

ISSN : 2527-5836

e-ISSN : 2528-0074

Vol. 10 No. 3, September 2025

JISKa

Jurnal Informatika Sunan Kalijaga

Jurusan Teknik Informatika
Fakultas Sains dan Teknologi
UIN Sunan Kalijaga Yogyakarta



JISKa Editorial Team

September 2025 Edition

Editor in Chief

Muhammad Taufiq Nuruzzaman, UIN Sunan Kalijaga Yogyakarta, Indonesia

Editorial Board

Aang Subiyakto, UIN Syarif Hidayatullah Jakarta, Indonesia

Agung Fatwanto, UIN Sunan Kalijaga Yogyakarta, Indonesia

Andang Sunarto, UIN Fatmawati Sukarno Bengkulu, Indonesia

Deokjai Choi, Chonnam National University, South Korea

Elyor Kodirov, Opentrons Labworks Inc., United Kingdom

Hamdani, Universitas Mulawarman Samarinda, Indonesia

Muhammad Anshari, Universiti Brunei Darussalam, Brunei Darussalam

Muhammad Syafrudin, Sejong University, South Korea

Nashrul Hakiem, UIN Syarif Hidayatullah Jakarta, Indonesia

Noor Akhmad Setiawan, Universitas Gadjah Mada, Indonesia

Copy Editor and Layout Editor

Sekar Minati, Victoria University of Wellington, New Zealand

Journal Manager and Technical Support

Eko Hadi Gunawan, UIN Sunan Kalijaga Yogyakarta, Indonesia

Muhammad Galih Wonoseto, UIN Sunan Kalijaga Yogyakarta, Indonesia

Reviewers

Agung Dewandaru, Institut Teknologi Bandung, Indonesia

Agus Mulyanto, UIN Sunan Kalijaga Yogyakarta, Indonesia

Ahmad Fathan Hidayatullah, Universitas Islam Indonesia Yogyakarta, Indonesia

Alam Rahmatulloh, Universitas Siliwangi Tasikmalaya, Indonesia

Anggi Rizky Windra Putri, Universitas Aisyiyah Yogyakarta, Indonesia

Ardiansyah Musa Efendi, Singapore Chipset Algorithm Design Lab, Huawei, Singapore

Bambang Sugiantoro, UIN Sunan Kalijaga Yogyakarta, Indonesia

Enny Itje Sela, Universitas Teknologi Yogyakarta, Indonesia

Ganjar Alfian, Universitas Gadjah Mada, Indonesia

Mandahadi Kusuma, UIN Sunan Kalijaga Yogyakarta, Indonesia

Maria Ulfah Siregar, UIN Sunan Kalijaga Yogyakarta, Indonesia

Millati Pratiwi, Pusan National University, South Korea

Muhammad Dzulfikar Fauzi, Telkom University Surabaya, Indonesia

Muhammad Habibi, Universitas Jenderal Achmad Yani Yogyakarta, Indonesia

Muhammad Rifqi Maarif, Universitas Jenderal Achmad Yani Yogyakarta, Indonesia

Mohd. Fikri Azli bin Abdullah, Multimedia University, Malaysia

M. Alex Syaekhoni, Who's Good, South Korea

Niki Min Hidayati Robbi, Universitas Gadjah Mada, Indonesia

Norma Latif Fitriyani, Sejong University Seoul, South Korea

Okfalisa, UIN Sultan Syarif Kasim Riau, Indonesia

Oman Somantri, Politeknik Negeri Cilacap, Indonesia

Puguh Jayadi, Universitas PGRI Madiun, Indonesia

Puji Winar Cahyo, Universitas Jenderal Achmad Yani Yogyakarta, Indonesia

Qorry Aina Fitroh, UIN KH. Abdurrahman Wahid Pekalongan, Indonesia

Ridho Surya Kusuma, Universitas Siber Muhammadiyah, Yogyakarta, Indonesia

Rischan Mafrur, Macquarie University, Sydney, Australia

Shofwatul Uyun, UIN Sunan Kalijaga Yogyakarta, Indonesia

Sumarsono, UIN Sunan Kalijaga, Indonesia

Sunu Wibirama, Universitas Gadjah Mada, Indonesia

Tundo, Sekolah Tinggi Ilmu Komputer Cipta Karya Informatika (STIKOM CKI), Indonesia

Windra Swastika, Universitas Ma Chung, Indonesia

Yudistira Dwi Wardhana Asnar, Institut Teknologi Bandung, Indonesia

ISSN : 2527-5836

e-ISSN: 2528-0074

JISKa (Jurnal Informatika Sunan Kalijaga)

Vol. 10, No. 3, SEPTEMBER 2025

TABLE OF CONTENT

Peramalan Nilai Saham BBCA Melalui Pendekatan Time Series Menggunakan Teknik Exponential Smoothing Febri Liantoni, Ondihon Simanjuntak	259-266
Enhancing Diabetes Classification Using a Relaxed Online Maximum Margin Algorithm Dyan Avando Meliala, Arum Kurnia Sulistyawati, Mohammad Diqi, Marselina Endah Hiswati, Tadem Vergi Kristian	267-278
Analisis Ketertarikan Pengguna Microsoft Excel Online untuk Pengolahan Data Silsilah Keluarga Menggunakan TAM dan TPB Fathur Rachman Nufaily, Maria Ulfah Siregar	279-293
Algoritma K-Means dan Analisis Komponen Utama untuk Mengatasi Multikolinearitas pada Pengelompokan Kabupaten Tertinggal Firna Aviliana, Putriaji Hendikawati	294-306
Spatial Decision Support System to Determine the Feasibility of Evacuation Posts in Natural Disasters Nuril Afni Alviola, Agung Teguh Wibowo Almais, A'la Syauqi, Totok Chamidy, Puspa Miladin Nuraida Safitri A Basid, Anisa Anisa, M. Dafa Wardana	307-318
Implementasi Metode YOLOv8 Mendeteksi Komputer Aktif dengan Subjek Layar Monitor Frisky Wijaya, Dedy Hermanto	319-330
Klasifikasi Hewan Anjing, Kucing, dan Harimau Menggunakan Metode Convolutional Neural Network (CNN) Murdifin Murdifin, Shofwatul Uyun	331-340

Arsitektur Microservice untuk Optimalisasi Aplikasi Eco-Maps dalam Mendukung Kampus Ramah Lingkungan	341-350
Alam Rahmatulloh, Rohmat Gunawan, Randi Rizal	
Analisis Efektivitas Metode Filtering dan Intersection dalam Analisis Data Permukaan Bangunan dengan QGIS	351-363
Prana Wijaya Pratama Nandana, Muhammad Faisal	
Analisis Sistem Deteksi Citra untuk Optimalisasi Pengawasan Lalu Lintas Udara Menggunakan Algoritma YOLOv5	364-376
Astika Ayuningtyas, Imam Riadi, Anton Yudhana	

Peramalan Nilai Saham BBKA Melalui Pendekatan Time Series Menggunakan Teknik Exponential Smoothing

Febri Liantoni ^{(1)*}, Ondihon Simanjuntak ⁽²⁾

Departemen Pendidikan Teknik Informatika dan Komputer, Universitas Sebelas Maret,
Surakarta, Indonesia

e-mail : {febri.liantoni,ondihonsimanjuntak5902}@gmail.com.

* Penulis korespondensi.

Artikel ini diajukan 21 Mei 2025, direvisi 20 Juni 2025, diterima 2 Juli 2025, dan dipublikasikan 30 September 2025.

Abstract

Forecasting stock prices plays a crucial role in shaping investment strategies within the financial market. This article aims to predict the stock prices of Bank Central Asia (BBKA), a prominent entity in the Indonesian banking sector. Employing a time series methodology, this study utilizes the Exponential Smoothing technique to anticipate the fluctuations in BBKA's share prices. Meanwhile, the dataset used is the BBKA share price data from April 2001 to early January 2023. The final error rate in this forecast is 10%.

Keywords: *Stock Price Forecasting, Time Series, Exponential Smoothing, Bank Central Asia (BBKA), Financial Markets, Investment, Risk*

Abstrak

Meramal harga saham memainkan peran penting dalam membentuk strategi investasi di pasar keuangan. Artikel ini bertujuan untuk meramalkan harga saham Bank Central Asia (BBKA), salah satu perusahaan terkemuka di sektor perbankan Indonesia. Penelitian ini mengadopsi pendekatan *time series* dan menerapkan metode Exponential Smoothing guna memprediksi harga saham BBKA. Sedangkan dataset yang digunakan adalah data harga saham BBKA dari bulan April 2001 sampai awal Januari 2023. Tingkat kesalahan akhir dalam peramalan ini adalah 10%.

Kata Kunci: *Peramalan Harga Saham, Time Series, Exponential Smoothing, Bank Central Asia (BBKA), Pasar Keuangan, Investasi, Risiko*

1. PENDAHULUAN

Pasar saham telah lama menjadi tempat yang menarik bagi investor dan pelaku pasar keuangan. Investasi saham adalah salah satu cara terkemuka untuk mengembangkan kekayaan. Berinvestasi dalam saham menawarkan potensi pendapatan yang signifikan namun diiringi oleh resiko yang cukup tinggi (Sofiyati et al., 2024; Tang et al., 2020). Investor berinvestasi untuk memperoleh keuntungan yang diukur dari besaran return atau tingkat pengembalian saham tersebut. Banyak yang ingin mencari keuntungan yang tinggi tetapi semakin tinggi pula risiko yang diperoleh. Meskipun saat ini banyak sekali penduduk di Indonesia yang berpikir bahwa saham adalah sesuatu yang rumit dan memiliki risiko yang besar (Juniarti et al., 2024). Sebagai seorang *investor* harus memahami fundamental terhadap saham yang akan dibeli dan harus melakukan diversifikasi saham yang ingin di beli (Tambunan, 2020). Maka dari itu, prediksi nilai saham menjadi sangat penting dalam membantu para investor dan pedagang saham untuk mengambil keputusan yang cerdas.

