka secara bebas berdasarkan penyelidikan yang mereka lakukan serta pengolahan informasi yang mereka peroleh. Kegiatan ini juga dapat melatih mahasiswa membangun konsep dari pemikirannya sendiri khususnya pada materi pH larutan. Hal tersebut sesuai yang diungkapkan Hidayati bahwa pemecahan masalah memberikan kesempatan mahasiswa untuk mempelajari, mencari, dan menemukan sendiri informasi untuk diolah menjadi konsep, prinsip, teori atau membuat keputusan tertentu.

Tahap 5. Menarik kesimpulan. Dalam tahap ini mahasiswa diberi kesempatan menyimpulkan hasil temuan bersama kelompoknya untuk menyelesaikan masalah yang ditemukan pada tahap 1. Tahap ini, mahasiswa diberi kebebasan untuk membuat keputusan atas jawaban dari masalahnya. Kemudian setiap kelompok mempresentasikan hasil temuan untuk memecahkan masalahnya.

Berbeda halnya dengan mahasiswa pada kelas B yang menggunakan metode ceramah dan diskusi kelas dalam pembelajarannya. Rasa ingin tahu mahasiswa di kelas B lebih rendah dibandingkan kelas A. Hal ini disebabkan mahasiswa di kelas B terbiasa memperoleh pengetahuan hanya dari penjelasan dosen semata tidak diarahkan untuk terlebih dahulu mencari pengetahuan dari sumber lain sehingga secara pengetahuan dan pengalaman belajar sangat jauh berbeda jika dibandingkan dengan kelas A, seperti pengalaman mendapatkan sumber informasi yang tidak dapat dipercaya. Mahasiswa di kelas B akan sulit membedakan apakah sumber itu dapat dipercaya atau tidak sebab sumber informasi yang mereka dapat selama ini hanya dari dosen, sehingga apabila dosen sewaktu-waktu salah memberikan informasi kepada mahasiswa maka mahasiswa tidak tahu bahwa informasi tersebut salah. Saat diskusi kelas berlangsung kurang terlihat adanya interaksi balik dari mahasiswa ke dosen. Saat dosen bertanya, mahasiswa yang sebenarnya mengetahui jawabannya tidak berani langsung menjawab, perlu waktu lama membujuk mahasiswa tersebut untuk mengutarakan pendapatnya, sementara mahasiswa lain yang tidak mengerti lebih memilih untuk diam dan mengobrol dengan teman sebangku.

Kenyataan di atas jelas akan memberikan pencapaian yang berbeda dengan kelas B yang tidak mengalami tahap demi tahap seperti pada kelas A. Hal ini terbukti dengan lebih baiknya pencapaian nilai postes kelas A dibandingkan dengan nilai postes kelas B, dalam hal kemampuan memutuskan sumber yang dapat dipercaya pada materi pH larutan. Diperkuat dengan hasil uji statistik yaitu uji-t yang menunjukkan bahwa rata-rata nilai *n-Gain* kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya pada kelas A yang menggunakan model *problem solving* lebih tinggi daripada rata-rata nilai *n-Gain* kemampuan mahasiswa dalam memutuskan sumber yang dapat dipercaya pada kelas B yang tidak menggunakan model *problem solving* pada materi pH larutan. Selain itu, data peningkatan perkembangan rata-rata afektif mahasiswa selama proses pembelajaran (rasa ingin tahu, komunikatif, bekerjasama, antusias menjawab pertanyaan dan berani mengemukakan pendapat) di setiap pertemuan pada kelas A lebih tinggi dibandingkan kelas B.

Berdasarkan hal tersebut, pembelajaran model *problem solving* dapat dikatakan efektif dalam meningkatkan kemampuan memutuskan sumber yang dapat dipercaya karena terdapat perbedaan rata-rata nilai *n-Gain* yang signifikan secara statistik antara kelas A dan kelas B. Hal ini sesuai dengan teori efektivitas pembelajaran menurut [6] yaitu suatu pembelajaran dikatakan efektif apabila adanya perbedaan yang signifikan secara statistik terhadap hasil belajar mahasiswa di kelas A dan kelas B yang ditunjukkan dengan peningkatan nilai pretes-postes mahasiswa kelas A lebih tinggi dibandingkan peningkatan nilai pretes-postes siswa di kelas B.

4. KESIMPULAN

Berdasarkan hasil penelitian dan pembahasan dapat disimpulkan bahwa model *problem solving* efektif dalam meningkatkan kemampuan berfikir lancar maha-siswa dalam hal memutuskan sumber yang dapat dipercaya pada materi pH larutan.

DAFTAR PUSTAKA

- [1] Suryani, L. A. 2012. *Strategi Belajar Mengajar*. Ombak.Yogyakarta.
- [2] Ennis, R. H. 1989. *Critical Thinking*. University of Illinois. Urbana-Campaign.
- [3] Kurniawati, E. 2011. Penerapan Metode *Problem Solving* untuk Meningkatkan Keterampilan Berpikir Kritis dan Penguasaan Konsep Sistem Koloid. *Skripsi* (tidak diterbitkan). Universitas Lampung. Bandar Lampung.
- [4] Saputra, A. 2012. Model Pembelajaran *Problem Solving* pada Materi Pokok Keseimbangan Kimia untuk Meningkatkan Keterampilan Berpikir Kritis Siswa. *Skripsi* (tidak diterbitkan). Universitas Lampung. Bandar Lampung.
- [5] Creswell, J. W. 2003. *Research Design Qualitative, Quantitative and Mixed Methods Approaches Second Edition*. Sage Publications. New Delhi.
- [6] Mergendoller, John R dan Nan L Maxwell. 2006. The Effectiveness of Problem-Based Instruction : A Comparative Study Of Instructional Methods and Student Characteristics. *The Interdisciplinary Journal Of Problem Based Learning/ Vol.1, No.2*. 1-69.
- [7] Rismalinda, A. 2014. Pembelajaran Menggunakan Pendekatan Ilmiah dalam Meningkatkan Keterampilan Berpikir Lancar pada Materi Keseimbangan Kimia. *Skripsi* (tidak diterbitkan). Universitas Lampung. Bandar Lampung.
- [8] Sudjana. 2005. *Metode Statistika*. Tarsito. Bandung.
- [9] Dahar, R. W. 1996. *Teori-Teori Belajar*. Erlangga. Jakarta.

THE OPTIMAL BANDWIDTH FOR KERNEL DENSITY ESTIMATION OF SKEWED DISTRIBUTION: A CASE STUDY ON SURVIVAL TIME DATA OF CANCER PATIENTS

Netti Herawati, Khoirin Nisa, Eri Setiawan

Department of Mathematics University of Lampung

Jl. Prof. Dr. Soemantri Brodjonegoro No. 1 Gedung Meneng Bandar Lampung

e-mail: netti.herawati@fmipa.unila.ac.id, khoirin.nisa@fmipa.unila.ac.id, eri.setiawan@fmipa.unila.ac.id .

ABSTRACT

In this paper, optimal bandwidth selection for skewed distribution is studied through data simulation. The data are generated from Exponential (1), Exponential (5), Gamma (1,6), Gamma (1,9), Weibull (1,5), and Weibull (1,10) distributions having parameter(s) that produces a skewed density function with $n = 100$. The Gaussian kernel density functions of the generated distributions are estimated using Scott (Nrd), Silverman's rule of thumb (Nrd0), Silverman's Long-Tailed distribution (Silverman-LT), Biased Cross validation (BCV) and Sheater-Jones (SJ) bandwidth methods. The kernel density estimates are compared to the corresponding probability density functions. The selected optimal bandwidth then is applied to kernel density estimation of survival time data of cancer patients. Result indicates that, overall, Silverman's Rule of Thumb (Nrd0) method outperformed the other methods.

Keywords: kernel density estimation, optimal bandwidth, survival time data

1. INTRODUCTION

Information about data distribution and its probability density function (PDF) is important in various statistical analysis. However, in some cases the information about the probability density function is unknown. One approach to density estimation is parametric. Another approach is non parametric in that less rigid assumptions will be made about the distribution of the observed data. The objective of density estimation in nonparametric approach is to obtain the density function curve which is a smooth curve with minimum sampling variance and has important information of the data. The simplest way to estimate the density estimation is using histogram. But it has a weakness in its shape which is influenced by the selection of the starting point and the width of the class interval. Different starting points will produce different histograms, as well as different interval classes will result in different histogram shapes.

In this study, we use nonparametric kernel density estimation to estimate the probability density function of skewed distributions. The most important part of kernel density estimation is the selection of kernel functions and the selection of bandwidth [1,3]. Bandwidth is a scale factor that controls how large the probability of spreading point on the curve. Selection of bandwidth will determine whether the obtained density function curve will be undersmoothing or oversmoothing. The value of the bandwidth that is too small will produce an undersmoothing density function curve, and vice versa, the value of the bandwidth that is too large will produce an oversmoothing density function curve. Therefore an optimal bandwidth is required to obtain a density function curve corresponding to the actual data distribution. To illustrate this issue, we provide an example of undersmoothing and oversmoothing density functions presented in Figure 1.

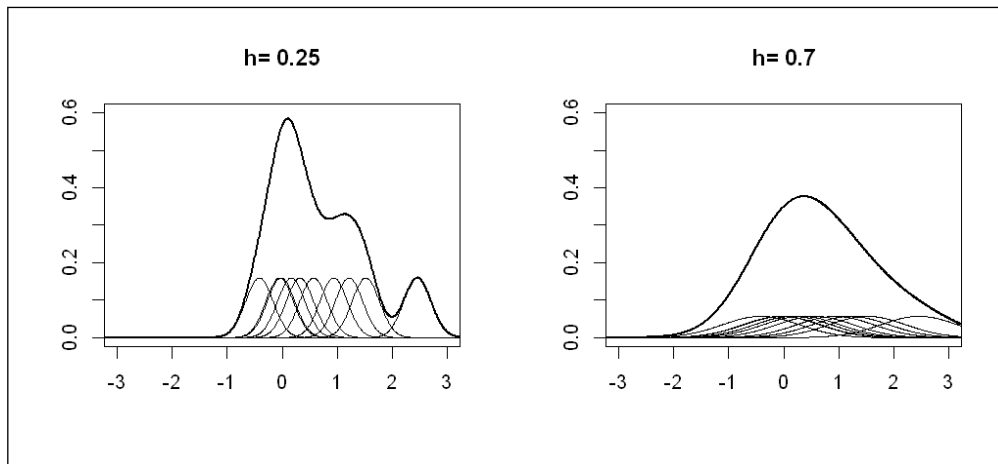


Figure 1. Undersmoothing (left) and Oversmoothing (right).

There are various well known methods available for obtaining optimal bandwidth, for example the Asymptotically first-order optimal (AFO) bandwidths[5], Unbiased Cross Validation (UCV) method, the Biased Cross Validation (BCV) method, Silverman's rule of thumb (Nrd0) method, Silverman's Long-Tailed distribution (Silverman-LT), Scott (Nrd) method, and Sheater-Jones (SJ) method. Each of these methods is known as the optimal bandwidth but produces different bandwidth values. The bandwidth selection in kernel density estimation becomes an interesting topic to study and has been investigated by many authors.

Studies on the optimal bandwidth selection have been done by many authors, one can see e.g. [2]-[6] for various techniques and issues related to bandwidth selections. In this paper we empirically study the optimal bandwidth estimations for skewed distribution by data simulation. We consider several optimal bandwidths mentioned above. The selected one then is applied to real data on survival time of lung cancer patients. The rest of the paper is organized as follow, in Section 2 we review the kernel method for density estimation and present the available optimal bandwidths. In Section 3 we describe our research methodology. The simulation results are presented in Section 4 and the application on survival data is given in Section 5.

2. KERNEL DENSITY ESTIMATION

Kernel Density Estimation is a method to estimate the frequency or probability function of a given value given a random sample. Given a set of observations (x_i) with $1 \leq i \leq n$, it is assumed that the observations are a random sampling of a probability distribution f . The kernel estimator of f is given as follow:

$$\hat{f}(x; h) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - X_i}{h}\right)$$

where:

- X_i is an identical independently distributed random variable
- $K: R^p \rightarrow R$ is the kernel, a function centered on 0 and that integrates to 1.

- n is the number of observations
- h is the bandwidth, a positive valued smoothing parameter that would typically tend to 0 when the number of samples tend to ∞

The kernel estimator depends on two parameters, i.e. the kernel function K and the bandwidth h . There are several kernel functions $K(\cdot)$ which can be used for density estimation, some of the most used kernel functions are Epanechnikov, biweight, triangular, Gaussian and rectangular kernels. For more details on the kernel function one can see e.g. [1] and [9]. However, the selection of kernel function $K(\cdot)$ used for the estimation does not really effect the accuracy of the estimation, furthermore the bias of density estimation using kernel estimator does not rely on the sample size but it depends only on the bandwidth h choice. [1]

Selecting an appropriate bandwidth for a kernel density estimator is of crucial importance, and the purpose of the estimation may be an influential factor in the selection method. In many situations, it is sufficient to subjectively choose the smoothing parameter by looking at the density estimates produced by a range of bandwidths. One can start with a large bandwidth, and decrease the amount of smoothing until reaching a "reasonable" density estimate. However, there are situations where several estimations are needed, and such an approach is impractical. An automatic procedure is essential when a large number of estimations are required as part of a more global analysis.