Peramalan harga saham merupakan upaya sistematis untuk memahami dan memperkirakan arah pergerakan harga saham di masa depan. Proses ini dilakukan dengan menganalisis data historis serta mengaitkannya dengan berbagai faktor yang dapat memengaruhi pasar, seperti kondisi ekonomi, performa perusahaan, serta perubahan kebijakan atau situasi global (Furizalz et al., 2024; Hubbi & Mustofa, 2024). Tujuannya adalah memberikan gambaran yang lebih jelas bagi investor dalam mengambil keputusan yang bijak dan terinformasi. Peramalan harga saham sangat penting karena mempengaruhi dalam pengambilan keputusan (Hubbi & Mustofa, 2024).



Diperkuat dalam penelitian tentang peramalannya tentang pakaian, bahwa peramalan dapat mempengaruhi pengambilan keputusan terkait penentuan jumlah produksi barang yang akan disediakan. Bank Central Asia (BBCA), sebagai salah satu lembaga perbankan terbesar dan paling kokoh di Indonesia, menjadi sorotan utama bagi para investor di negeri ini. Dengan kapitalisasi pasar yang signifikan dan peran pentingnya dalam ekonomi Indonesia, peramalan harga saham BBCA menjadi topik yang sangat relevan dan bermanfaat.

Dalam upaya untuk menghadapi fluktuasi harga saham BBCA, pendekatan *time series* menjadi alat yang sangat efektif. Pendekatan ini memperlakukan data harga saham sebagai suatu rangkaian waktu yang menggambarkan perkembangan harga saham dari waktu ke waktu. Exponential Smoothing merupakan salah satu teknik yang sering diterapkan secara umum dalam pendekatan *time series*. Memperkuat argumen tersebut dalam penelitiannya terkait peramalan jumlah mahasiswa baru, dimana ia menyimpulkan bahwa penggunaan metode Exponential Smoothing mampu mengatasi kekurangan yang ada pada *moving average* (Landia, 2020). Metode Exponential Smoothing memberikan bobot yang menurun secara eksponensial pada observasi masa lalu (Cetin & Yavuz, 2021; Svetunkov et al., 2022). Data yang lebih baru memiliki pengaruh yang lebih besar terhadap hasil peramalan dibandingkan dengan observasi yang lebih lama. Teknik ini terbukti efektif dalam menghasilkan prediksi jangka pendek, terutama ketika data menunjukkan sedikit pola tren dan musiman (Svetunkov et al., 2023).

Time series adalah sekumpulan observasi data yang memiliki frekuensi waktu berupa detik, menit, hingga tahun. Exponential Smoothing adalah suatu algoritma yang evolusioner dari konsep rata-rata bergerak (*moving average*) dan secara umum digunakan untuk mengatasi tantangan dalam data *time series* (Baharaeen & Masud, 1986). Keunggulan Exponential Smoothing terletak pada ketergantungan operasional dan kemampuannya untuk menyesuaikan data melalui pengaturan parameter *alpha* (Al Ihsan et al., 2020). Meskipun demikian, dalam penelitiannya diungkapkan bahwa formulasi metode Exponential Smoothing kurang sesuai untuk melakukan peramalan terhadap data yang mengalami perubahan secara mendadak dan drastis (Cunha & Pereira, 2024; Gibran et al., 2021; Kumar et al., 2024).

Dalam membandingkan metode Autoregressive Integrated Moving Average (ARIMA) dan Double Exponential Smoothing untuk peramalan harga saham tiga perusahaan dari nilai *Earning Per Share* (EPS) tertinggi dari saham-saham yang tergabung di LQ45, disimpulkan bahwa pendekatan Exponential Smoothing unggul dibandingkan pendekatan ARIMA karena memiliki *Mean Absolute Percentage Error* (MAPE) yang lebih rendah. Konsisten dengan temuan penelitian terkait peramalan harga saham, disimpulkan bahwa kinerja metode Exponential Smoothing sangat unggul dibandingkan pendekatan ARIMA (Zahrunnisa et al., 2021).

Dari hasil penelitian diharapkan dapat memberikan wawasan yang berguna untuk investor, pedagang saham, dan para praktisi pasar keuangan yang tertarik dalam meramalkan harga saham BBCA dan mengelola risiko investasi mereka. Dengan pemahaman yang lebih baik tentang pergerakan harga saham BBCA, para pemangku kepentingan dapat mengambil keputusan investasi yang lebih baik, berdasarkan analisis data historis yang kuat dan metode peramalan yang efektif.

2. METODE PENELITIAN

Sebelum memulai tahap penelitian, acuan data yang diterapkan mencakup rentang data harga saham BBCA mulai dari bulan Januari hingga Desember pada tahun 2022. Data ini diperoleh melalui platform www.kaggle.com. Dalam proses penelitian, peneliti memanfaatkan dukungan dari aplikasi perangkat lunak Microsoft Excel untuk memudahkan dan membantu dalam pengelolaan data dan analisis.

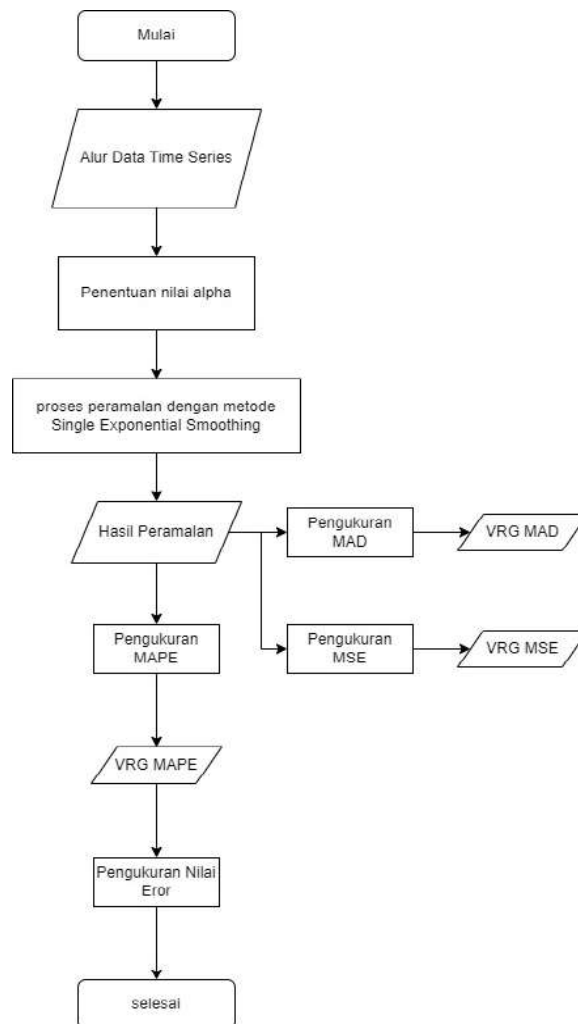
Penelitian ini mengaplikasikan *forecasting* dengan metode Exponential Smoothing. Metode ini menerapkan rumus yang ditunjukkan pada Pers. (1). Dengan F_{t+1} adalah nilai *forecasting* berikutnya. Alpha (α) adalah faktor pembobotan yang disebut sebagai konstanta pemulusan dan F_t adalah perkiraan yang telah ditentukan sebelumnya untuk periode sekarang. Dengan



memanfaatkan rumus tersebut, pelaksanaan proses peramalan menggunakan metode Single Exponential Smoothing (SES) dapat dijalankan.

$$F_{t+1} = \alpha D_t + (1 - \alpha) F_t \quad (1)$$

Peramalan metode SES merupakan salah satu pendekatan yang cukup populer di kalangan peneliti (Arini et al., 2023; Aziz et al., 2025; Putra et al., 2024). Kepopulerannya tidak lepas dari karakteristiknya yang sederhana namun efektif, serta kemudahan dalam penerapannya. SES hanya memerlukan satu parameter utama, yaitu *alpha*, yang berfungsi untuk mengatur seberapa besar pengaruh data masa lalu terhadap hasil peramalan terkini. Semakin tinggi nilai *alpha*, semakin besar bobot data terbaru dalam menentukan nilai ramalan, sehingga metode ini sangat fleksibel untuk digunakan dalam berbagai kondisi data.



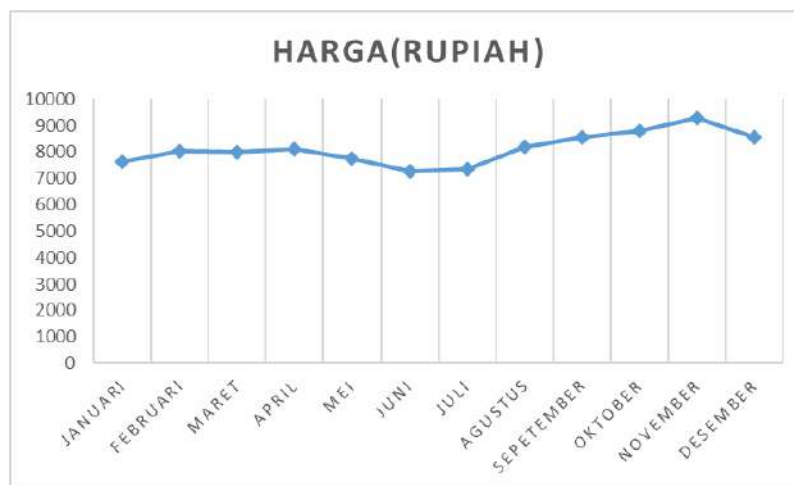
Gambar 1 Flowchart Single Exponential Smoothing

Gambar 1 merupakan *flowchart* Metode peramalan Single Exponential Smoothing (SES) yang digunakan untuk memprediksi nilai masa depan dalam data *time series*. Dalam proses ini, data *time series* dimasukkan dan dibagi menjadi dua bagian, yaitu data historis yang digunakan untuk melatih model dan data aktual yang akan dijadikan patokan untuk peramalan. Pengguna kemudian menentukan nilai *alpha*, sebuah parameter yang mengontrol seberapa besar bobot yang diberikan pada data historis dibandingkan dengan data aktual. Dengan menggunakan



α , peramalan SES dilakukan dengan cara menghitung nilai ramalan baru berdasarkan data historis sebelumnya dan hasil peramalan sebelumnya.

Selanjutnya, dilakukan perhitungan *error* dengan menghitung MAD, MSE, dan MAPE, yang merupakan metrik evaluasi performa peramalan. MAD mengukur rata-rata deviasi absolut antara nilai ramalan dan nilai aktual, sedangkan MSE mengukur rata-rata dari kuadrat deviasi antara nilai ramalan dan nilai aktual (Ibrahim et al., 2023). Sementara MAPE mengukur rata-rata persentase deviasi absolut antara nilai ramalan dan nilai aktual terhadap nilai aktual (Liantoni & Agusti, 2020). Setelah proses peramalan selesai, hasil peramalan berserta metrik evaluasinya ditampilkan kepada pengguna. Dengan pendekatan yang sederhana namun efektif, metode SES menjadi salah satu pilihan yang populer dalam melakukan peramalan pada data *time series*. Contoh data yang digunakan dalam penelitian ini ditunjukkan pada Gambar 2.



Gambar 2 Grafik Harga Saham BBKA per Januari - Desember 2022

3. HASIL DAN PEMBAHASAN

Dalam era digital saat ini, kemampuan perangkat lunak seperti Microsoft Excel tidak hanya terbatas pada fungsi pengolahan angka semata, tetapi juga berkembang menjadi alat analisis data yang cukup andal. Salah satu fitur yang menunjukkan kecanggihan Excel adalah penerapan metode Exponential Smoothing, yang digunakan dalam peramalan data deret waktu. Melalui metode ini, pengguna dapat memodelkan tren historis untuk memperkirakan nilai-nilai masa depan dengan lebih terstruktur dan efisien.

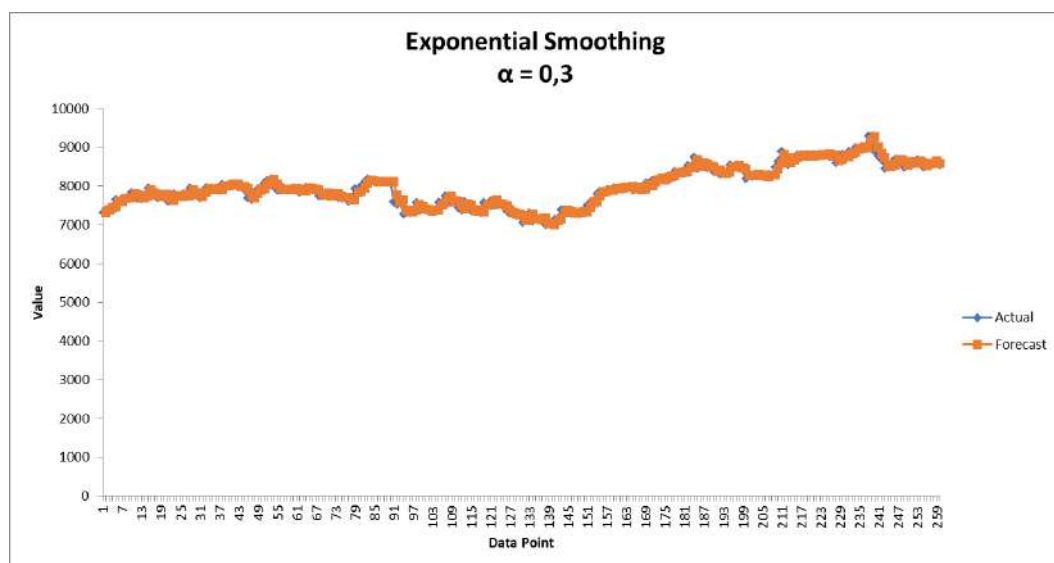
Dalam proses penerapannya, nilai α menjadi parameter penting yang menentukan seberapa besar pengaruh data terbaru terhadap hasil peramalan. Pada analisis ini, digunakan tiga variasi nilai α , yaitu 0,3, 0,5, dan 0,9, yang masing-masing mencerminkan tingkat sensitivitas yang berbeda terhadap fluktuasi data. Pemilihan nilai-nilai tersebut tidak dilakukan secara sembarangan, melainkan mempertimbangkan kebutuhan akan akurasi serta kemampuan model dalam menangkap perubahan tren yang terjadi dalam data.

Penggunaan nilai α yang beragam memberikan fleksibilitas bagi model untuk menyesuaikan diri dengan karakteristik unik dari data yang dianalisis. Dengan demikian, hasil peramalan yang diperoleh menjadi lebih adaptif dan mencerminkan dinamika sebenarnya dari pola yang diamati. Pendekatan ini membantu pengguna, baik peneliti maupun praktisi, untuk memperoleh pemahaman yang lebih dalam terhadap arah perkembangan data di masa mendatang dan mengambil keputusan yang lebih tepat berdasarkan hasil analisis tersebut.

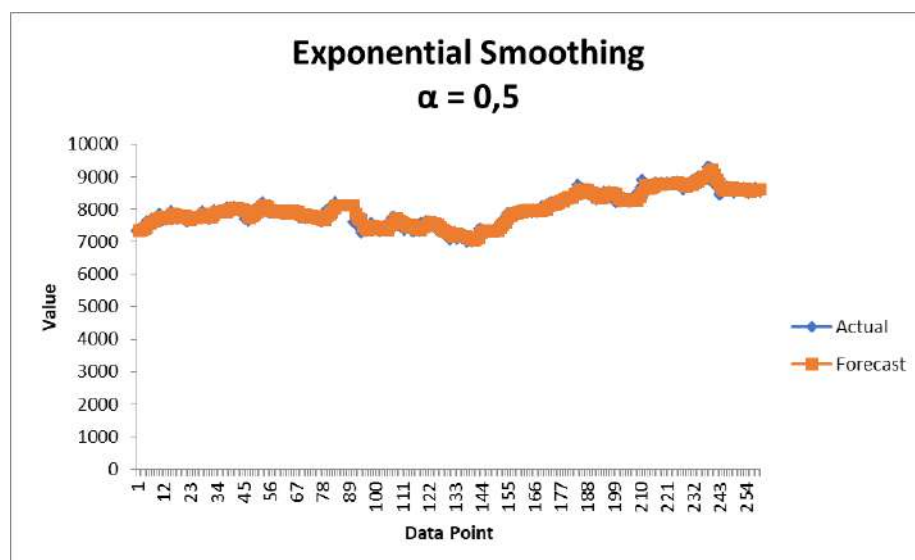
Hasil analisis data menggunakan metode Single Exponential Smoothing dengan berbagai nilai parameter pelincinan (α) menunjukkan perbedaan yang cukup signifikan dalam pola



peramalan yang dihasilkan. Pada Gambar 3, diperlihatkan hasil peramalan menggunakan nilai alpha sebesar 0,3, di mana model lebih responsif terhadap tren historis jangka panjang dan cenderung menghasilkan kurva yang lebih halus. Hal ini mengindikasikan bahwa dengan nilai alpha yang lebih rendah, bobot terhadap data masa lalu lebih besar sehingga perubahan harga terkini tidak terlalu memengaruhi hasil estimasi secara drastis. Sementara itu, pada Gambar 4 ditampilkan hasil peramalan dengan alpha sebesar 0,5 yang mencerminkan tingkat respons sedang terhadap fluktuasi data terbaru—menyeimbangkan pengaruh data historis dan nilai observasi terkini. Selanjutnya, Gambar 5 menyajikan hasil peramalan dengan alpha sebesar 0,9, di mana model menjadi sangat peka terhadap perubahan nilai terbaru, sehingga grafik peramalan menunjukkan respons yang lebih tajam terhadap setiap pergeseran harga. Perbandingan antar ketiga grafik ini memberikan gambaran yang komprehensif mengenai dampak pemilihan parameter smoothing terhadap akurasi dan karakteristik pola peramalan, yang pada akhirnya dapat membantu dalam menentukan konfigurasi model yang paling sesuai dengan sifat data saham yang dianalisis.

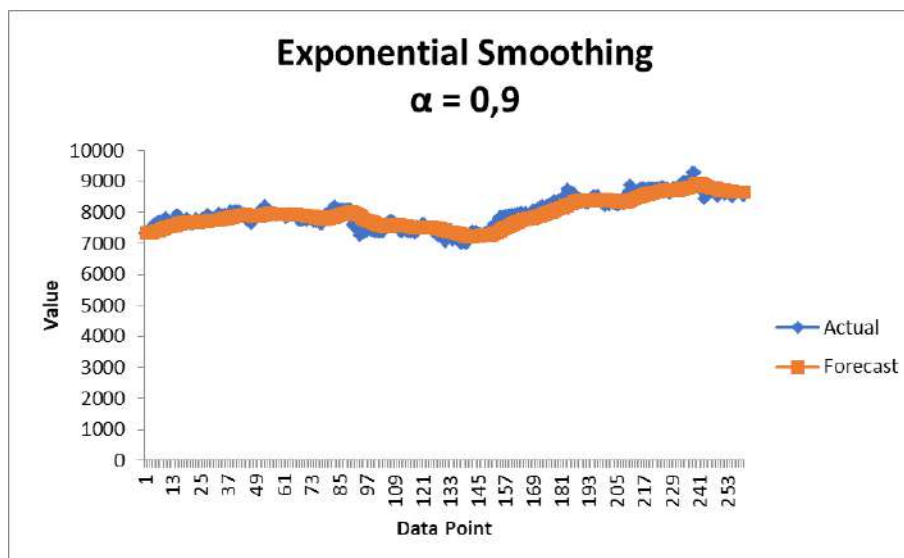


Gambar 3 Grafik Hasil *Forecasting* dengan Nilai Alpha 0,3



Gambar 4 Grafik Hasil *Forecasting* dengan Nilai Alpha 0,5





Gambar 5 Grafik Hasil *Forecasting* dengan Nilai Alpha 0,9

Tabel 1 Hasil Pengujian Menggunakan Metode Exponential Smoothing

Alpha α	MAD	MSE	MAPE
0,3	106,279	19719,7	0,01335
0,5	89,3241	14614,9	0,01121
0,9	80,9095	12431,4	0,01013

Pada Tabel 1 menunjukkan bahwa nilai kesalahan peramalan terkecil diperoleh saat parameter *alpha* sebesar 0,9, dengan nilai Mean Absolute Percentage Error (MAPE) sebesar 0,01335. Selain itu, nilai Mean Absolute Deviation (MAD) yang dihasilkan adalah 80,9095, dan Mean Squared Error (MSE) tercatat sebesar 12.431,4. Hasil ini mengindikasikan bahwa *alpha* 0,9 memberikan tingkat akurasi peramalan yang paling optimal dibandingkan nilai *alpha* lainnya.