The problem with using the optimal bandwidth is that it depends on the unknown quantity f'' which measures the speed of fluctuations in the density f , i.e., the roughness of f . Many methods have been proposed to select a bandwidth that leads to good performance in the estimation. The followings methods are some optimal bandwidths selection methods available in R software for kernel density estimation:

A. Scott (Nrd) bandwidth method

A bandwidth that optimize the Integrated Mean Square Error (IMSE) introduced by Scott is given in the following simple formula :

$$h = 1.06 \sigma n^{-1/5},$$

Where σ is population standard deviation (estimated by the sample standard deviation) and n is the sample size. The Scott bandwidth is usually used for normal symmetric and unimodal data [7].

B. Silverman's rule of thumb (Nrd0) bandwidth method

If the data is unimodal but not symmetric, then the following Silverman's rule of thumb bandwidth will optimize the IMSE :

$$h_{opt} = 0.9 \min\{S, IQR/1.34\} n^{-1/5},$$

where S is the sample standard deviation, IQR is the Inter Quartile Range ($Q3 - Q1$) [7].

C. *Silverman's Long-Tailed distribution (Silverman-LT)*

Silverman [1] introduced a bandwidth estimator for skewed and long-tailed distribution given in the following formula:

$$h = 0.79 (\text{IQR}) n^{-1/5}.$$

D. *Unbiased Cross Validation(UCV)bandwith method*

The UCV bandwidth (h_{UCV}) is the bandwidth h that minimize the following function

$$\text{UCV}(h) = \frac{1}{2nh\sqrt{\pi}} + \frac{1}{n^2h\sqrt{\pi}} \sum_{1 \leq i < j \leq n} \left[\exp\left(-\frac{(x_i - x_j)^2}{4h^2}\right) - \sqrt{8} \exp\left(-\frac{(x_i - x_j)^2}{2h^2}\right) \right].$$

The h_{UCV} value is obtained iteratively [8].

E. *Biased Cross validation (BCV)bandwith method*

The BCV bandwidth (h_{BCV}) is the bandwidth h that minimize the following function

$$\text{BCV}(h) = \frac{R(K)}{nh} + \frac{\mu_2(K)^2}{2n^2h} \sum_{1 \leq i < j \leq n} \int K''(w) K''\left(w + \frac{(x_i - x_j)}{h}\right) dw.$$

The h_{BCV} value is also obtained iteratively [8].

F. *Sheater-Jones(SJ)bandwith method*

Sheater-Jones bandwidth method is given as follow

$$h = \left[\frac{R(K)}{\mu_2(K)^2 \hat{S}_D(a_2(h))} \right]^{1/5} n^{-1/5},$$

where $a_2 = \hat{c}_1 h^{5/7}$ with $\hat{c}_1$ is an appropriate constant [4].

3. RESEARCH METHOD

A simulation study using R program was conducted to compare the several optimal bandwidth selection methods: Scott (Nrd), Silverman's rule of thumb (Nrd0), Silverman's Long-Tailed distribution (Silverman-LT), Unbiased Cross Validation(UCV), Biased Cross validation (BCV) and Sheater-Jones (SJ) of $n=100$ which were artificially repeated from skewed distributions: Exponential (1), Exponential (5), Gamma (1,6), Gamma (1,9), Weibull (1,5), and Weibull (1,10). Gaussian kernel with selected optimal bandwidth selection method is applied to estimate the density function of survival time data of lung cancer patients consisting of 62 patients measured in

days, ie patient duration ranging from 100 days before treatment until the patient died. This data is part of the research data of Veteran Administration USA [10].

4. RESULTS AND DISCUSSION

In this section we present and discuss the result from our data simulation. The simulation results of the generated skewed distributions (Exponential (1), Exponential (5), Gamma (1,6), Gamma (1,9), Weibull (1,5), and Weibull (1,10)) using Nrd0, Nrd, UCV, BCV, SJ, and Silverman-LT bandwidth method scan been seen clearly in Figure 2-7.

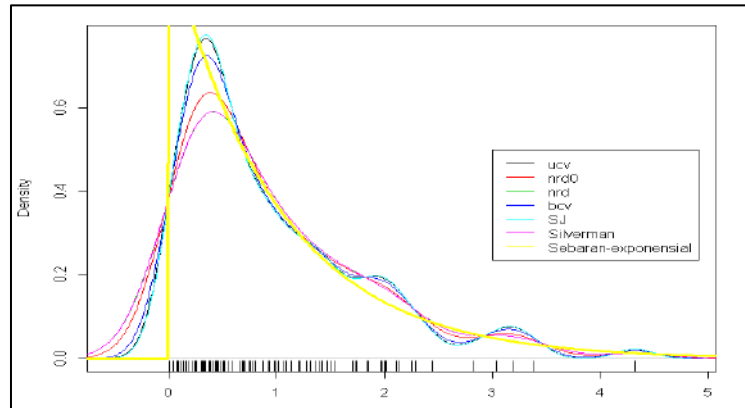


Figure 2. Density estimation curves of Exponential(1) distribution with different bandwidth methods

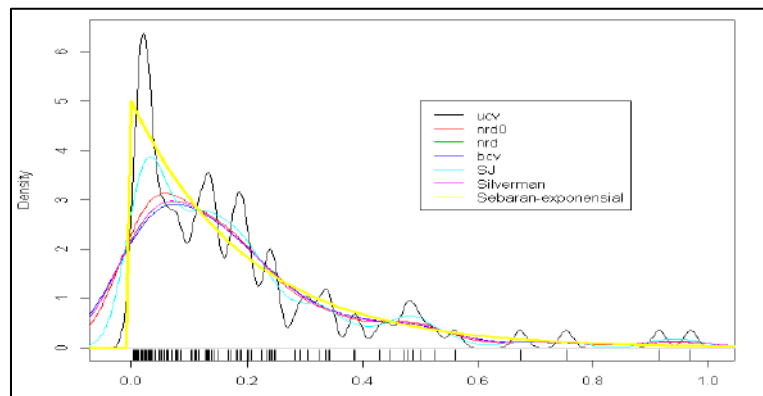


Figure 3. Density estimationcurve of Exponential(5) distribution with different bandwidth methods.

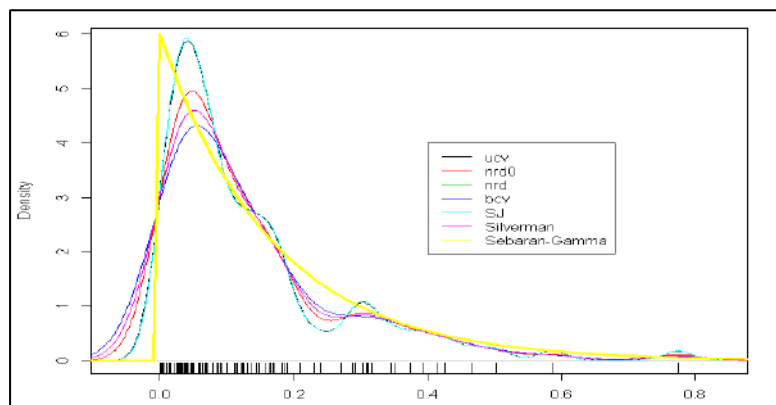


Figure 4. Density estimationcurve of Gamma(1,6) distribution with different bandwidth methods.

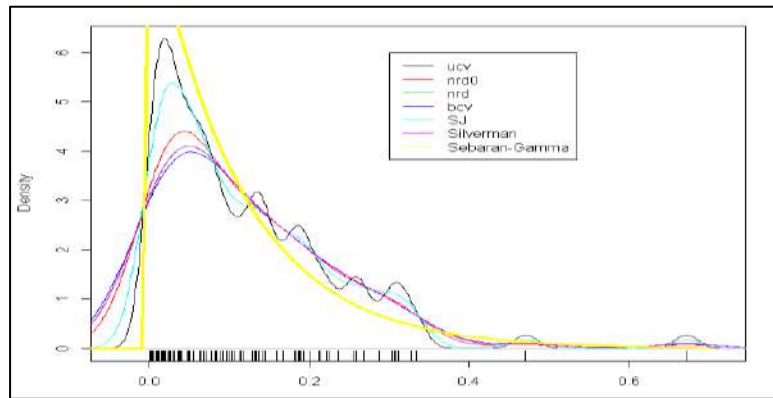


Figure 5. Density estimation curve of Gamma(1,9) distribution with different bandwidth methods.

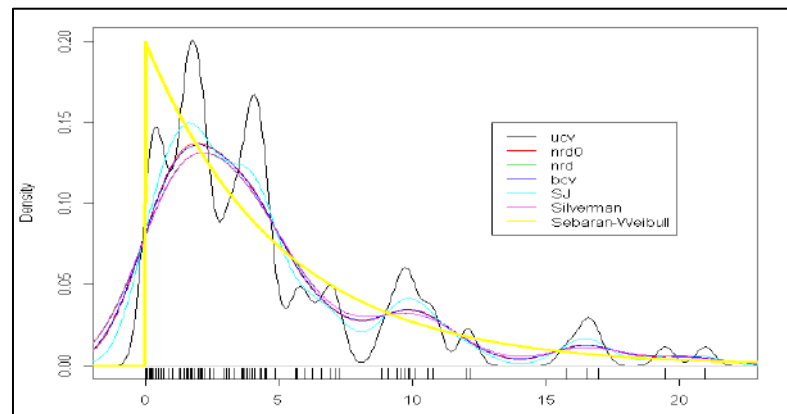


Figure 6. Density estimation curve of Weibull (1,5) distribution with different bandwidth selection methods.

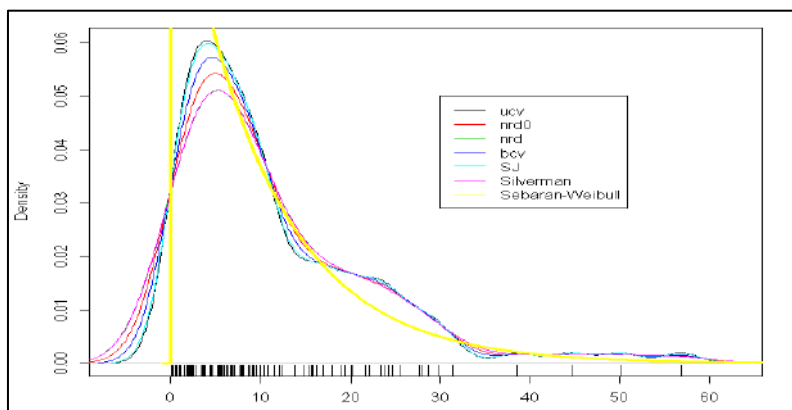


Figure 7. Density estimation curve of Weibull (1,10) distribution with different bandwidth selection methods.

It is apparent from Figure 2-7 that UCV and SJ bandwidths provide unfavorable results for estimating the density curve of the data since the curves are significantly distorted from the real distribution (PDF) curve (ie, yellow curves), particularly in Fig. 3, Fig. 5 and Fig. 6. The kernel estimation using UCV bandwidth is the worst one followed by SJ bandwidth. The bandwidths behaviour of Nrd0, Nrd, Silverman-LT and BCV methods are

seen to meet the real PDF curve. The curve of Nrd0 method has a shape similar to the curve of Silverman-LT, while the curve of the Nrd method has a shape similar to the BCV curve. From these two groups, it is obvious that the Nrd0 and Silverman-LT curves approximate the actual data curve better than the Nrd and BCV curves. In order to ensure the best bandwidth estimation methods, specifically we provide the curve of the density of the two bandwidth selection methods (Nrd0 and Silverman-LT) for Exponential (1), Exponential (5), Gamma (1,6), Gamma (1,9), Weibull (1,5), and Weibull (1,10) distributions as shown Fig. 8.