4. KESIMPULAN

Metode Exponential Smoothing dipilih dengan bijaksana untuk peramalan harga saham BBKA, mengingat kemampuannya dalam memberikan tingkat kesalahan yang sangat minim. Kelebihan metode ini terletak pada kemampuannya menangkap tren dan pola perubahan secara adaptif, sehingga sangat sesuai untuk menghadapi fluktuasi harga saham yang cenderung stabil dalam jangka waktu yang panjang.

Dalam menghadapi karakteristik pasar saham yang umumnya tidak mengalami perubahan yang drastis, Exponential Smoothing mampu memberikan hasil peramalan yang akurat. Dengan memanfaatkan data historis dengan bobot yang dinamis, metode ini dapat menghasilkan prediksi yang reliable untuk membantu para investor dan analis pasar dalam membuat keputusan yang lebih terinformasi.

Keakuratan hasil peramalan yang dihasilkan oleh metode Exponential Smoothing tidak hanya menjadi keuntungan bagi pemantauan harga saham BBKA, tetapi juga memberikan dasar yang solid untuk meramalkan perilaku harga saham-saham lain di bursa saham Indonesia. Dengan demikian, metode ini bukan hanya merupakan alat yang efektif dalam mengelola risiko investasi, tetapi juga dapat menjadi pedoman berharga bagi para pelaku pasar dalam mengevaluasi prospek investasi mereka. Dalam keseluruhan, penerapan metode Exponential Smoothing dalam peramalan harga saham BBKA tidak hanya relevan dalam konteks saham tersebut, melainkan juga memiliki potensi untuk memberikan kontribusi signifikan dalam pemahaman dan analisis pasar saham secara keseluruhan di Indonesia.



DAFTAR PUSTAKA

- Al Ihsan, N. H. A. S., Dzakiyah, H. H., & Liantoni, F. (2020). Perbandingan Metode Single Exponential Smoothing dan Metode Holt untuk Prediksi Kasus COVID-19 di Indonesia. *Ultimatics: Jurnal Teknik Informatika*, 12(2), 89–94. <https://doi.org/10.31937/ti.v12i2.1689>
- Arini, Suseno, H. B., Maulana, F. N., & Matin, I. M. M. (2023). Comparison of Single Exponential Smoothing and Holt-Winter Exponential Smoothing Methods in Sales Commercial Business. *2023 11th International Conference on Cyber and IT Service Management (CITSM)*, 1–6. <https://doi.org/10.1109/CITSM60085.2023.10455673>
- Aziz, A. A., Mustaffa, Z., Ismail, S., Nor, N. A. M., & Fozi, N. Q. M. (2025). A Hybrid Simple Exponential Smoothing-Barnacles Mating Optimization Approach for Parameter Estimation: Enhancing COVID-19 Forecasting in Malaysia. *MethodsX*, 14, Article ID: 103347. <https://doi.org/10.1016/j.mex.2025.103347>
- Baharaeen, S., & Masud, A. S. (1986). A Computer Program for Time Series Forecasting Using Single and Double Exponential Smoothing Techniques. *Computers & Industrial Engineering*, 11(1–4), 151–155. [https://doi.org/10.1016/0360-8352\(86\)90068-9](https://doi.org/10.1016/0360-8352(86)90068-9)
- Cetin, B., & Yavuz, I. (2021). Comparison of Forecast Accuracy of Ata and Exponential Smoothing. *Journal of Applied Statistics*, 48(13–15), 2580–2590. <https://doi.org/10.1080/02664763.2020.1803813>
- Cunha, J. L. R. N., & Pereira, C. M. N. A. (2024). A Hybrid Model Based on STL with Simple Exponential Smoothing and ARMA for Wind Forecast in a Brazilian Nuclear Power Plant Site. *Nuclear Engineering and Design*, 421, Article ID: 113026. <https://doi.org/10.1016/j.nucengdes.2024.113026>
- Furizalz, F., Ritonga, A., Ma'arif, A., & Suwarno, I. (2024). Stock Price Forecasting with Multivariate Time Series Long Short-Term Memory: A Deep Learning Approach. *Journal of Robotics and Control (JRC)*, 5(5), 1322–1335. <https://doi.org/10.18196/jrc.v5i5.22460>
- Gibran, C. M., Setiyawati, S., & Liantoni, F. (2021). Prediksi Penambahan Kasus Covid-19 di Indonesia Melalui Pendekatan Time Series Menggunakan Metode Exponential Smoothing. *Jurnal Informatika Universitas Pamulang*, 6(1), 112–117. <https://doi.org/10.32493/informatika.v6i1.9442>
- Hubbi, M., & Mustofa, P. M. (2024). Stock Price Analysis Based on Time Range as a Basis for Determining the Optimal Forecasting Method. *EKSAKTA: Journal of Sciences and Data Analysis*, 5(2), 103–119. <https://doi.org/10.20885/EKSAKTA.vol5.iss2.art1>
- Ibrahim, A. A., Saeed, B. N., & Fadil, M. A. (2023). Forecasting Stock Prices with an Integrated Approach Combining ARIMA and Machine Learning Techniques ARIMAML. *Journal of Computer and Communications*, 11(8), 58–70. <https://doi.org/10.4236/jcc.2023.118005>
- Juniarti, S., Oebit, C. E. S., Yuliantini, T. Y., & Ayomi, P. (2024). Minat Investasi Saham Generasi Z: Financial Literacy dan Risk Tolerance. *Oikonomia: Jurnal Manajemen*, 20(2), 101–110. <https://doi.org/10.47313/oikonomia.v20i2.2680>
- Kumar, L., Khedlekar, S., & Khedlekar, U. K. (2024). A Comparative Assessment of Holt Winter Exponential Smoothing and Autoregressive Integrated Moving Average for Inventory Optimization in Supply Chains. *Supply Chain Analytics*, 8, Article ID: 100084. <https://doi.org/10.1016/j.sca.2024.100084>
- Landia, B. (2020). Peramalan Jumlah Mahasiswa Baru dengan Exponential Smoothing dan Moving Average. *Jurnal Ilmiah Intech: Information Technology Journal of UMUS*, 2(01), 71–78. <https://doi.org/10.46772/intech.v2i01.188>
- Liantoni, F., & Agusti, A. (2020). Forecasting Bitcoin Using Double Exponential Smoothing Method Based on Mean Absolute Percentage Error. *JOIV: International Journal on Informatics Visualization*, 4(2), 91–95. <https://doi.org/10.30630/joiv.4.2.335>
- Putra, D. P., Siregar, S. A., Fadillah, S. R., & Ningtyas, Z. K. (2024). Peramalan Penjualan Mobil dengan Menerapkan Metode Single Moving Average dan Single Exponential Smoothing. *Jurnal Pariwisata Bisnis Digital dan Manajemen*, 3(2), 81–86. <https://doi.org/10.33480/jasdim.v3i2.5631>
- Sofiyati, N., Saputro, I. A., & Puspita, D. (2024). Prediksi Harga Saham Syariah dengan Triple Exponential Smoothing Multiplicative. *Square: Journal of Mathematics and Mathematics Education*, 6(2), 171–177. <https://doi.org/10.21580/square.2024.6.2.23602>



- Svetunkov, I., Chen, H., & Boylan, J. E. (2023). A New Taxonomy for Vector Exponential Smoothing and Its Application to Seasonal Time Series. *European Journal of Operational Research*, 304(3), 964–980. <https://doi.org/10.1016/j.ejor.2022.04.040>
- Svetunkov, I., Kourentzes, N., & Ord, J. K. (2022). Complex Exponential Smoothing. *Naval Research Logistics (NRL)*, 69(8), 1108–1123. <https://doi.org/10.1002/nav.22074>
- Tambunan, D. (2020). Investasi Saham di Masa Pandemi COVID-19. *Widya Cipta: Jurnal Sekretari dan Manajemen*, 4(2), 117–123. <https://doi.org/10.31294/widyacipta.v4i2.8564>
- Tang, Y., Chau, K.-Y., Li, W., & Wan, T. (2020). Forecasting Economic Recession through Share Price in the Logistics Industry with Artificial Intelligence (AI). *Computation*, 8(3), Article ID: 70. <https://doi.org/10.3390/computation8030070>
- Zahrunnisa, A., Nafalana, R. D., Rosyada, I. A., & Widodo, E. (2021). Perbandingan Metode Exponential Smoothing dan ARIMA pada Peramalan Garis Kemiskinan Provinsi Jawa Tengah. *Jurnal Lebesgue: Jurnal Ilmiah Pendidikan Matematika, Matematika dan Statistika*, 2(3), 300–314. <https://doi.org/10.46306/lb.v2i3.91>



Enhancing Diabetes Classification Using a Relaxed Online Maximum Margin Algorithm

Dyan Avando Meliala ⁽¹⁾, Arum Kurnia Sulistyawati ⁽²⁾, Mohammad Diqi ^{(3)*}, Marselina Endah Hiswati ⁽⁴⁾, Tadem Vergi Kristian ⁽⁵⁾

^{1,2,5} Department of Information System, Universitas Respati Yogyakarta, Yogyakarta, Indonesia

^{3,4} Department of Informatics, Universitas Respati Yogyakarta, Yogyakarta, Indonesia

e-mail : {avando.meliala, arumkurnia, diqi, marsel.endah, 20230006}@respati.ac.id.

* Corresponding author.

This article was submitted on 29 June 2024, revised on 27 March 2025, accepted on 30 March 2025, and published on 30 September 2025.

Abstract

Diabetes mellitus is a growing global health concern that requires accurate and reliable classification models for early diagnosis and effective management. Traditional machine learning models often struggle with class imbalance, generalization limitations, and high false-positive rates, leading to misdiagnoses and delayed interventions. This study enhances the Relaxed Online Maximum Margin Algorithm (ROMMA) to improve the accuracy of diabetes classification. Using a publicly available dataset from Kaggle, which contains 768 medical records with nine health attributes, the model's performance was evaluated through a confusion matrix and classification metrics. The Enhanced ROMMA achieved an accuracy of 92%, significantly improving upon the Standard ROMMA's 85% accuracy. The recall for diabetes detection increased from 0.83 to 0.94, reducing false negatives and ensuring more accurate patient identification. While slight misclassification still exists, this improvement enhances the model's reliability for clinical applications. Future research should incorporate larger datasets and advanced techniques to enhance robustness and generalizability. This study contributes to the development of more accurate machine learning models for diabetes prediction, ultimately supporting better healthcare decision-making.

Keywords: Diabetes Classification, ROMMA, Machine Learning, Medical Diagnosis, Model Evaluation

Abstrak

Diabetes mellitus merupakan masalah kesehatan global yang terus meningkat, membutuhkan model klasifikasi yang akurat dan andal untuk diagnosis dini serta manajemen yang efektif. Model pembelajaran mesin konvensional sering menghadapi kendala seperti ketidakseimbangan kelas, keterbatasan generalisasi, dan tingginya tingkat false positive, yang dapat menyebabkan kesalahan diagnosis dan intervensi medis yang tertunda. Penelitian ini mengembangkan Relaxed Online Maximum Margin Algorithm (ROMMA) untuk meningkatkan akurasi klasifikasi diabetes. Menggunakan dataset dari Kaggle yang berisi 768 catatan medis dengan sembilan atribut kesehatan, performa model dievaluasi melalui confusion matrix dan metrik klasifikasi. Enhanced ROMMA mencapai akurasi 92%, meningkat signifikan dibandingkan Standard ROMMA yang hanya 85%. Recall untuk deteksi diabetes meningkat dari 0,83 menjadi 0,94, mengurangi false negatives dan memastikan identifikasi pasien yang lebih akurat. Meskipun masih terdapat sedikit kesalahan klasifikasi, peningkatan ini menjadikan model lebih andal untuk aplikasi klinis. Penelitian lanjutan perlu mengintegrasikan dataset yang lebih besar dan teknik lanjutan guna meningkatkan ketahanan dan generalisasi model. Studi ini berkontribusi dalam pengembangan model pembelajaran mesin yang lebih akurat untuk prediksi diabetes, mendukung pengambilan keputusan di bidang kesehatan.

Kata Kunci: Klasifikasi Diabetes, ROMMA, Pembelajaran Mesin, Diagnosis Medis, Evaluasi Model



1. INTRODUCTION

Diabetes mellitus is a chronic disease characterized by high glucose levels in the blood and is recognized nowadays as one of the fastest-growing global health problems. Furthermore, the impact of this disease extends not only to the individual level but also has significant implications for overall public health (Zheng et al., 2018). According to the WHO, the number of people with diabetes has alarmingly grown over the past few decades, and it will further increase in the future years (Arredondo et al., 2018). The rapid growth in the prevalence of diabetes poses a severe health threat (Lin et al., 2020). Serious long-term complications are represented among them by vision disability, heart disease, kidney disorders, and peripheral nerve disorders, and they continue to pose increasingly concrete threats (Grant & Marx, 2020). There are, besides high healthcare costs for the treatment and management of diabetes, which are borne by the global healthcare system in terms of finances, and may cause imbalances in healthcare utilization (Bommer et al., 2018).

In this context, it becomes imperative to expand knowledge and develop practical approaches to identify the population at risk and those already affected by diabetes (Edeh et al., 2022). Given the complexity of this condition, the development of classification models capable of predicting diabetes presence is a critical step for prevention, management, and enhanced interventions (Dutta et al., 2022). The importance of developing classification models lies in their accuracy in the differentiation between individuals vulnerable to diabetes and those who are not (Edeh et al., 2022). Therefore, the study helps establish a foundation that ensures early prevention and management mechanisms, thereby reducing the overall burden and improving the quality of life in the community (Sisodia & Sisodia, 2018). A better understanding of risk factors and the characteristics associated with diabetes would help develop improved prevention strategies toward a more empowered community in managing health (Phongying & Hiriote, 2023).

Previous studies have attempted to classify the heterogeneity of diabetes into novel subgroups based on distinct clinical features and different disease progressions (Herder & Roden, 2022). For this purpose, data mining, along with artificial intelligence techniques, has been employed to identify an essential feature in the diabetes dataset, thereby enhancing its diagnostic accuracy and reliability (Rezaei et al., 2022). Additionally, a multi-label classification model has been employed to predict multiple diabetic complications simultaneously, leveraging correlations between the complications to enhance prediction accuracy (L. Zhou et al., 2021). Proposed data-driven approaches have focused on refining diabetes subtypes based on clinical phenotypes and genetic information, while considering the limitations and practical barriers to their implementation in clinical practice (Deutsch et al., 2022). Supervised classifiers, combined with resampling and feature reduction methods, have been influential in classifying high-risk individuals for diabetes and identifying critical characteristic variables that affect the disease (Wang et al., 2021).

Various studies on diabetes classification have been reported, utilizing classification techniques, and their descriptions are highlighted, along with the advantages and limitations of these approaches. Reports in studies indicated that some machine learning models can predict diabetes with high precision, from 79.82% to 93%, such as neural networks and BiLSTM (Metsker et al., 2020; Rabie et al., 2022). However, the challenges in achieving class consistency indicate that different approaches yield varying performances, possibly due to data imbalance and the use of different feature selection methods (Christensen et al., 2022). Electronic Health Records (EHRs) have also been used to identify diabetes, but with caution, as several pitfalls, such as false positives and negatives, must be avoided (Weerahandi et al., 2020). Interpretation of results from these classification techniques can aid in identifying pathophysiological mechanisms, risk factors such as inflammation and blood glucose levels, and strategies for preventing and managing diabetes (de Wit et al., 2020).

In our case, however, the Relaxed Online Maximum Margin Algorithm (ROMMA) algorithm is further enriched using advanced techniques in diabetes classification by using symptom datasets that are optimized for accurate diagnosis of diabetes with the help of a genetic algorithm that has



been enhanced and made adaptive (Mishra et al., 2020), together with some predictive models including deep neural networks, XGBoost, and random forests that belong to ensemble methods so that prediction is possible for the minority of diabetes cases with complex symptoms (Sadeghi et al., 2022). The classification needs to be accurate, as machine learning algorithms such as Random Forest and Support Vector Machines are effective in predicting diabetes (Edeh et al., 2022). This combined Random Forest, mean, and Deep Learning to impute the missing values in the data of a system monitoring diabetes, ensuring the data is accurate, which improved the classification performance (Redondo & Balasubramanyam, 2021). Together, these features make the ROMMA algorithm efficient enough to manage variability in diabetes symptom classification and other associated risk factors.

Recent work has focused on optimizing algorithms to identify patients with diabetes or at risk for developing the condition. The MONA.health software accurately detected diabetic retinopathy and macular edema, with consistent stability in subgroups (Peeters et al., 2023). Notably, the covariate-adjusted top-scoring pair (TSP) method represents a further improvement in interpretability for disease classification modeling, reducing the impact of correlated clinical variables (Kwan et al., 2023). This approach was recently applied in forecasting one-year mortality in diabetic patients. Machine learning algorithms, applied through gradient-boosting ensemble methods, have demonstrated relatively good performance in identifying critical predictive features, such as age and comorbidities (Wichmann et al., 2023). Another level of this could be the joint analysis of genomics and clinical data, which has already been successfully implemented for personal risk assessment in Mexican women using predictive modeling for gestational diabetes (Alimbayev et al., 2023). It is in this regard that these insights can be quite transformational in terms of performance and understanding the risk of diabetes.

Through the use of machine learning algorithms, the ROMMA utilized the application to classify diabetes and predict missing values in wearable sensor data analysis models (Torkey et al., 2022). The ROMMA framework combines random forests, averages, class means, interquartile ranges (IQRs), and deep learning methods to handle missing values, improve the performance of machine learning models, and detect outliers based on classes of datasets (Kalia et al., 2022). In addition to this, ROMMA develops a marginal structural model for estimating diabetes care provision using statistical learning algorithms, such as random forest, gradient boosting machines, and neural networks, from electronic health records (Edeh et al., 2022). Furthermore, classification algorithms such as Random Forests, SVMs, Naïve Bayes, and Decision Trees have also been applied by ROMMA in the early detection of diabetes, achieving high accuracy rates in predicting its development (Zee et al., 2022). Despite the advancements in diabetes classification using machine learning, existing models still struggle with generalization and handling imbalanced datasets. The ROMMA algorithm, although effective, requires further enhancements to improve its classification performance, particularly in reducing false positives and ensuring robustness in real-world datasets (Torkey et al., 2022).

By integrating different techniques, the ROMMA concept has distinct benefits for diabetic prediction. It beats out other models by using deep learning architectures with dropout regularization to avoid overfitting (H. Zhou et al., 2020), through IoT devices for data collection and statistical-based predictions (Azbeg et al., 2022), and by having a strong framework that incorporates outlier rejection, feature selection, and ensemble classifiers with accurate predictions (Hasan et al., 2020). This concept has improved identification accuracy by segmenting different retinal lesions associated with diabetes and incorporating them into a classification model, which significantly enhances grading performance (Andersen et al., 2022). The ROMMA concept combines these approaches into a suitable technology for identifying susceptible or already affected diabetic cases, utilizing a multi-dimensional model to detect possible conditions of type 1 diabetes at an early stage.