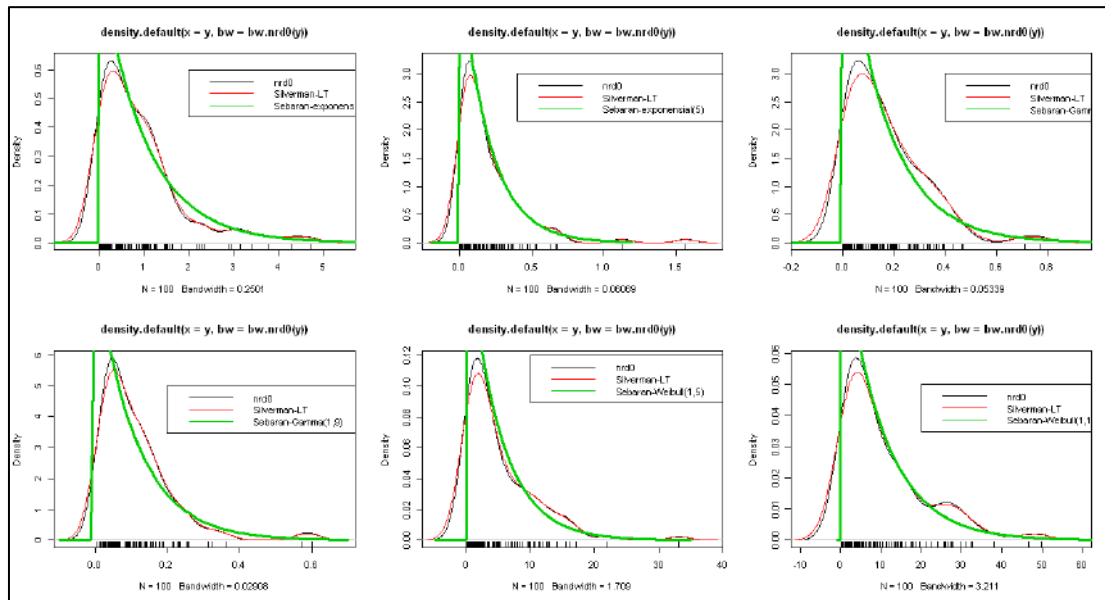


Figure 8. Comparison of Nrd0 and Silverman-LT bandwidth selection methods for Exponential (1), Exponential (5), Gamma (1,6), Gamma (1,9), Weibull (1,5), and Weibull (1,10) distributions.

Figure 8 shows that the density curves of the Nrd0 and Silverman-LT methods are similar. However, at its peak, the density curve of the Nrd0 bandwidth method approximates to the real data distribution (PDF) curve better than the density curve of the Silverman-LT method. Based on the above description, it can be concluded that the Nrd0 bandwidth estimation method gives the best density curve estimation among other methods.

5. APPLICATION ON SURVIVAL DATA OF CANCER PATIENTS

To see the performance of the selected bandwidth method resulting from the simulation study above, we apply it to the survival time data of cancer patients. First, we predict the data distribution using histogram and polygon as shown in Figure 9.

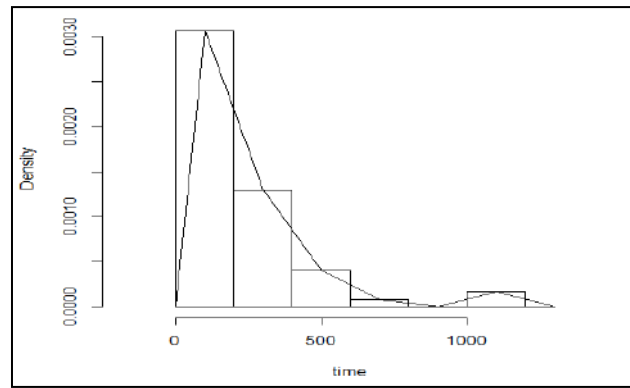


Figure 9. Histogram of survival time data of cancer patients

The histogram above shows that the data is skewed to the right. The values of all optimal bandwidths discussed in previous section are presented in Table 1.

Table 1. Optimal bandwith values of all bandwidth selection methods

Methods	Bandwidth
Silverman's Rule of Thumb (Nrd0)	34.35
Scott (Nrd)	40.46
Unbiased Cross Validation (UCV)	14.37
Biased Cross Validation (BCV)	44.85
Sheater – Jones (SJ)	20.39
Silverman Long-Tailed distributions (Silverman-LT)	36.72

We obtained the bandwidth values of the two best methods (Nrd0 and Silverman-LT) are 34.35 and 36.72 respectively. We present the plot of the polygon curve and kernel density estimates using Nrd0 and Silverman-LT bandwidths in Figure 10.

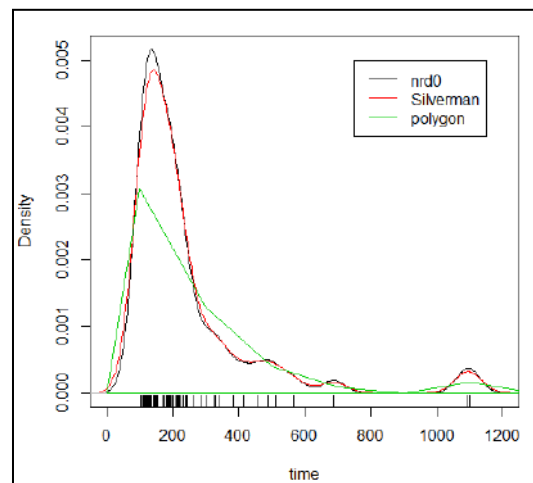


Figure 10. Comparison kernel density estimates using Nrd0 and Silverman-LT bandwidths and polygon of survival time data of cancer patients.

From Figure 10, it can be seen that Nrd0 and Silverman-LT bandwidth methods yields the density curves that best suits the distribution of survival time data of cancer patients. The Nrd0 and Silverman-LT density curves look similar. However, at the peak and tail of the densities are quite different. For final decision, we choose the best bandwidth to fit the data in accordance to the result of data simulation, i.e. the Nrd0 bandwidth. This results is similar to the skewed distributions simulation result as shown in Figure 8.

6. CONCLUSION

In this paper, we have shown that one can estimate the density curve using kernel density estimation method using various bandwidth selection methods. For skewed distribution considered in this paper, overall, Silverman's Rule of Thumb (Nrd0) bandwidth outperformed all other methods. It provides the best density estimates curve for both skewed distribution and survival time data of cancer patients compared to others.

REFERENCES

- [1] Silverman, B.W. 1986. *Density Estimation for Statistics and Data Analysis*. Chapman and Hall, London.
- [2] Samworth, R.J. & Wand M.P. 2010. Asymptotics and Optimal Bandwidth Selection for Highest Density Region Estimation. *The Annals of Statistics*, **38**(3), 1767–1792.
- [3] Chen., S. 2015. Optimal Bandwidth Selection for Kernel Density Functional Estimation. *Journal of Probability and Statistics*, **2015**:1-21.
- [4] Sheather, S.J. & Jones, M.C. 1991. A Reliable Data-Based Bandwidth Selection Method for Kernel Density Estimation. *Journal of the Royal Statistical Society*, **53** (3), 683-690.
- [5] Arai, Y. & Ichimura, H. 2013. Optimal Bandwidth Selection for Differences of Nonparametric Estimators with an Application to Sharp Regression Discontinuity Design. GRIPS discussion paper. National Graduate Institute for Policy Studies, Tokyo Japan.
- [6] Hansen, B.E. 2004. *Bandwidth Selection for Nonparametric Distribution Estimation*. Research supported by the National Science Foundation. Department of Economics University of Wisconsin, Madison.
- [7] Rizzo, M.L. 2008. *Statistical Computing with R*. Chapman & Hall/CRC. Boca Raton.
- [8] Vrahimis, A. 2010. Smoothing Methodology With Applications To Nonparametric Statistics. thesis submitted to the University of Manchester, School of Mathematics.
- [9] Sheather, S.J. 2004. Density Estimation. *Statistical Science*, **19**(4), 588-597.
- [10] Prentice, R.L. 1973. Exponential Survival with Censoring and Explanatory Variables. *Biometrika*, **60**, 279-288.

KARAKTERISTIK LARUTAN KIMIA DI DALAM AIR DENGAN MENGGUNAKAN SISTEM PERSAMAAN LINEAR

Titik Suparwati

Dosen Jurusan Matematika FMIPA Universitas Cenderawasih, Jayapura

e-mail : tix_az@ymail.com

At this writing is given by picture about system application equatioan of linear differensial. There are two system namely open system and closed system. But in this artichel will be discuss open system only. Open system the studied is system which there are stream enter from outside into tank A and or tank B as well as stream go out system of tank B. Amount of chemicals in both tank after time of t, can be counted by making pemodelan of mathematics. From research can be seen that system equatian of formed by differensial is open system is system equation of linear differensial of nonhomogen . To the number of chemicals at tank A and B at open system to be determined by including the problem of value early into common solution of system equation of formed of differensial the system.

Keyword : condensation chemicals, System equation of linear differensial, and open system.

1. PENDAHULUAN

Setiap cabang ilmu Matematika mempunyai aplikasi atau penerapan di berbagai bidang ilmu pengetahuan, di antaranya matematika terapan yang menjadikan kalkulus modern sebagai persamaan dan sistem persamaan differensial yang merupakan bentuk atau model matematika yang cukup mendominasi dalam matematika terapan.

Persamaan differensial merupakan persamaan yang memuat turunan satu buah fungsi yang tidak diketahui. Meskipun persamaan seperti itu seharusnya disebut persamaan turunan, namun istilah persamaan differensial yang diperkenalkan oleh Leibniz pada tahun 1676 sudah umum digunakan. Persamaan differensial hanya memuat satu fungsi yang tak tak diketahui, sehingga penyelesaiannya sampai saat ini berupa satu persamaan differensial yang menagndung satu fungsi yang tak diketahui. Dalam perkembangannya yang berhubungan dengan penerapan persamaan differensial, akan dipelajari pula n buah persamaan differensial dengan n buah fungsi yang tak diketahui, dimana n merupakan bilangan positif lebih dari sama dengan 2 ($n \geq 2$). Inilah yang disebut dengan sistem persamaan differensial [1].

Terdapat banyak aplikasi dari sistem persamaan differensial, diantaranya adalah di bidang ilmu kimia. Reaksi kimia biasanya dapat berlangsung pada dua campuran zat. Salah satu contoh yang lazim dari campuran adalah larutan . Di alam, kebanyakan reaksi berlangsung dalam air. Cairan tubuh, baik tumbuhan maupun hewan adalah larutan yang terdiri dari berbagai zat. Jenis reaksi di samudera, danau dan sunagi juga melibatkan larutan. Kuantitas relatif suatu zat tertentu dalam suatu larutan disebut konsentrasi, sehingga konsentrasi merupakan faktor penting dalam menentukan cepat lambatnya suatu reaksi kimia berlangsung [2].

Dalam penulisan ini akan dikaji aplikasi persamaan differensial linear dalam menentukan jumlah bahan kimia di dalam dua buah tangki setelah dilarutkan ke dalam zat pelarut air, yang mana larutan di dalam tangki A dapat mengalir ke dalam tangki B, dan sebaliknya, dengan konsentrasi yang telah ditentukan.

2. PERUMUSAN MASALAH

Rumusan masalahnya adalah bagaimana mengaplikasikan sistem persamaan differensial dalam menentukan jumlah bahan kimia ke dalam air yang terdapat di dalam dua buah tangki, bagaimana cara menentukan model matematikanya dan bagaimana rumus menghitung banyaknya bahan kimia dalam tangki A dan B pada sistem terbuka setiap saat.

3. TINJAUAN PUSTAKA

Dalam penulisan ini, sistem persamaan differensial yang terbentuk akan diselesaikan dengan menggunakan metode matriks.

3.1. Matriks dan Operasinya

Definisi 3.1 [3]

Bentuk yang paling umum dari sebuah matriks adalah bilangan-bilangan yang berbentuk persegi panjang yang dapat disajikan sebagai berikut:

$$A = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & A_{22} & A_{23} & A_{2n} \\ \dots & \dots & \dots & \dots \\ A_{m1} & A_{m2} & \dots & A_{mn} \end{bmatrix} \quad (3.1.1)$$

Bilangan $A_{11}, A_{12}, \dots, A_{mn}$ yang menyusun rangkaian tersebut disebut elemen dari matriks. Sedangkan indeks pertama dari elemen menunjukkan baris dan indeks kedua menunjukkan kolom.

Ordo sebuah matriks ditentukan oleh banyaknya baris dan kolom dari suatu matriks. Matriks A pada Persamaan (2.1.1) mempunyai ordo $m \times n$.

Matriks bujur sangkar adalah matriks yang mempunyai jumlah baris dan jumlah kolomnya sama ($m = n$) dan dikatakan matriks bujur sangkar berordo n .

Jika dua buah matriks A dan B mempunyai ordo yang sama, maka jumlah $A+B$ adalah matriks yang diperoleh dengan menambahkan bersama-sama elemen yang bersesuaian dalam kedua matriks tersebut. Pengurangan matriks mempunyai syarat yang sama dengan penjumlahan.

Misalkan A adalah matriks $m \times n$ dan B adalah matriks $n \times r$, maka hasil kali AB adalah matriks $m \times r$ yang elemen-elemennya ditentukan sebagai berikut: untuk mencari elemen dalam baris ke i dan kolom ke j dari AB , pilih baris i dari matriks A dan kolom j dari matriks B. Kalikan elemen-elemen yang bersesuaian dari baris dan kolom tersebut bersama-sama, kemudian tambahkan hasil kali yang dihasilkannya.