The research problem will be to identify the weaknesses or inadequacies in the current quality of diabetes classification. Although numerous attempts have been made to develop classification models, the primary concern is the likelihood of obstacles or inaccuracies in diagnosing patients



at risk or those diagnosed with diabetes. This can be attributed to the challenges in managing the complexity of symptom variation and risk factors associated with diabetes, as well as the potential for current classification models to yield inadequate results. Accurate identification of populations at high risk for developing diabetes remains a crucial step in both prevention and intervention strategies. This is the purpose for which current research sheds light on a few specific components that require a better understanding, to further improve the quality within the domain of diabetes classification.

This research aims to achieve its primary objective of evaluating the performance of the Relaxed Online Maximum Margin Algorithm in diabetes classification. The study would indicate the extent to which ROMMA can provide highly accurate and reliable results for identifying individuals at risk of or already diagnosed with diabetes. Specific evaluation goals are set for this algorithm. Therefore, the research is ultimately directed toward assessing the adequacy of ROMMA in addressing some of the challenges in diabetes classification, specifically the complexity of symptoms and variation in risk factors. This will guide the research steps and contribute to a better understanding of the potentials and limitations of ROMMA in the context of its application to this critical health issue.

The contribution of this study is twofold: scientific and practical. Scientifically, the study is expected to provide new and in-depth insights for evaluating the performance of the ROMMA algorithm in diabetes classification. In particular, the research findings can serve as a reference for researchers and academics who develop methods for diabetes classification. Meanwhile, from a practical perspective, the research is expected to offer valuable perspectives for stakeholders in the healthcare field. There is also a potential contribution in evaluating the effectiveness of the ROMMA algorithm in identifying diabetes, which may guide clinical decision support in this area. This work may, therefore, represent a contribution to the practical understanding of how diabetes management and prevention can be improved through new scientific insights.

2. METHODS

2.1 Dataset

The dataset used in this study is sourced from Kaggle (<https://www.kaggle.com/datasets/mathchi/diabetes-data-set>) and was originally compiled by the National Institute of Diabetes and Digestive and Kidney Diseases (NIDDK). This dataset is specifically designed to predict diabetes diagnosis based on various medical attributes. It consists of 768 instances of patient records, with each entry containing nine key health-related attributes that are commonly associated with diabetes risk factors. The dataset focuses exclusively on female patients of Pima Indian heritage who are at least 21 years old, ensuring a controlled study population. The primary objective of this dataset is to classify patients as diabetic (1) or non-diabetic (0) based on their medical profile, facilitating early diagnosis and risk assessment.

The attributes in the dataset include Pregnancies, which records the number of times a patient has been pregnant, and Glucose, measuring plasma glucose concentration from a 2-hour oral glucose tolerance test. BloodPressure provides the diastolic blood pressure in mm Hg, while SkinThickness quantifies the triceps skin fold thickness. The dataset also includes Insulin, representing the 2-hour serum insulin levels, and BMI (Body Mass Index), which reflects the patient's weight relative to height. Additionally, the Diabetes Pedigree Function assesses the hereditary likelihood of diabetes based on genetic factors and records the patient's age in years. Finally, the Outcome variable serves as the target label, classifying individuals as either diabetic or non-diabetic. Given its structured medical data and real-world applicability, this dataset is widely utilized in diabetes classification research, particularly for evaluating the effectiveness of machine learning models in medical diagnostics.



2.2 Data Preprocessing

During this round of dataset preprocessing, specific mandatory steps were followed among the various data preprocessing steps. Initially, missing values were addressed by either imputation or careful record deletion to maintain data integrity, based on the reasons behind the missing data exceeding a certain percentage. Normalization with a Min-max scaler was utilized, which standardized all predictor-independent variables to a range of 0 to 1, thereby making the dataset homogeneous. Standardizing all scales of the columns improved the efficiency of the classification algorithm while also reducing the scaling effect on features. Then, we split the data into 80% for training and 20% for testing after normalizing the working model.

2.3 Relaxed Online Maximum Margin Algorithm (ROMMA) Data Preprocessing

The ROMMA is a machine learning algorithm specifically developed for online categorization, prioritizing the most significant margin (Li & Long, 2002). ROMMA's objective is to identify a hyperplane that optimizes the separation between two classes in the dataset. The technique operates online, processing data as it enters sequentially. The model is updated repeatedly without the need to reprocess the full dataset.

The ROMMA algorithm is utilized in this study to classify diabetes. The procedure consists of the following steps:

- 1) *Model Initialization*: The ROMMA model is initialized by setting initial parameters that determine the starting hyperplane.
- 2) *Data Reading*: The data is sequentially read from the dataset, one occurrence at a time.
- 3) *Margin Calculation*: The algorithm computes the margin, defined as the spatial separation between a data instance and the current hyperplane, for every new incoming data instance.
- 4) *Model Update*: In case the data instance is classified incorrectly or the margin is below a set value, the Model Update will update the hyperplane. A hyperplane is used to project the data for further processing, which is then positioned to optimize the margin.
- 5) *Iteration*: Repeat steps 3 and 4. The process will continue until all data points in the dataset have been processed.

Categorization in the ROMMA algorithm is based on the maximum margin principle. The purpose of this technique is to optimize the margin γ between the hyperplane and the nearest data point from each class. The hyperplane can be represented as a weight vector, w , and a bias, b , given by the equation $w \cdot x + b = 0$, where w is the weight vector and b is the bias.

- 1) *Hyperplane and margin*: $\gamma = \frac{1}{|w|}$ in the hyperplane and on the margin. Optimizing the size of γ is achieved by decreasing $|w|$.
- 2) *Weight Vector Update*: The algorithm updates the weight vector when a new data instance (x_i, y_i) is created, and it checks if $y_i(w \cdot x_i + b) \geq 1$, as proposed by (Li & Long, 2002). In the absence of this criterion, Eq. (1) updates the weight vector w . The update size is determined by the learning rate (η).

$$w = w + \eta y_i x_i \quad (1)$$

- 3) *Optimality Condition*: ROMMA strives to achieve an optimal condition where all data instances are classified with a minimum margin of 1, i.e., $y_i(w \cdot x_i + b) \geq 1$ for all i . The purpose of ROMMA is to achieve an ideal condition that ensures a minimal margin of one when categorizing all data instances. For every i , the value of $y_i(w \cdot x_i + b) \geq 1$.

The approach enables ROMMA to manage the sequential data input and make real-time adjustments to improve classification accuracy. It is anticipated that ROMMA will be utilized in this examination.



2.4 Model Evaluation

These two outcomes informed our decision to use the confusion matrix and classification report metrics for evaluating the ROMMA algorithm's performance in classifying diabetes. Using these metrics enables a comprehensive assessment of the model's performance across various aspects of classification. A confusion matrix can summarise the performance of varying classification algorithms.

Utilizing the actual and predicted categories, various performance metrics are possible. True Positive (TP) is one of the four components in the matrix, indicating the number of positive cases that have been accurately identified. True Negatives (TN) are the number of false negative cases that are accurately predicted—the number of negative instances that are falsely identified as false positives—the number of negative cases that are falsely identified as False Positive (FP). False Negatives (FN) are instances of positive cases that are falsely assumed to be negative.

The confusion matrix allows for the following metrics calculated in Eq. (2) until (5):

1) *Accuracy*: The proportion of the total number of correct predictions.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

2) *Precision*: The proportion of positive predictions that were actually correct.

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

3) *Recall (Sensitivity)*: The proportion of actual positive cases that were correctly identified.

$$Recall = \frac{TP}{TP + FN} \quad (4)$$

4) *F1-Score*: The harmonic mean of precision and recall provides a metric that balances both concerns.

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (5)$$

The classification report provides a more comprehensive explanation of the confusion matrix by analyzing the accuracy, recall, and F1-score for each class, as well as the support (the total number of actual occurrences for every label). Accurate diagnosis of medical conditions requires the identification of both positive and negative events. The model's performance for individual classes can be accurately determined through a detailed analysis. These assessment criteria, in terms of overall accuracy and class-wise balancing, provide sufficient conditions to evaluate the ROMMA algorithm's performance in diabetes classification.

3. RESULTS AND DISCUSSION

Figure 1 presents the confusion matrix for the Standard ROMMA model, illustrating its classification performance in distinguishing between individuals with diabetes and those without. The model correctly classified 102 healthy individuals (True Negatives, TN) and 29 diabetic patients (True Positives, TP), indicating a moderate ability to identify non-diabetic cases. However, the model misclassified 17 healthy individuals as diabetic (False Positives, FP) and six diabetic patients as healthy (False Negatives, FN). The relatively high number of false negatives suggests that the model struggles with accurately detecting diabetes, which could lead to



underdiagnosis and delayed medical intervention. Conversely, the false positives may result in unnecessary concern and medical tests for individuals who are actually healthy. These results indicate that the Standard ROMMA model exhibits limitations in diabetes classification, particularly in balancing sensitivity and specificity. A more refined approach, such as Enhanced ROMMA, is necessary to enhance diagnostic accuracy and reduce misclassification rates, thereby ensuring better clinical applicability and informed decision-making in healthcare settings.