3.2. Persamaan differensial

Definisi 3.2 [4]

Persamaan differensial adalah suatu persamaan yang meliputi turunan fungsi dari satu atau lebih variabel terikat terhadap satu atau lebih variabel bebas. Selanjutnya jika turunan fungsi itu hanya tergantung pada satu variabel bebas disebut persamaan differensial biasa, dan jika tergantung pada lebih dari satu variabel bebas disebut persamaan differensial parsial.

Persamaan diferensial biasa orde n dikatakan linear bila dapat dinyatakan dalam bentuk

$$a_0(x)y^{(n)} + a_1(x)y^{(n-1)} + \dots + a_n(x)y = F(x), \quad \text{dimana } a_0(x) \neq 0$$

Persamaan diferensial orde satu adalah persamaan yang berbentuk :

$$\frac{dy}{dx} = f(x, y) \quad (3.2.1)$$

(Persamaan di atas dikutip dari [1]).

Definisi 3.3 [4]

Sistem persamaan diferensial linear orde satu disajikan sebagai berikut:

$$\begin{aligned} \frac{dx_1}{dt} &= a_{11}(t)x_1(t) + a_{12}(t)x_2(t) + \dots + a_{1n}(t)x_n(t) + b_1(t) \\ \frac{dx_2}{dt} &= a_{21}(t)x_1(t) + a_{22}(t)x_2(t) + \dots + a_{2n}(t)x_n(t) + b_2(t) \\ &\vdots \\ \frac{dx_n}{dt} &= a_{n1}(t)x_1(t) + a_{n2}(t)x_2(t) + \dots + a_{nn}(t)x_n(t) + b_n(t) \end{aligned} \quad (3.2.3)$$

dengan $a_{ij}(t)$ dan $b_i(t)$ adalah fungsi khusus pada interval I .

Jika $b_1 = b_2 = \dots = b_n = 0$, untuk semua t pada interval I , maka sistem ini dinamakan homogen.

Jika satu atau lebih dari $b_i(t)$ tidak nol, maka sistem Persamaan di atas dinamakan nonhomogen.

3.3. Nilai eigen dan vektor eigen

Definisi 3.4 [5]

Misalkan A adalah matriks $n \times n$, vektor tak nol x di dalam R^n dinamakan vektor Eigen (eigenvector) dari A jika Ax adalah kelipatan skalar dari x ; yakni

$$Ax = \lambda x$$

untuk suatu skalar λ . Skalar λ dinamakan nilai eigen (eigen value) dari A dan x dikatakan vektor eigen yang bersesuaian dengan λ .

Untuk mencari nilai eigen matriks A yang berukuran $n \times n$ maka $Ax = \lambda x$ dituliskan kembali sebagai

$$Ax = \lambda Ix$$

atau

$$(\lambda I - A)x = 0$$

Supaya λ menjadi nilai Eigen, maka harus ada penyelesaian tak nol dari persamaan ini, yaitu jika dan hanya jika

$$\det(\lambda I - A) = 0 \quad (3.3.1)$$

Inilah yang dinamakan persamaan karakteristik A dan skalar yang memenuhi persamaan ini adalah nilai Eigen dari A . Jika λ adalah suatu parameter, maka $\det(\lambda I - A)$ adalah suatu polinom λ yang dinamakan polinom karakteristik dari A .

Vektor Eigen A yang bersesuaian dengan nilai Eigen λ adalah vektor tak nol x yang memenuhi $Ax = \lambda x$

3.4. Sifat Dasar Larutan

Suatu larutan adalah campuran homogen dari molekul, atom, ataupun ion dari dua zat atau lebih. Suatu larutan disebut suatu campuran karena susunannya dapat berubah-ubah. Disebut homogen karena susunannya begitu seragam sehingga tak dapat diamati adanya bagian-bagian yang berlainan, bahkan dengan mikroskop optis sekalipun. Dalam campuran heterogen, permukaan-permukaan tertentu dapat dideteksi antara bagian-bagian atau fase-fase yang terpisah.

Biasanya dengan larutan yang dimaksudkan adalah fase cair. Salah satu komponen penyusun larutan semacam itu adalah suatu cairan sebelum campuran itu dibuat. Cairan ini disebut medium pelarut atau zat pelarut (*solvent*). Zat yang terlarut disebut zat terlarut (*solute*). Air disebut sebagai pelarut karena air tetap mempertahankan keadaan fisiknya. Sedangkan zat terlarut adalah zat yang berubah keadaan fisiknya setelah dicampurkan dengan zat pelarut.

Komposisi zat terlarut dan pelarut dalam larutan dinyatakan dalam *konsentrasi* larutan. Konsentrasi umumnya dinyatakan dalam perbandingan jumlah zat terlarut dengan jumlah total zat dalam larutan, atau dalam perbandingan jumlah zat terlarut dengan jumlah pelarut. Sedangkan proses pencampuran zat terlarut dan pelarut membentuk larutan disebut pelarutan atau *solvasi*. Contoh larutan yang umum dijumpai adalah padatan yang dilarutkan dalam cairan. Seperti garam atau gula dilarutkan dalam air [2].

4. PEMBAHASAN

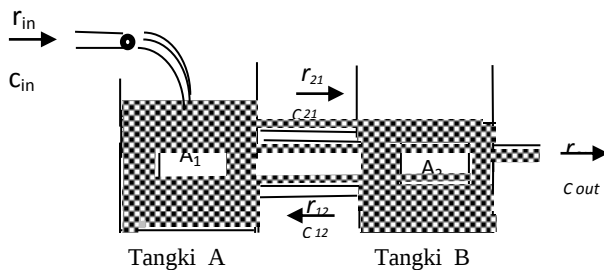
Aplikasi sistem persamaan diferensial linear dalam menentukan jumlah bahan kimia yang dilarutkan ke dalam air yang terdapat di dalam dua buah tangki, yang mana dua buah tangki tersebut saling berhubungan dengan

posisi horisontal sehingga membentuk dua sistem yaitu sistem terbuka dan sistem tertutup, tetapi pada makalah ini hanya akan dijelaskan untuk sistem terbuka sebagai berikut :

Sistem Terbuka

Ada dua bentuk sistem terbuka yang akan dibahas. Salah satu bentuk sistem terbuka dapat dilihat pada Gambar 4.1.1.

Dua tangki berisi larutan yang mengandung garam. Larutan (lain) yang mengandung konsentrasi c_{in} gram/liter garam mengalir ke dalam tangki A dengan laju r_{in} liter/menit dan larutan dengan konsentrasi c_{out} gram/liter mengalir keluar sistem melalui tangki B dengan laju r_{out} liter/menit. Kemudian larutan dengan konsentrasi c_{12} gram/liter mengalir ke dalam tangki A dari tangki B dengan laju r_{12} liter/menit dan larutan dengan konsentrasi c_{21} gram/liter mengalir ke dalam tangki B dari tangki A dengan laju r_{21} liter/menit. Yang akan dihitung adalah $A_1(t)$ dan $A_2(t)$, yakni jumlah garam di dalam tangki A dan tangki B setelah waktu t secara berturut-turut.



Gambar 4.1.1

Diketahui pada Gambar 4.1.1

r_{in} = laju larutan masuk dari luar sistem ke dalam sistem

c_{in} = konsentrasi larutan masuk dari luar sistem ke dalam sistem

r_{21} = laju larutan yang masuk ke dalam tangki B dari tangki A

c_{21} = konsentrasi larutan yang masuk ke dalam tangki B dari tangki A

r_{12} = laju larutan yang masuk ke dalam tangki A dari tangki B

c_{12} = konsentrasi larutan yang masuk ke dalam tangki A dari tangki B

r_{out} = laju larutan keluar meninggalkan sistem

c_{out} = konsentrasi larutan keluar meninggalkan sistem

$A_1(t)$ = jumlah garam di dalam tangki A setelah waktu t

$A_2(t)$ = jumlah garam di dalam tangki B setelah waktu t .

Diasumsikan bahwa larutan pada setiap tangki bercampur sempurna sehingga berdasarkan definisi konsentrasi

diperoleh
$$c_{12} = c_{out} = \frac{A_2}{V_2}, \text{ dan } c_{21} = \frac{A_1}{V_1},$$

dengan V_i merupakan volume larutan di dalam tangki i pada saat t .

Pada interval waktu Δt jumlah bahan kimia yang memasuki tangki A adalah

$$(c_{in} r_{in} + c_{12} r_{12}) \Delta t \text{ gram}$$

dan jumlah bahan kimia yang keluar dari tangki A pada interval waktu yang sama adalah

$$c_{21} r_{21} \Delta t \text{ gram.}$$

Perubahan jumlah bahan kimia di dalam tangki A pada interval waktu Δt , ditunjukkan oleh ΔA_1 , yaitu

$$\Delta A_1 \approx \left(c_{in} r_{in} + r_{12} \frac{A_2}{V_2} - r_{21} \frac{A_1}{V_1} \right) \Delta t. \quad (4.1.1)$$

Jumlah total bahan kimia yang memasuki tangki B pada interval waktu Δt adalah

$$c_{21} r_{21} \Delta t \text{ gram}$$

Jumlah total bahan kimia yang keluar dari tangki B pada interval waktu Δt adalah

$$(r_{12} c_{12} + r_{out} c_{out}) \Delta t \text{ gram}$$

Perubahan jumlah bahan kimia di dalam tangki B pada interval waktu Δt , ditunjukkan oleh ΔA_2 yaitu

$$\Delta A_2 \approx [r_{21} c_{21} - (r_{12} c_{12} + r_{out} c_{out})] \Delta t,$$

$$\Delta A_2 \approx \left[r_{21} \frac{A_1}{V_1} - (r_{12} + r_{out}) \frac{A_2}{V_2} \right] \Delta t.$$

sehingga diperoleh :

$$\begin{aligned} \frac{dA_1}{dt} &= -r_{21} \frac{A_1}{V_1} + r_{12} \frac{A_2}{V_2} + c_{in} r_{in}, \\ \frac{dA_2}{dt} &= r_{21} \frac{A_1}{V_1} - (r_{12} + r_{out}) \frac{A_2}{V_2}. \end{aligned} \quad (4.1.2)$$

Selanjutnya Persamaan (4.1.2) dapat dinyatakan dalam bentuk matriks, yaitu:

$$\dot{x} = A x + b$$

sehingga diperoleh suatu penyelesaian sebagai berikut :

$$x_1 = e^{\lambda_1 t} v_1, \quad x_2 = e^{\lambda_2 t} v_2$$

dengan v_1 dan v_2 , sehingga diperoleh

$$x = c_1 e^{\lambda_1 t} v_1 + c_2 e^{\lambda_2 t} v_2 + x_p$$

Sedangkan untuk menghitung c_1 dan c_2 digunakan syarat awal.

Contoh kasus:

Diketahui : dua buah tangki

$$V_1 = 20 \text{ liter}, V_2 = 20 \text{ liter}, A_1(0) = 40 \text{ gram}, A_2(0) = 20 \text{ gram}$$

$$c_{in} = 4 \text{ g/liter}, r_{in} = 3 \text{ liter/menit},$$

$$r_{12} = 1 \text{ liter/menit}, r_{21} = 4 \text{ liter/menit}, r_{out} = 3 \text{ liter/menit}.$$

$$c_{12} = c_{out} = \frac{A_2}{V_2} = \frac{A_2}{20}$$

$$c_{21} = \frac{A_1}{V_1} = \frac{A_1}{20}$$

Pada interval waktu Δt , perubahan jumlah garam di dalam tangki A adalah

$$\begin{aligned} \Delta A_1 &\approx \left(12 + \frac{A_2}{20} \right) \Delta t - \frac{1}{5} A_1 \Delta t \\ &\approx \left(12 + \frac{1}{20} A_2 - \frac{1}{5} A_1 \right) \Delta t \end{aligned} \quad (4.2.1)$$

Dengan analisis yang sama, perubahan jumlah garam di dalam tangki B pada interval waktu Δt adalah

$$\Delta A_2 \approx \left[\frac{1}{5} A_1 - \frac{1}{5} A_2 \right] \Delta t \quad (4.2.2)$$

Sehingga diperoleh

$$\begin{aligned} \frac{dA_1}{dt} &= -\frac{1}{5} A_1 + \frac{1}{20} A_2 + 12 \\ \frac{dA_2}{dt} &= \frac{1}{5} A_1 - \frac{1}{5} A_2 \end{aligned} \quad (4.2.3)$$

Sistem di atas dapat ditulis sebagai persamaan matriks

$$x' = A x + b$$

$$\text{dengan } x = \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}, \quad A = \begin{bmatrix} -\frac{1}{5} & \frac{1}{20} \\ \frac{1}{5} & -\frac{1}{5} \end{bmatrix}, \quad b = \begin{bmatrix} 12 \\ 0 \end{bmatrix} \text{ Nilai eigen dari } A \text{ adalah}$$

$$\lambda_1 = -\frac{1}{10}, \quad \lambda_2 = -\frac{3}{10}.$$

Sehingga diperoleh

$$x_1 = e^{-t/10} \begin{bmatrix} 1 \\ 2 \end{bmatrix}, \text{ dan } x_2 = e^{-3t/10} \begin{bmatrix} 1 \\ -2 \end{bmatrix}$$