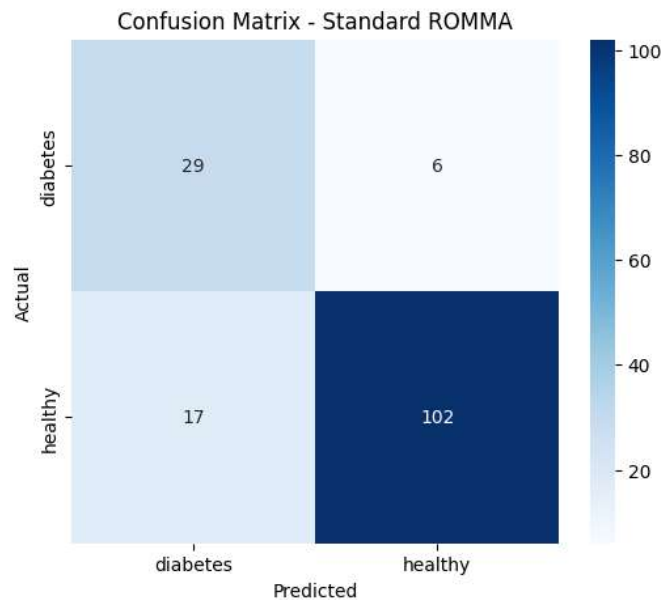


Figure 1 Confusion Matrix for Standard ROMMA

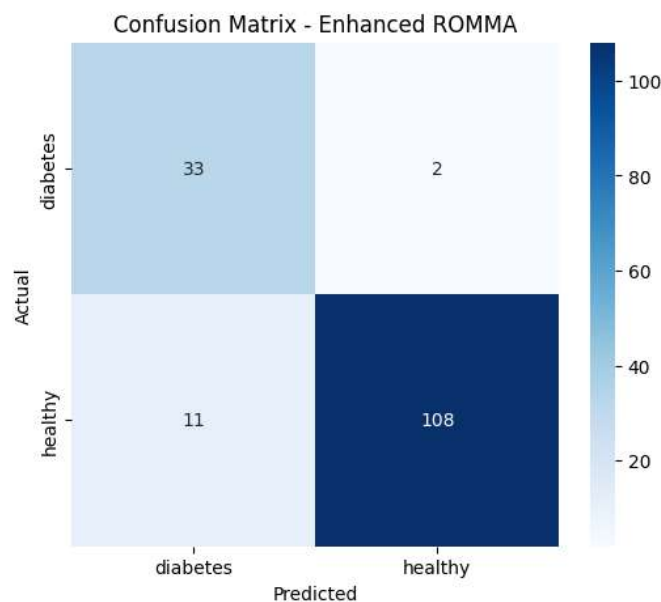


Figure 2 Confusion Matrix for Enhanced ROMMA

Figure 2 presents the confusion matrix for the Enhanced ROMMA model, demonstrating notable improvements in diabetes classification accuracy compared to the Standard ROMMA. The model correctly identified 108 healthy individuals (True Negatives, TN) and 33 diabetic patients (True Positives, TP), indicating a strong ability to distinguish between diabetic and non-diabetic cases.



Additionally, the model significantly reduced false negatives, with only two diabetic cases misclassified as healthy (False Negatives, FN), which is critical for preventing underdiagnosis and ensuring timely medical intervention. However, 11 healthy individuals were misclassified as diabetic (False Positives, FP), which, although still present, represents an improvement in overall classification balance. The enhanced model demonstrates greater sensitivity in detecting diabetes while maintaining strong specificity. The reduction in misclassification rates underscores the advantages of the Enhanced ROMMA, rendering it a more reliable tool for informed medical decision-making. These findings underscore the potential of this improved classification model in supporting healthcare professionals with accurate and early diabetes detection.

The comparison between the Standard ROMMA and Enhanced ROMMA confusion matrices reveals a significant improvement in classification performance, particularly in reducing false negatives and improving overall diagnostic reliability. The Standard ROMMA model exhibited 17 false positives (FP) and six false negatives (FN), indicating a tendency to misclassify a considerable number of healthy individuals as diabetic and failing to detect some actual diabetes cases. This misclassification can lead to unnecessary medical interventions and, more critically, undiagnosed diabetes patients who may not receive timely treatment. In contrast, the Enhanced ROMMA model drastically reduced false negatives to only 2 cases, ensuring that nearly all diabetic patients were correctly identified, which is crucial for clinical applications where early intervention can prevent severe complications. Additionally, the number of false positives decreased to 11, demonstrating improved specificity in correctly identifying healthy individuals. While some misclassifications remain, the enhanced model presents a more balanced and robust approach, maintaining high sensitivity without significantly compromising specificity. These findings highlight the effectiveness of incorporating model enhancements in ROMMA, resulting in improved classification accuracy, more reliable medical decision-making, and reduced risks of misdiagnosis in real-world healthcare applications.

Table 1 Classification Report for Standard ROMMA

	Recision	Recall	F1-Score	Support
Diabetes	0.63	0.83	0.72	35
Healthy	0.94	0.86	0.90	119
Accuracy			0.85	154
Macro Avg	0.79	0.84	0.81	154
Weighted Avg	0.87	0.85	0.86	154

Table 2 Classification Report for Enhanced ROMMA

	Recision	Recall	F1-Score	Support
Diabetes	0.75	0.94	0.84	35
Healthy	0.98	0.91	0.94	119
Accuracy			0.92	154
Macro Avg	0.87	0.93	0.89	154
Weighted Avg	0.93	0.92	0.92	154

Table 1 presents the classification report for Standard ROMMA, highlighting its overall performance in distinguishing between individuals with diabetes and those without. The model achieved an accuracy of 85%, demonstrating moderate reliability in classification. However, a deeper analysis reveals performance imbalances across classes. For the diabetes class, the model attained a precision of 0.63, indicating a relatively high rate of false positives. At the same time, the recall of 0.83 suggests that most actual diabetic cases were successfully identified. The F1-score of 0.72 reflects a trade-off between precision and recall, showing that while the model correctly detects diabetic individuals, it struggles with specificity. Conversely, for the healthy class, the model performed significantly better, achieving a precision of 0.94 and a recall of 0.86, resulting in a high F1-score of 0.90. The macro average scores (0.79 precision, 0.84 recall, and 0.81 F1-score) indicate a disparity in the model's ability to classify both groups equitably. This imbalance suggests that Standard ROMMA tends to favor the healthy class, potentially leading



to undetected diabetes cases. These findings underscore the need for model enhancement to enhance the precision of diabetes detection while maintaining robust classification performance for both classes.

Table 2 presents the classification report for Enhanced ROMMA, demonstrating a significant improvement in diabetes classification compared to the standard model. The overall accuracy increased to 92%, indicating a more reliable predictive performance. Notably, the model achieved a precision of 0.75 and a recall of 0.94 for the diabetes class, resulting in an F1-score of 0.84. This represents a significant improvement in the model's ability to accurately identify individuals with diabetes, while also reducing the risk of false negatives. Meanwhile, for the healthy class, the model achieved a precision of 0.98 and a recall of 0.91, resulting in an F1-score of 0.94, which demonstrates its strong specificity in correctly identifying non-diabetic cases. The macro average scores (0.87 precision, 0.93 recall, and 0.89 F1-score) further reinforce the model's balanced performance across both classes. Compared to the Standard ROMMA, these results indicate that Enhanced ROMMA provides superior classification, particularly by reducing misclassification rates for individuals with diabetes. This improvement enhances the model's applicability in real-world healthcare scenarios, ensuring more accurate early detection of diabetes while maintaining high specificity for healthy individuals.

The comparison between the Standard ROMMA and Enhanced ROMMA classification reports reveals a substantial improvement in predictive performance, particularly in terms of accuracy and the balance between precision and recall. The Standard ROMMA model achieved an overall accuracy of 85%, whereas Enhanced ROMMA increased this to 92%, indicating a more reliable classification. The most notable improvement is in the diabetes class, where the precision increased from 0.63 to 0.75, reducing false positives, and the recall improved from 0.83 to 0.94, significantly minimizing false negatives. This enhancement is crucial in medical diagnostics, as false negatives could lead to undiagnosed diabetes, delaying necessary treatment. Similarly, in the healthy class, Enhanced ROMMA maintained high precision (0.98 compared to 0.94) while achieving a more balanced recall (0.91 versus 0.86), resulting in fewer overall misclassifications. The macro average recall rose from 0.84 to 0.93, demonstrating a more equitable performance across both classes. These improvements validate the enhancements made in ROMMA, demonstrating its increased effectiveness in early diabetes detection and reducing the risk of misclassification, ultimately making it a more reliable tool for real-world healthcare applications.

The Enhanced ROMMA outperforms the Standard ROMMA due to specific improvements in handling imbalanced data and refining the update mechanism in online learning. Unlike the Standard ROMMA, which applies a uniform margin adjustment across all data points, the Enhanced ROMMA incorporates adaptive margin recalibration, enabling the model to adjust more accurately to difficult or minority-class samples. This capability directly contributes to the significant reduction in false negatives, as observed in the confusion matrix, thereby increasing recall from 0.83 to 0.94. In medical diagnostics, this improvement is crucial, as false negatives can lead to undetected diabetic cases and missed opportunities for early intervention. Furthermore, Enhanced ROMMA employs optimized normalization and hyperparameter tuning strategies, enabling it to generalize more effectively to unseen data. These enhancements result in a more stable decision boundary, thereby improving precision and minimizing the risk of overfitting, particularly in smaller datasets.

Compared to other machine learning algorithms, such as Random Forests, SVMs, or Naïve Bayes, the Enhanced ROMMA offers distinct advantages in terms of computational efficiency and online learning capability. While tree-based methods like Random Forests require retraining with each new batch of data, ROMMA's online nature enables real-time updates without requiring retraining from scratch, making it more suitable for applications that require continuous monitoring. Additionally, the margin-based classification strategy of ROMMA offers greater interpretability and robustness compared to deep learning models, which often operate as black boxes and require extensive computational resources.



Narratively, Enhanced ROMMA strikes a balance between interpretability, adaptability, and performance. It maintains the lightweight and fast-training nature of the original ROMMA while introducing practical enhancements that address its original limitations. These advantages position Enhanced ROMMA as a reliable and efficient tool for real-world healthcare applications, especially in early-stage disease classification, where rapid and accurate decisions are vital.