Jadi, solusi umum $x' = Ax$ adalah

$$x_c = c_1 e^{-t/10} \begin{bmatrix} 1 \\ 2 \end{bmatrix} + c_2 e^{-3t/10} \begin{bmatrix} 1 \\ -2 \end{bmatrix}$$

Setelah dicari nilai dari x_p solusi umum untuk $x' = Ax + b$ adalah

$$x = c_1 e^{-t/10} \begin{bmatrix} 1 \\ 2 \end{bmatrix} + c_2 e^{-3t/10} \begin{bmatrix} 1 \\ -2 \end{bmatrix} + \begin{bmatrix} 80 \\ 80 \end{bmatrix}$$

Dengan menggunakan kondisi awal yaitu $A_1(0) = 40, A_2(0) = 20$.

diperoleh $x(0) = \begin{bmatrix} 40 \\ 20 \end{bmatrix}$

$$\begin{aligned} 2c_1 &= -70 \\ \text{dan } c_1 &= -35 \\ c_2 &= -5 \end{aligned}$$

Sehingga diperoleh solusi dari masalah nilai awalnya yaitu

$$x = -35e^{-t/10} \begin{bmatrix} 1 \\ 2 \end{bmatrix} - 5e^{-3t/10} \begin{bmatrix} 1 \\ -2 \end{bmatrix} + \begin{bmatrix} 80 \\ 80 \end{bmatrix}$$

Jadi jumlah garam di dalam tangki A dan B setiap saat yaitu

$$A_1(t) = (80 - 35e^{-t/10} - 5e^{-3t/10}) \text{ gram}$$

$$A_2(t) = (80 - 70e^{-t/10} + 10e^{-3t/10}) \text{ gram}$$

5. KESIMPULAN

Pada penulisan ini dapat disimpulkan bahwa

1. Sistem persamaan diferensial dapat digunakan untuk menghitung jumlah bahan kimia yang dilarutkan ke dalam air yang terdapat di dalam dua buah tangki. Sistem persamaan diferensial yang terbentuk dari sistem terbuka adalah sistem persamaan diferensial linier nonhomogen. Sistem ini diselesaikan dengan menggunakan metode matriks.
2. Banyaknya bahan kimia dalam tangki A dan B pada sistem terbuka setiap saat dapat ditentukan dengan menentukan nilai Eigen dan vektor Eigen, kemudian menentukan solusi dan memasukkan masalah nilai awal ke dalam solusi tersebut.

6. SARAN

Pada penulisan ini hanya membahas tentang sistem terbuka. Secara sama dapat diselesaikan juga untuk sistem tertutup.

KEPUSTAKAAN

- [1] Finizio, N & G. Ladas. 1998, *Persamaan Diferensial Biasa Dengan Penerapan Modern*. Terjemahan oleh Dra. Widiarti Santoso. Jakarta: Erlangga.
- [2] Kleinfelter, Wood. 1991, *Kimia Untuk Universitas*. Terjemahan oleh Aloysius Hadyana Pudjaatmaka, Ph.D. Jakarta: Erlangga.
- [3] Gere, James, M & William Weaver Jr. 1987, *Aljabar Matriks Untuk Para Insinyur*. Terjemahan oleh Drs. G. Tejosutikno, Jakarta: Erlangga.
- [4] Goode, Stephen, W. 1991, *An Introduction To Differential Equations And Linear Algebra*. California: Prentice Hall International Inc.
- [5] Anton, Howard & Rorres Chris, 2004, *Aljabar Linear Elementer Versi Aplikasi*. Jakarta: Erlangga.

BENTUK SOLUSI GELOMBANG BERJALAN PERSAMAAN $\Delta\Delta$ mKdV YANG DIPERUMUM

Notiragayu, R. Ruswandi, dan L. Zakaria

Jurusan Matematika FMIPA, Universitas Lampung
Lampung-Indonesia
notiragaru@gmail.com

Abstract

Bentuk solusi gelombang berjalan dari sebuah sistem dinamik diskrit merupakan bentuk persamaan diskrit biasa (ODE) yang diturunkan dari bentuk parsialnya melalui sebuah transformasi. Persamaan $\Delta\Delta$ -mKdV merupakan sebuah persamaan diskrit parsial (PDE) yang diturunkan dari persamaan mKdV versi kontinu. Dalam artikel ini akan dideskripsikan penurunan bentuk solusi gelombang berjalan dari bentuk persamaan $\Delta\Delta$ -mKdV yang diperumum.

Subject Classification: 37J10, 37J35, 39A11, 70K43

Keywords: Persamaan $\Delta\Delta$ -mKdV yang diperumum, Matriks Lax, Solusi Gelombang Berjalan.

1 Pendahuluan

Sebuah upaya untuk dapat mengkaji lebih banyak dinamika yang terjadi dari sebuah sistem dinamik diskrit dapat dilakukan dengan memperumum bentuk standard sistem dengan cara memperbanyak parameternya. Dalam kertas kerja ini, selain memperlihatkan proses memperumum sistem dinamik $\Delta\Delta$ -mKdV melalui modifikasi parameter pada pasangan matrik Lax juga diperlihatkan proses penurunan persamaan $\Delta\Delta$ -mKdV yang diperumum untuk sebuah solusi gelombang berjalan serta bentuk-bentuk invarian (integral) yang dinormalkan. Dalam artikel Quispel dan kawan-kawan ([1]), sebuah persamaan $\Delta\Delta$ -mKdV pada latis 2D ($\mathbb{Z}^2$) didefinisikan sebagai

$$q(V_{l,m+1}V_{l+1,m+1} - V_{l,m}V_{l+1,m}) = p(V_{l+1,m}V_{l+1,m+1} - V_{l,m}V_{l,m+1}), \quad (1)$$

dimana medan-medan V didefinisikan pada sisi-sisi latis $l, m \in \mathbb{Z}$ yang merupakan dua peubah diskrit. Misalkan $\xi_{l,m}(k)$ menyatakan vektor yang mengandung fungsi gelombang yang bergantung kepada sebuah parameter spektral k . Persamaan di atas dapat diturunkan melalui pemetaan-pemetaan berikut ini

$$\begin{aligned}\xi_{l+1,m}(k) &= \frac{1}{p-k} M_{l,m}^{\text{hor}} \xi_{l,m}(k) \\ \xi_{l,m+1}(k) &= \frac{1}{q-k} M_{l,m}^{\text{vert}} \xi_{l,m}(k)\end{aligned}$$

dengan

$$M_{l,m}^{\text{hor}} = \begin{pmatrix} p & -V_{l+1,m} \\ -\left(\frac{k^2}{V_{l,m}}\right) & p\left(\frac{V_{l+1,m}}{V_{l,m}}\right) \end{pmatrix} \text{ dan } M_{l,m}^{\text{vert}} = \begin{pmatrix} q & -V_{l,m+1} \\ -\frac{q}{V_{l,m}} & q\frac{V_{l,m+1}}{V_{l,m}} \end{pmatrix}.$$

merupakan matriks pasangan Lax. Pemetaan ini terdefinisi dengan baik apabila dipenuhi kondisi berikut.

$$(M_{l+1,m}^{\text{vert}} M_{l,m}^{\text{hor}} - M_{l,m+1}^{\text{hor}} M_{l,m}^{\text{vert}}) \xi_{l,m} = 0, \quad (2)$$

untuk semua $(l, m) \in \mathbb{Z}^2$.

Kondisi 2 dikenal dengan sebutan *compatibility condition*.

2 HASIL DAN PEMBAHASAN

2.1 Memperumum Persamaan $\Delta\Delta$ -mKdV

Tuwankotta dan Quispel (2012), telah melakukan upaya memperumum sebuah sistem dinamik diskrit melalui upaya memperbanyak parameter pada pasangan matriks Lax sistem tersebut. Hal ini bertujuan agar dalam mengkaji lebih banyak dinamika yang terjadi dari sebuah sistem dinamik diskrit sifat keterintegralan sistem senantiasa dipertahankan, (lihat [3] dan [4]). Dengan prosedur yang sama, berikut diperlihatkan upaya memperumum (1).

Pandang pasangan matriks Lax berikut ini

$$P_{l,m}^{\text{hor}} = \begin{pmatrix} \alpha_1 p & -\alpha_2 V_{l+1,m} \\ -\alpha_3 \left(\frac{k^2}{V_{l,m}}\right) & \alpha_4 p \left(\frac{V_{l+1,m}}{V_{l,m}}\right) \end{pmatrix} \text{ dan } P_{l,m}^{\text{vert}} = \begin{pmatrix} \beta_1 q & -\beta_2 V_{l,m+1} \\ -\beta_3 \frac{k^2}{V_{l,m}} & \beta_4 q \frac{V_{l,m+1}}{V_{l,m}} \end{pmatrix}. \quad (3)$$

Dengan *compatibility condition*, empat persamaan nonlinear berikut akan diperoleh

$$\begin{cases} k^2 (\alpha_3 \beta_2 - \alpha_2 \beta_3) V_{1+l,1+m} & = 0 \\ -k^2 (\alpha_3 \beta_2 - \alpha_2 \beta_3) & = 0 \\ p\alpha_1 \beta_2 V_{l,m} V_{l,1+m} - q\alpha_2 \beta_1 V_{l,m} V_{1+l,m} + \\ q\alpha_2 \beta_4 V_{l,1+m} V_{1+l,1+m} - p\alpha_4 \beta_2 V_{1+l,m} V_{1+l,1+m} & = 0 \\ -k^2 (p\alpha_1 \beta_3 V_{l,m} V_{l,1+m} - q\alpha_3 \beta_1 V_{l,m} V_{1+l,m}) - \\ k^2 (p(q\alpha_3 \beta_4 V_{l,1+m} V_{1+l,1+m} - p\alpha_4 \beta_3 V_{1+l,m} V_{1+l,1+m})) & = 0. \end{cases} \quad (4)$$

Agar konsisten satu dengan lainnya, maka parameter α_j dan β_j dengan $j = 1, 2, 3, 4$ dalam persamaan (4) harus konsisten.

Akibatnya, dari dua persamaan pertama diperoleh

$$\alpha_3 \beta_2 - \alpha_2 \beta_3 = 0. \quad (5)$$

Selain itu, dari dua persamaan terakhir dalam (4) diperoleh:

$$q(\alpha_3\beta_2 - \alpha_2\beta_3)(\beta_1V_{l,m}V_{1+l,m} - \beta_4V_{l,1+m}V_{1+l,1+m}) = 0. \quad (6)$$

Dari hubungan (5), persamaan (6) menjadi konsisten apabila $\alpha_2 = \alpha_3$ dan $\beta_2 = \beta_3$. Dengan demikian sistem dengan matriks Lax (3) akan konsisten jika $\alpha_2 = \alpha_3$ dan $\beta_2 = \beta_3$. Misalkan,

$$(\alpha, \beta) = (\alpha_1, \alpha_2, \alpha_2, \alpha_4, \beta_1, \beta_2, \beta_2, \beta_4).$$

Akibatnya, matriks Lax untuk sistem (1) yang diperumum dapat ditulis sebagai

$$P_{l,m}^{\text{hor}} = \begin{pmatrix} \alpha_1 p & -\alpha_2 V_{l+1,m} \\ -\alpha_2 \left(\frac{k^2}{V_{l,m}}\right) & \alpha_4 p \left(\frac{V_{l+1,m}}{V_{l,m}}\right) \end{pmatrix} \text{ dan } P_{l,m}^{\text{ver}} = \begin{pmatrix} \beta_1 q & -\beta_2 V_{l,m+1} \\ -\beta_2 \frac{k^2}{V_{l,m}} & \beta_4 q \frac{V_{l,m+1}}{V_{l,m}} \end{pmatrix}.$$

Dengan mensubstitusikan $P_{l,m}^{\text{hor}}$ dan $P_{l,m}^{\text{ver}}$ ke dalam kondisi kompatibel (2) maka akan diperoleh bentuk pemetaan-pemetaan yang diturunkan dari persamaan $\Delta\Delta$ -mKdV yang diperumum (generalized $\Delta\Delta$ -mKdV) yang tidak lain merupakan sebuah bagian dari keluarga pemetaan empat parameter, yakni

$$\theta_1 V_{l,m} V_{l,m+1} - \theta_2 V_{l+1,m} V_{l+1,m+1} - \theta_3 V_{l,m} V_{l+1,m} + \theta_4 V_{l,m+1} V_{l+1,m+1} = 0, \quad (7)$$

dengan $\theta_1 = \alpha_1 \beta_2 p$, $\theta_2 = \alpha_4 \beta_2 p$, $\theta_3 = \alpha_2 \beta_1 q$ dan $\theta_4 = \alpha_2 \beta_4 q$.

2.2 Solusi Gelombang Berjalan $\Delta\Delta$ -mKdV yang Diperumum

Dalam artikel ([4]) telah dideskripsikan secara lengkap penurunan persamaan $\Delta\Delta$ -sine Gordon untuk keadaan solusi gelombang berjalan. Dengan prosedur serupa, solusi gelombang berjalan dari (7) dapat diperoleh.

Pandang, hubungan solusi gelombang berjalan diskrit berikut:

$$V_{l,m} = V_n, \text{ dengan } n = z_1 l + z_2 m.$$

dengan z_1 dan z_2 merupakan bilangan bulat yang relatif prima. Mensubstitusikan ketentuan tersebut ke dalam (7) diperoleh

$$\theta_1 V_n V_{n+z_2} - \theta_2 V_{n+z_1} V_{n+z_1+z_2} - \theta_3 V_n V_{n+z_1} + \theta_4 V_{n+z_2} V_{n+z_1+z_2} = 0 \quad (8)$$

Persamaan (8) merupakan bentuk solusi gelombang berjalan dari $\Delta\Delta$ -mKdV yang diperumum. Untuk z_1 dan z_2 yang ditetapkan, persamaan (8) merupakan sebuah pemetaan dari $\mathbb{R}^{z_1+z_2} \rightarrow \mathbb{R}^{z_1+z_2}$. Dapat diperiksa bahwa persamaan (8) invarian untuk suatu transformasi $z_1 \rightarrow -z_1, p \rightarrow -p$, dan $z_1 \leftrightarrow z_2$.

Selain itu ia juga memenuhi sifat keperiodikan, yakni $(i + z_2, j - z_1)$.
Persamaan (8) ekuivalen dengan pemetaan

$$\begin{aligned} V'_{z_1+z_2-1} &= \frac{V_0(\theta_3 V_{z_1} - \theta_1 V_{z_2})}{(\theta_4 V_{z_2} - \theta_2 V_{z_1})} \\ V'_{z_1+z_2-2} &= V_{z_1+z_2-1} \\ &\vdots \\ V'_1 &= V_2 \\ V'_0 &= V_1 \end{aligned} \tag{9}$$

Dapat dicatat bahwa pemetaan dalam [1] dapat diperoleh dari (9) dengan menetapkan $\theta_1 = \theta_2 = p$ dan $\theta_3 = \theta_4 = q$.

3 Kesimpulan

Dari hasil dan pembahasan yang telah dikemukakan dalam bagian sebelumnya dapat disimpulkan bahwa bentuk solusi gelombang berjalan yang diperoleh dari pemetaan-pemetaan yang diturunkan dari persamaan $\Delta\Delta$ -mKdV yang diperumum (*generalized* $\Delta\Delta$ -mKdV) dapat diturunkan dengan terlebih dahulu menurunkan bentuk persamaan $\Delta\Delta$ -mKdV yang diperumum (*generalized* $\Delta\Delta$ -mKdV) yang merupakan sebuah bagian dari keluarga pemetaan empat parameter sebagai sebuah hasil pengembangan bentuk standar $\Delta\Delta$ -mKdV yang melibatkan dua parameter.

Acknowledgment

Penulis mengucapkan terima kasih kepada Dekan FMIPA Unila dan Ketua LPPM Unila atas dukungan dana yang diberikan melalui DIPA FMIPA Unila Tahun 2017.

References

- [1] Quispel, G.R.W., Capel, H.W., Papageorgiou, V.G., Nijhoff, F.W. (1991) *Integrable mappings derived from soliton equations*, Physica A **173** , pp. 243–266.
- [2] Roberts, J.A.G., Iatrou A., and Quispel,G.R.W., (2002) *Interchanging parameters and integrals in dynamical systems: the mapping case*, J.Phys.A: Math. Gen., **35**, 2309-2325.
- [3] Tuwankotta J., and Quispel, G., Dynamics Of 2-Dimensional Maps Derived From A Discrete Sine-Gordon Equations, 2012, Unpublished.
- [4] Zakaria L., and Tuwankotta, J.M., (2016): Dynamics and Bifurcations in a Two-Dimensional Maps Derived From a Generalized $\Delta\Delta$ sine-Gordon Equation, *Far East Journal of Dynamical Systems*, **28(3)**, pp 165–194.

PENDUGAAN BLUP DAN EBLUP (SUATU PENDEKATAN SIMULASI)

Nusyirwan
Jurusan Matematika
Universitas Lampung

ABSTRAK

Artikel ini bertujuan membandingkan model pendugaan BLUP dan EBLUP. Metode yang digunakan dalam penelitian ini adalah metode simulasi dengan menggunakan software R. Hasil simulasi menunjukkan Meningkatnya ukuran sampel dari $n=4$, $n=16$, dan $n=32$, berdasarkan nilai ARB dimana kondisi ragam dalam desa (Di) relatif homogen dan ragam antar desa (A) membesar ($A=10$, $A=20$, dan $A=30$) maka penduga BLUP lebih baik dari penduga EBLUP yang artinya penduga BLUP menjadi tak bias dan konsisten.

Kata kunci : BLUP; EBLUP ; General Linear Model.

A. PENDAHULUAN

Pendugaan yang akurat dari karakteristik populasi untuk desa merupakan tujuan penting dari banyak kalangan statistikawan. Pendugaan parameter secara langsung (*direct estimation*) berdasarkan data survei untuk menduga statistik desa kurang dapat diandalkan karena memiliki keragaman yang besar akibat contoh yang relatif sedikit.

Small area estimation merupakan salah satu solusi untuk memperbaiki hal tersebut, yaitu melakukan pendugaan dengan informasi-informasi tambahan yang bisa didapatkan dari desa lain yang serupa, survei terdahulu yang dilakukan di desa yang sama dan peubah lain yang berhubungan dengan peubah yang ingin diduga. Pendugaan seperti ini disebut dengan pengoshdugaan tak langsung (*indirect estimation*).

Ada beberapa metode yang sering digunakan di antaranya BLUP (*Best Linear Unbiased Prediction*) dan EBLUP (*Empirical Best Linear Unbiased Prediction*). BLUP merupakan pendugaan parameter yang meminimumkan *mean square error* diantara kelas-kelas pendugaan parameter *linear unbiased* lainnya [1]. BLUP dihasilkan dengan asumsi bahwa komponen ragam telah diketahui. Dalam prakteknya, komponen ragam sangat sulit diketahui, untuk itu diperlukan pendugaan terhadap komponen ragam ini melalui data contoh. Metode EBLUP mensubstitusi komponen ragam yang tidak diketahui ini dengan penduganya [2]. Simulasi ini bertujuan untuk mengkaji pendekatan metode BLUP dan EBLUP dalam pendugaan tak langsung pada desa dan membandingkan kedua metode tersebut secara empirik berdasarkan sifat-sifat penduga sampel yaitu ketakbiasan, efisiensi, dan konsisten.

B. LANDASAN TEORI

1. Penduga BLUP dan EBLUP

Best Linear Unbiased Predictor (BLUP) dahulu dikeluarkan dengan mengasumsikan bahwa komponen keragaman telah diketahui. Dalam prakteknya, komponen keragaman sangat sulit untuk diketahui. Untuk itu diperlukan pendugaan terhadap komponen keragaman ini melalui data contoh. Metode *Empirical Best Linear*

Unbiased Predictor (EBLUP) menggantikan komponen keragaman yang tidak diketahui ini dengan menduganya [2].

[3] memperkenalkan tiga metode (I, II dan III) dalam pendugaan komponen keragaman. Ketiga metode ini banyak digunakan sampai dengan tahun 1970an, ketika pendugaan dengan metode *maximum likelihood* (ML) dan *residual maximum likelihood* (REML) diperkenalkan. [4] memperlihatkan bahwa menggantikan komponen keragaman didalam BLUP dengan penduganya dapat menimbulkan bias. Tetapi [5] memperlihatkan bahwa 2 pendekatan (pertama, menduga komponen keragaman kemudian menggunakannya untuk menduga dan memprediksi parameter-parameter tetap dan komponen-komponen acak) dapat menghasilkan penduga yang tidak berbias [2].

[6] mengembangkan model $y_i = x_i' \beta + v_i + e_i$ sebagai dasar dalam pengembangan pendugaan desa. Selanjutnya diasumsikan bahwa β dan A tidak diketahui, tetapi D_i ($i = 1, 2, \dots, m$) diketahui. Penduga terbaik (BP) bagi $\theta_i = x_i' \beta + v_i$ jika β dan A diketahui adalah

$$\begin{aligned}\hat{\theta}_i^{BP} &= \hat{\theta}_i(y_i | \beta, D_i) \\ &= x_i' \beta + (1 - B_i)(y_i - x_i' \beta)\end{aligned}$$

dengan $B_i = D_i / (A + D_i)$ untuk $i = 1, 2, \dots, m$ dan

$$MSE(\hat{\theta}_i^{BP}) = \text{Var}(\theta_i | y_i, \beta, A) = (1 - B_i) D_i = g_{1i}(A).$$

Dalam praktek, baik β maupun A biasanya tidak diketahui sehingga untuk kasus A diketahui, β dapat diduga dengan metode kemungkinan maksimum atau metode momen $\beta^* = \hat{\beta}_i(A) = (X'V^{-1}X)^{-1} X'V^{-1}Y$ dengan $V = \text{Diag}(A + D_1, A + D_2, \dots, A + D_m)$. Kemudian dengan mensubstitusi β dengan β^* pada $\hat{\theta}_i^{BP}$, maka diperoleh

$$\begin{aligned}\hat{\theta}_i^{BLUP} &= \hat{\theta}_i(y_i | A) \\ &= x_i' \beta^* + (1 - B_i)(y_i - x_i' \beta^*)\end{aligned}$$

Menurut [1] $MSE(\hat{\theta}_i^{BLUP}) = g_{1i}(A) + g_{2i}(A)$, dengan $g_{2i}(A) = (D_i)^2 / (A + D_i) [X_i'(X'V^{-1}X)^{-1}X_i]$. Jika terlebih dahulu A diduga oleh $\hat{A}$ baik menggunakan metode ML, REML ataupun momen sehingga dengan mensubstitusi β oleh $\hat{\beta}$ dan A oleh $\hat{A}$ terhadap penduga BLUP ($\hat{\theta}_i^{BLUP}$), maka akan diperoleh suatu penduga baru

$$\begin{aligned}\hat{\theta}_i^{EBLUP} &= \hat{\theta}_i(y_i | \hat{A}) \\ &= x_i' \hat{\beta} + (1 - \hat{B}_i)(y_i - x_i' \hat{\beta})\end{aligned}$$

Jika didefinisikan MSE dari $\hat{\theta}_i^{EBLUP}$ adalah

$$\begin{aligned} \text{MSE}(\hat{\theta}_i^{EBLUP}) &= E(\hat{\theta}_i^{EBLUP} - \theta)^2 \\ &= \text{Var}(\hat{\theta}_i^{EBLUP}) + (\text{Bias } \hat{\theta}_i^{EBLUP})^2 \end{aligned}$$

persamaan tersebut dapat diuraikan menjadi

$$\begin{aligned} \text{MSE}(\hat{\theta}_i^{EBLUP}) &= \text{MSE}(\hat{\theta}_i^{BLUP}) + E(\hat{\theta}_i^{EBLUP} - \hat{\theta}_i^{BLUP})^2 \\ &= H1i(A) + H2i(A) \end{aligned}$$

$$\text{Dengan } H1i(A) = \text{MSE}(\hat{\theta}_i^{BLUP}) = g1i(A) + g2i(A) \quad H2i(A) = E(\hat{\theta}_i^{EBLUP} - \hat{\theta}_i^{BLUP})^2$$

Dengan menggunakan ekspansi deret Taylor untuk menduga $\text{MSE}(\hat{\theta}_i^{EBLUP})$ dan diperoleh $\text{MSE}(\hat{\theta}_i)_{PR} = g1i(\hat{A}) + g2i(\hat{A}) + 2g3i(\hat{A})$ dengan

$$g3i(\hat{A}) = \frac{2D_i}{m^2(\hat{A} + D_i)^2} \sum_{j=1}^m (\hat{A} + D_i)^2$$

2. SIFAT-SIFAT PENDUGA

Suatu penduga $\hat{\theta}$ dikatakan penduga takbias dari parameter θ jika $E(\hat{\theta}) = \theta$. Salah satu ukuran ketakbiasan adalah absolute relatif bias (ARB) dengan rumus sbb:

$$ARB = |1/k \sum \hat{Y}_d - Y_d| / Y_d \quad (1)$$

Sedangkan penduga $\hat{\theta}$ dikatakan penduga efisien dari suatu parameter θ jika penduganya tak bias dan mempunyai ragam minimum. Salah satu ukuran penduga yang baik dilihat dari mean square error (MSE) yang didefinisikan sebagai

$$\text{MSE}(\hat{\theta}) = \text{Var}(\hat{\theta}) + [\text{Bias}(\hat{\theta})]^2 \quad (2)$$

Dari persamaan (1) bila ada 2 penduga maka kita harus memilih MSE yang terkecil yang ekivalen memilih penduga yang efisien yaitu takbias dan mempunyai ragam minimum. Sedangkan akar dari MSE dibagi dengan Y adalah **relatif root mean square error (RRMSE)**.

C. Prosedur Simulasi

Adapun langkah-langkah simulasi sbb:

1. Menentukan x1, x2, dan x3 sebagai auxiliary variabel
2. Tentukan $\beta = (\beta_0, \beta_1, \beta_2, \beta_3)$
3. Hitung $X' \beta$
4. Bangkitkan v_i mengikuti $N(0, A)$
Shga $\Theta_i = X' \beta + v_i$

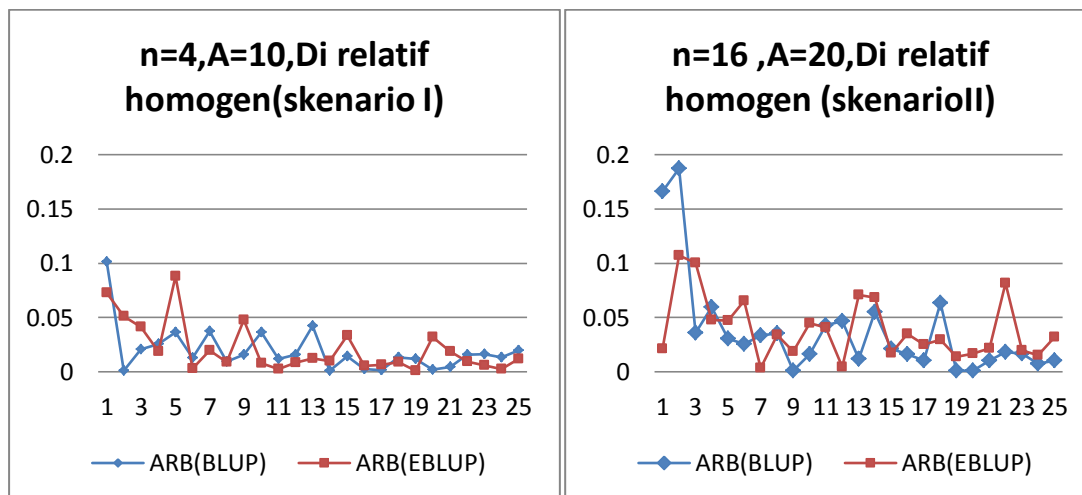
5. Bangkitkan untuk i e_{ij} ikut $N(0, D_i)$
 $Shga \ Y_i = \theta_i + e_i \rightarrow y_{ij} = \theta_i + e_{ij} \rightarrow \text{dapat } \bar{y}_i = \theta_i$ (penduga langsung)
6. Lakukan penduga thdp $\theta \rightarrow BLUP, EBLUP$ (diulang 10 kali)

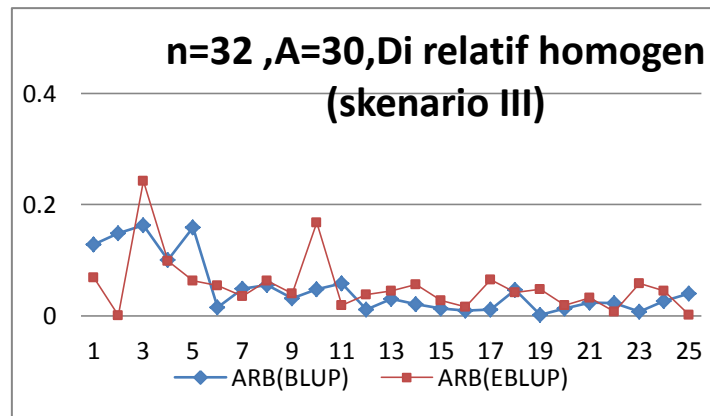
Simulasi menggunakan data hipotetik yang ditetapkan yaitu x_1, x_2 , dan x_3 untuk 25 desa, dengan nilai parameter β juga ditetapkan yaitu $\beta_0 = 0.5$, $\beta_1 = 1.0$, $\beta_2 = -1.5$, dan $\beta_3 = 2.0$. Ada beberapa skenario yang diambil untuk membangkitkan v_i ditentukan nilai ragam antar desa yang ditetapkan yaitu $A=10$, $A=20$ dan $A=30$ dengan ragam dalam desa (D_i) yang ditetapkan relatif homogen dan relatif heterogen. Sedangkan untuk membangkitkan e_i ditetapkan nilai sampel unit level dalam desa yaitu $n=4$, $n=16$ dan $n=32$. Skenario yang dibuat sebagai berikut : Skenario I yaitu $n=4$, $A=10$ dan D_i yang relatif homogen. Skenario II yaitu $n=16$, $A=20$ dan D_i yang relatif homogen. Skenario III yaitu $n=32$, $A=30$ dan D_i yang relatif homogen. Skenario IV yaitu $n=4$, $A=10$ dan D_i yang relatif heterogen. Skenario V yaitu $n=16$, $A=20$, dan D_i yang relatif heterogen. Dan skenario VI yaitu $n=32$, $A=30$ dan D_i yang relatif heterogen. Kemudian masing-masing skenario dilakukan pendugaan terhadap θ dengan metode BLUP dan EBLUP. Simulasi diulang sebanyak 10 kali. Hasil pendugaan θ dihitung nilai ARB, RMSE dan RRMSE.

D. Hasil Simulasi dan Pembahasan

1. Pembandingan Nilai ARB

Nilai ARB hasil simulasi skenario I, II, dan III untuk ragam dalam desa (D_i) relatif homogen dapat dilihat pada gambar 1 berikut ini :

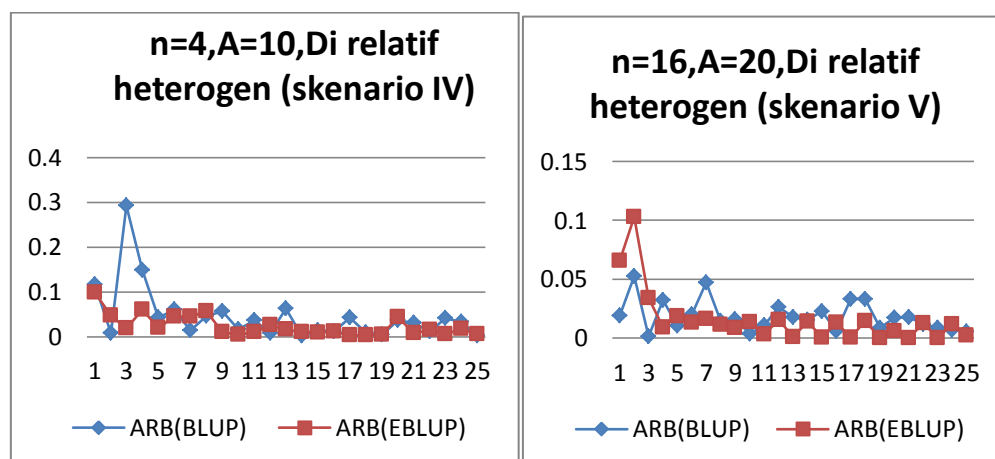


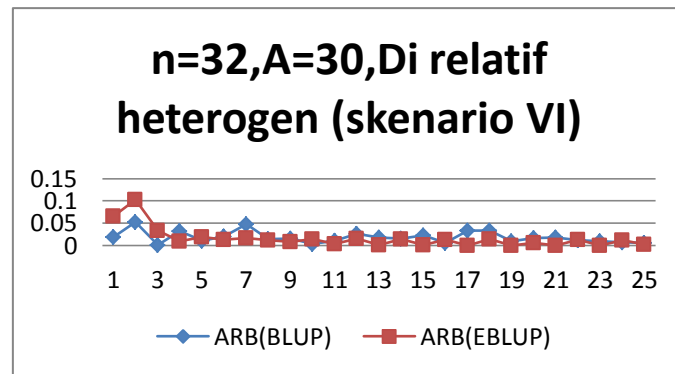


Gambar 1 Plot nilai ARB hasil simulasi skenario I, II, dan III

Pada gambar 1, skenario I, tampak untuk nilai ARB dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga EBLUP lebih baik dari BLUP. Skenario II untuk nilai ARB dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga BLUP lebih baik dari EBLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga BLUP masih lebih baik dari EBLUP.

Hal ini berarti dalam kondisi ragam dalam desa (Di) relatif homogen penduga EBLUP hanya lebih baik dari penduga BLUP pada ukuran sampel ($n=4$) berdasarkan nilai ARB. Sedangkan untuk ukuran sampel ($n=16$ dan $n=32$) penduga BLUP lebih baik dari penduga EBLUP berdasarkan nilai ARB. Nilai ARB hasil simulasi skenario IV, V, dan VI untuk ragam dalam desa (Di) relatif heterogen dapat dilihat pada gambar 2 berikut ini :



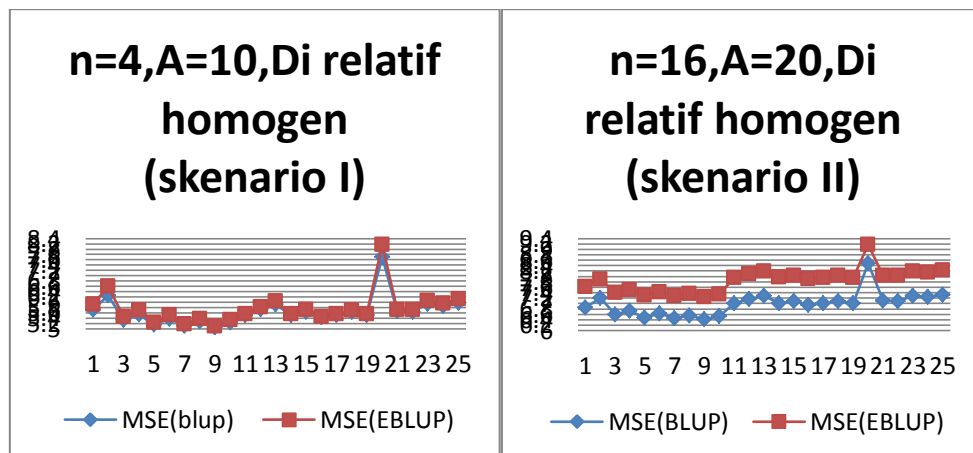


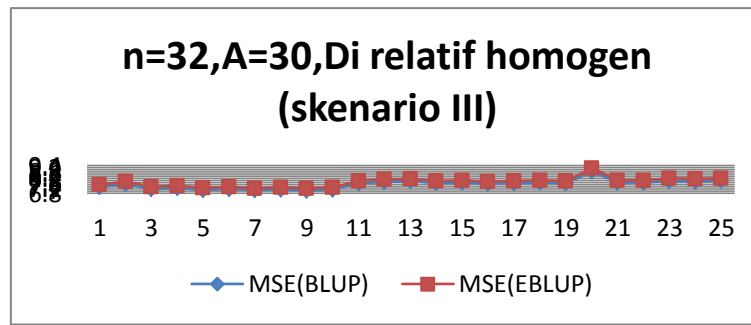
Gambar 2 Plot nilai ARB hasil simulasi skenario IV, V, dan VI

Pada gambar 2, skenario IV, tampak untuk nilai ARB dengan ragam dalam desa (Di) yang relatif heterogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga EBLUP lebih baik dari BLUP. Skenario V untuk nilai ARB dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga EBLUP masih lebih baik dari BLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga EBLUP masih lebih baik dari EBLUP. Hal ini berarti dalam kondisi ragam dalam desa (Di) relatif heterogen penduga EBLUP lebih baik dari penduga BLUP untuk ukuran sampel kecil ($n=4$ dan $n=16$) maupun ukuran sampel besar ($n=32$).

2. Pembandingan Nilai MSE

Nilai MSE hasil simulasi skenario I, II, dan III untuk ragam dalam desa (Di) relatif homogen dapat dilihat pada gambar 3 berikut ini:



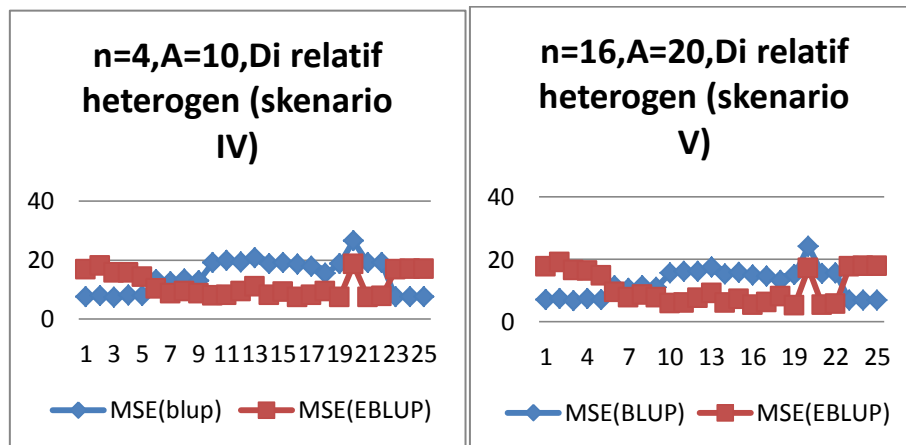


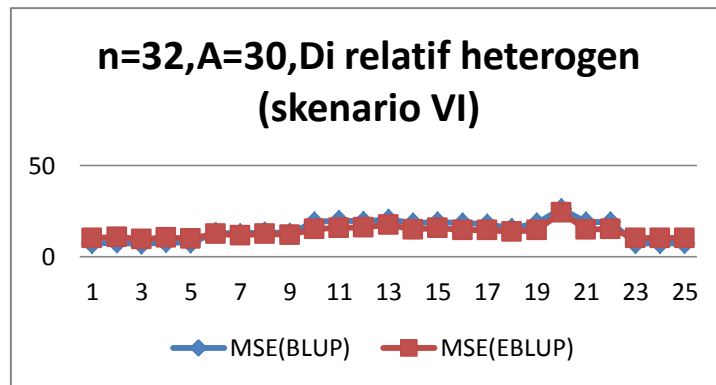
Gambar 3 Plot nilai MSE hasil simulasi skenario I, II, dan III

Pada gambar 3, skenario I, tampak untuk nilai MSE dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga BLUP lebih baik dari EBLUP. Skenario II untuk nilai MSE dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga BLUP masih lebih baik dari EBLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga BLUP masih lebih baik dari EBLUP.

Hal ini berarti dalam kondisi ragam dalam desa (Di) relatif homogen penduga BLUP lebih baik dari penduga EBLUP untuk ukuran sampel kecil ($n=4$ dan $n=16$) maupun ukuran sampel besar ($n=32$) berdasarkan ukuran MSE. Untuk $n=16$, penduga BLUP lebih berarti perbedaan dengan penduga EBLUP.

Nilai MSE hasil simulasi skenario IV, V, dan VI untuk ragam dalam desa (Di) relatif heterogen dapat dilihat pada gambar 4 berikut ini:





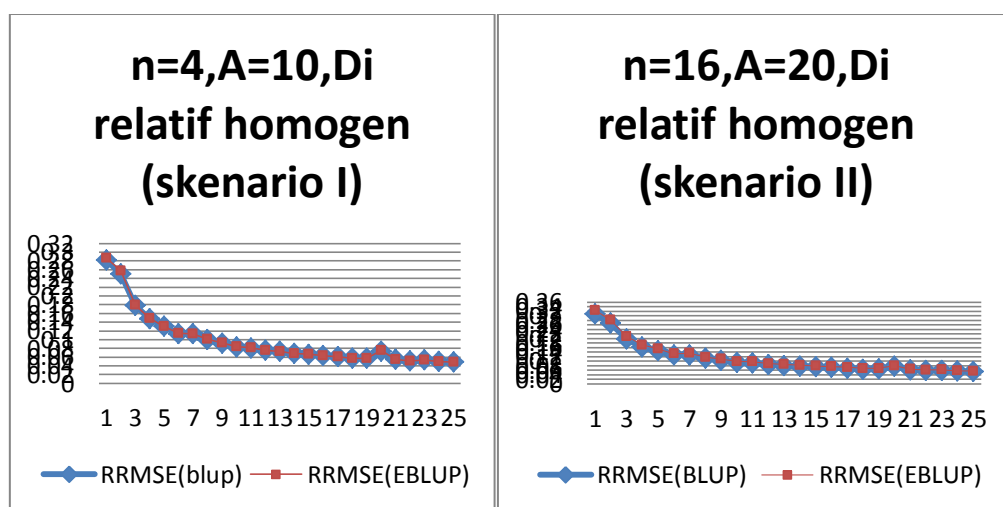
Gambar 4 Plot nilai MSE hasil simulasi skenario IV, V, dan VI

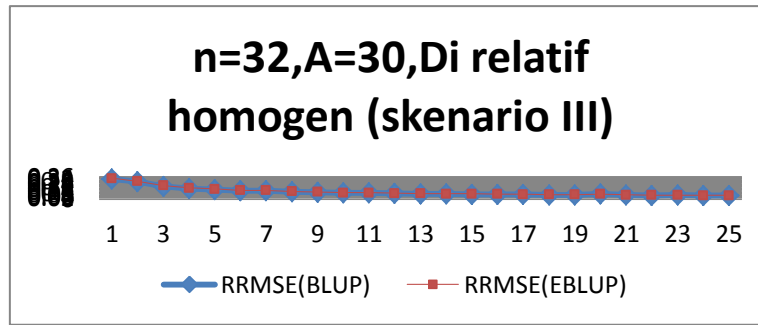
Pada gambar 4, skenario IV, tampak untuk nilai MSE dengan ragam dalam desa (D_i) yang relatif heterogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga EBLUP lebih baik dari BLUP. Skenario V untuk nilai MSE dengan ragam dalam desa (D_i) yang relatif heterogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga BLUP masih lebih baik dari EBLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga EBLUP masih lebih baik dari BLUP.

Hal ini berarti dalam kondisi ragam dalam desa (D_i) relatif heterogen dan ragam antar desa (A) membesar penduga EBLUP lebih baik dari penduga BLUP untuk ukuran sampel kecil ($n=4$ dan $n=16$) maupun ukuran sampel besar ($n=32$) berdasarkan ukuran MSE.

a. Perbandingan Nilai RRMSE

Nilai RRMSE hasil simulasi skenario I, II, dan III untuk ragam dalam desa (D_i) relatif homogen dapat dilihat gambar 5 berikut ini:

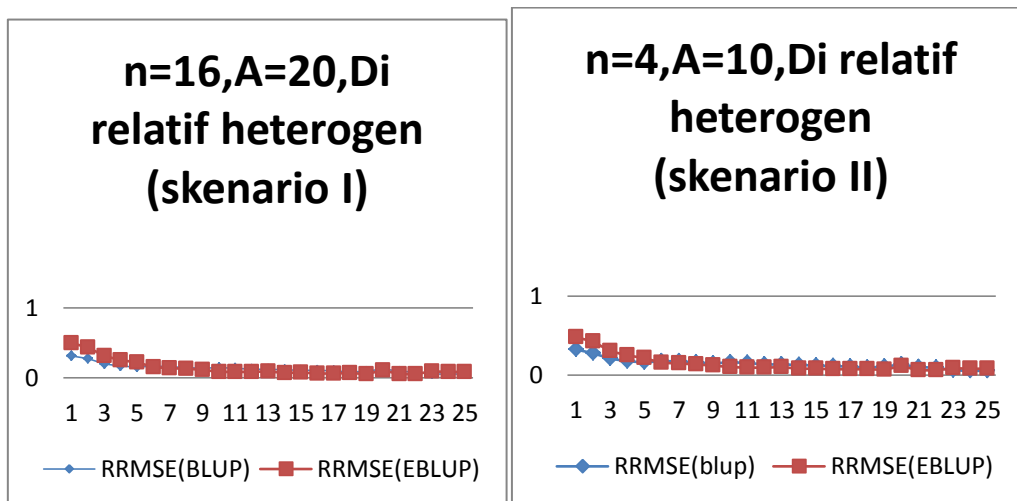


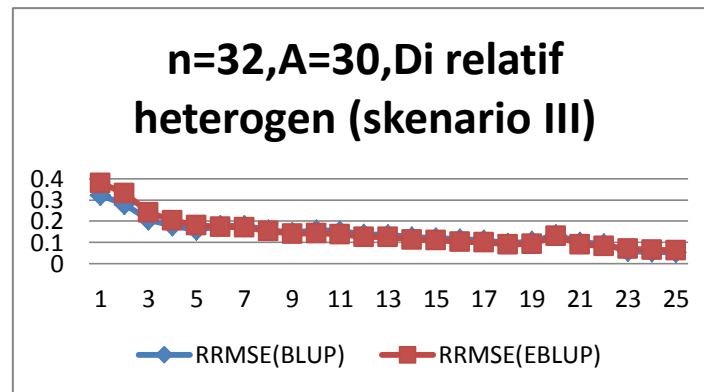


Gambar 5 Plot nilai RRMSE hasil simulasi skenario I, II, dan III

Pada gambar 5, skenario I, tampak untuk nilai RRMSE dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga BLUP lebih baik dari EBLUP. Skenario II untuk nilai MSE dengan ragam dalam desa (Di) yang relatif heterogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga BLUP masih lebih baik dari EBLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga BLUP masih lebih baik dari EBLUP.

Hal ini berarti dalam kondisi ragam dalam desa (Di) relatif homogen dan ragam antar desa (A) membesar penduga BLUP lebih baik dari penduga EBLUP untuk ukuran sampel kecil ($n=4$ dan $n=16$) maupun ukuran sampel besar ($n=32$) berdasarkan ukuran RRMSE. Nilai RRMSE hasil simulasi skenario VI, V, dan VI untuk ragam dalam desa (Di) relatif heterogen dapat dilihat pada gambar 6 berikut ini:





Gambar 6 Plot nilai RRMSE hasil simulasi skenario IV, V, dan VI

Pada gambar 6, skenario IV, tampak untuk nilai RRMSE dengan ragam dalam desa (Di) yang relatif homogen dengan ukuran sampel kecil ($n=4$) dan ragam antar desa ($A=10$) penduga EBLUP lebih baik dari BLUP. Skenario II untuk nilai MSE dengan ragam dalam desa (Di) yang relatif heterogen dengan ukuran sampel kecil ($n=16$) dan ragam antar desa ($A=20$) tampak penduga EBLUP masih lebih baik dari BLUP. Begitu juga untuk ukuran sampel besar ($n=32$) dengan $A=30$ bahwa penduga EBLUP masih lebih baik dari BLUP.

Hal ini berarti dalam kondisi ragam dalam desa (Di) relatif heterogen dan ragam antar desa (A) membesar penduga EBLUP lebih baik dari penduga BLUP untuk ukuran sampel kecil ($n=4$ dan $n=16$) maupun ukuran sampel besar ($n=32$) berdasarkan ukuran RRMSE.

E. KESIMPULAN

- Meningkatnya ukuran sampel dari $n=4$, $n=16$, dan $n=32$, berdasarkan nilai ARB dimana kondisi ragam dalam desa (Di) relatif homogen dan ragam antar desa (A) membesar ($A=10$, $A=20$, dan $A=30$) maka penduga BLUP lebih baik dari penduga EBLUP yang artinya penduga BLUP menjadi tak bias dan konsisten.
- Berdasarkan nilai MSE dan RRMSE pada ukuran sampel ($n=4$, $n=16$ dan $n=32$), dimana kondisi ragam dalam desa (Di) relatif homogen dan ragam antar desa membesar yaitu $A=10$, $A=20$, dan $A=3$ maka penduga BLUP lebih baik dari penduga BLUP yang artinya penduga BLUP menjadi efisien yaitu tak bias dan mempunyai ragam minimum
- Berdasarkan nilai MSE dan RRMSE pada ukuran sampel ($n=4$, $n=16$ dan $n=32$), dimana kondisi ragam dalam desa (Di) relatif heterogen dan ragam antar desa membesar ($A=10$, $A=20$, dan $A=30$) maka penduga EBLUP lebih baik dari penduga BLUP yang artinya penduga EBLUP menjadi efisien yaitu tak bias dan mempunyai ragam minimum
- Dalam kondisi ragam dalam desa (Di) relatif homogen untuk nilai ARB, MSE, dan RRMSE ternyata penduga BLUP lebih baik dari EBLUP, sebaliknya ragam dalam desa (Di) relatif heterogen penduga EBLUP lebih baik dari BLUP.

KEPUSTAKAAN

- [1] Gosh M., Rao J.N.K. (1994), *Small area estimation: an appraisal*, Statistical Science, Vol 9, No 1, pp. 55-93
- [2] Saei, A. Dan R. Chambers, (2003), “ *Small Area Estimation : A Review of Methods Based on the Application of Mixed Models*”, SRI Methodologi Working paper M03/16, University of Southampton, UK.
- [3] Henderson, C.R. (1963),” *Selection index and expected genetic advance*” Nas-NRC Publ. 982
- [4] Henderson, C.R. (1975),” *Best linear unbiased estimation and prediction under a selection model.*” Biometrics 32:69
- [5] Kacker, R.N. and Harville, D.A. (1981). *Unbiasedness of twostage estimation and prediction procedures for mixed linear models*. Communications in Statistics - Theory and Methods, Series a 10, 1249-1261.
- [6] Fay, R.E. and Herriot, R.A, (1979), “ *Estimates of income for small places : an application of James-Stein procedures to Censuss data:*, Journal pf the american Statistical Association, Vol. 74, p : 269-277