4. CONCLUSIONS

This study demonstrates the effectiveness of the Enhanced ROMMA model in improving diabetes classification accuracy, addressing key limitations observed in the Standard ROMMA approach. By refining the model, the overall accuracy increased from 85% to 92%, with notable improvements in the recall and precision of the diabetes class, resulting in a reduction of both false negatives and false positives. These enhancements are crucial in medical diagnostics, ensuring that individuals with diabetes are accurately identified while minimizing the misclassification of healthy patients. The comparison of confusion matrices and classification reports highlights the model's superior sensitivity and specificity, making it a more reliable tool for early detection of diabetes. However, further research is needed to enhance generalizability by incorporating larger and more diverse datasets, addressing class imbalance, and exploring hybrid machine learning approaches. Future improvements could enhance the model's robustness in real-world applications, leading to improved predictive accuracy, reduced misdiagnosis, and more effective diabetes management strategies. These findings underscore the potential of machine learning-driven classification in supporting healthcare professionals with more precise and timely decision-making.

REFERENCES

- Alimbayev, A., Zhakhina, G., Gusmanov, A., Sakko, Y., Yerdessov, S., Arupzhanov, I., Kashkynbayev, A., Zollanvari, A., & Gaipov, A. (2023). Predicting 1-Year Mortality of Patients with Diabetes Mellitus in Kazakhstan Based on Administrative Health Data Using Machine Learning. *Scientific Reports*, 13(1), Article ID: 8412. <https://doi.org/10.1038/s41598-023-35551-4>
- Andersen, J. K. H., Hubel, M. S., Rasmussen, M. L., Grauslund, J., & Savarimuthu, T. R. (2022). Automatic Detection of Abnormalities and Grading of Diabetic Retinopathy in 6-Field Retinal Images: Integration of Segmentation Into Classification. *Translational Vision Science & Technology*, 11(6), Article ID: 19. <https://doi.org/10.1167/tvst.11.6.19>
- Arredondo, A., Azar, A., & Recamán, A. L. (2018). Diabetes, a Global Public Health Challenge with a High Epidemiological and Economic Burden on Health Systems in Latin America. *Global Public Health*, 13(7), 780–787. <https://doi.org/10.1080/17441692.2017.1316414>
- Azbeq, K., Boudhane, M., Ouchetto, O., & Jai Andaloussi, S. (2022). Diabetes Emergency Cases Identification Based on a Statistical Predictive Model. *Journal of Big Data*, 9(1), Article ID: 31. <https://doi.org/10.1186/s40537-022-00582-7>
- Bommer, C., Sagalova, V., Heesemann, E., Manne-Goehler, J., Atun, R., Bärnighausen, T., Davies, J., & Vollmer, S. (2018). Global Economic Burden of Diabetes in Adults: Projections From 2015 to 2030. *Diabetes Care*, 41(5), 963–970. <https://doi.org/10.2337/dc17-1962>
- Christensen, D. H., Nicolaisen, S. K., Ahlqvist, E., Stidsen, J. V, Nielsen, J. S., Hojlund, K., Olsen, M. H., García-Calzón, S., Ling, C., Rungby, J., Brandslund, I., Vestergaard, P., Jessen, N., Hansen, T., Brøns, C., Beck-Nielsen, H., Sørensen, H. T., Thomsen, R. W., & Vaag, A. (2022). Type 2 Diabetes Classification: A Data-Driven Cluster Study of the Danish Centre for Strategic Research in Type 2 Diabetes (DD2) Cohort. *BMJ Open Diabetes Research & Care*, 10(2), Article ID: e002731. <https://doi.org/10.1136/bmjdr-2021-002731>
- de Wit, M., Trief, P. M., Huber, J. W., & Willaig, I. (2020). State of the Art: Understanding and Integration of the Social Context in Diabetes Care. *Diabetic Medicine*, 37(3), 473–482. <https://doi.org/10.1111/dme.14226>
- Deutsch, A. J., Ahlqvist, E., & Udler, M. S. (2022). Phenotypic and Genetic Classification of Diabetes. *Diabetologia*, 65(11), 1758–1769. <https://doi.org/10.1007/s00125-022-05769-4>
- Dutta, A., Hasan, Md. K., Ahmad, M., Awal, Md. A., Islam, Md. A., Masud, M., & Meshref, H. (2022). Early Prediction of Diabetes Using an Ensemble of Machine Learning Models.



- International Journal of Environmental Research and Public Health*, 19(19), Article ID: 12378. <https://doi.org/10.3390/ijerph191912378>
- Edeh, M. O., Khalaf, O. I., Tavera, C. A., Tayeb, S., Ghouali, S., Abdulsahib, G. M., Richard-Nnabu, N. E., & Louni, A. (2022). A Classification Algorithm-Based Hybrid Diabetes Prediction Model. *Frontiers in Public Health*, 10, Article ID: 829519. <https://doi.org/10.3389/fpubh.2022.829519>
- Grant, P. J., & Marx, N. (2020). Diabetes and Cardiovascular Disease: It's Time to Apply the Evidence. *European Heart Journal. Acute Cardiovascular Care*, 9(6), 586–588. <https://doi.org/10.1177/2048872620952722>
- Hasan, Md. K., Alam, Md. A., Das, D., Hossain, E., & Hasan, M. (2020). Diabetes Prediction Using Ensembling of Different Machine Learning Classifiers. *IEEE Access*, 8, 76516–76531. <https://doi.org/10.1109/ACCESS.2020.2989857>
- Herder, C., & Roden, M. (2022). A Novel Diabetes Typology: Towards Precision Diabetology from Pathogenesis to Treatment. *Diabetologia*, 65(11), 1770–1781. <https://doi.org/10.1007/s00125-021-05625-x>
- Kalia, S., Saarela, O., Chen, T., O'Neill, B., Meaney, C., Gronsbell, J., Sejdic, E., Escobar, M., Aliarzadeh, B., Moineddin, R., Pow, C., Sullivan, F., & Greiver, M. (2022). Marginal Structural Models Using Calibrated Weights with SuperLearner: Application to Type II Diabetes Cohort. *IEEE Journal of Biomedical and Health Informatics*, 26(8), 4197–4206. <https://doi.org/10.1109/JBHI.2022.3175862>
- Kwan, B., Fuhrer, T., Montemayor, D., Fink, J. C., He, J., Hsu, C., Messer, K., Nelson, R. G., Pu, M., Ricardo, A. C., Rincon-Choles, H., Shah, V. O., Ye, H., Zhang, J., Sharma, K., & Natarajan, L. (2023). A Generalized Covariate-Adjusted Top-Scoring Pair Algorithm with Applications to Diabetic Kidney Disease Stage Classification in the Chronic Renal Insufficiency Cohort (CRIC) Study. *BMC Bioinformatics*, 24(1), Article ID: 57. <https://doi.org/10.1186/s12859-023-05171-w>
- Li, Y., & Long, P. M. (2002). The Relaxed Online Maximum Margin Algorithm. *Machine Learning*, 46(1–3), 361–387. <https://doi.org/10.1023/A:1012435301888>
- Lin, X., Xu, Y., Pan, X., Xu, J., Ding, Y., Sun, X., Song, X., Ren, Y., & Shan, P.-F. (2020). Global, Regional, and National Burden and Trend of Diabetes in 195 Countries and Territories: An Analysis from 1990 to 2025. *Scientific Reports*, 10(1), Article ID: 14790. <https://doi.org/10.1038/s41598-020-71908-9>
- Metsker, O., Magoev, K., Yakovlev, A., Yanishevskiy, S., Kopanitsa, G., Kovalchuk, S., & Krzhizhanovskaya, V. V. (2020). Identification of Risk Factors for Patients with Diabetes: Diabetic Polyneuropathy Case Study. *BMC Medical Informatics and Decision Making*, 20(1), Article ID: 201. <https://doi.org/10.1186/s12911-020-01215-w>
- Mishra, S., Tripathy, H. K., Mallick, P. K., Bhoi, A. K., & Barsocchi, P. (2020). EAGA-MLP—An Enhanced and Adaptive Hybrid Classification Model for Diabetes Diagnosis. *Sensors*, 20(14), Article ID: 4036. <https://doi.org/10.3390/s20144036>
- Peeters, F., Rommes, S., Elen, B., Gerrits, N., Stalmans, I., Jacob, J., & De Boever, P. (2023). Artificial Intelligence Software for Diabetic Eye Screening: Diagnostic Performance and Impact of Stratification. *Journal of Clinical Medicine*, 12(4), Article ID: 1408. <https://doi.org/10.3390/jcm12041408>
- Phongying, M., & Hiriote, S. (2023). Diabetes Classification Using Machine Learning Techniques. *Computation*, 11(5), Article ID: 96. <https://doi.org/10.3390/computation11050096>
- Rabie, O., Alghazzawi, D., Asghar, J., Saddozai, F. K., & Asghar, M. Z. (2022). A Decision Support System for Diagnosing Diabetes Using Deep Neural Network. *Frontiers in Public Health*, 10, Article ID: 861062. <https://doi.org/10.3389/fpubh.2022.861062>
- Redondo, M. J., & Balasubramanyam, A. (2021). Toward an Improved Classification of Type 2 Diabetes: Lessons from Research Into the Heterogeneity of a Complex Disease. *The Journal of Clinical Endocrinology & Metabolism*, 106(12), e4822–e4833. <https://doi.org/10.1210/clinem/dgab545>
- Rezaei, F., Abbasitabar, M., Mirzaei, S., Kamari Direh, Z., Ahmadi, S., Azizi, Z., & Danialy, D. (2022). Improve Data Classification Performance in Diagnosing Diabetes Using the Binary Exchange Market Algorithm. *Journal of Big Data*, 9(1), Article ID: 43. <https://doi.org/10.1186/s40537-022-00598-z>



- Sadeghi, S., Khalili, D., Ramezankhani, A., Mansournia, M. A., & Parsaeian, M. (2022). Diabetes Mellitus Risk Prediction in the Presence of Class Imbalance Using Flexible Machine Learning Methods. *BMC Medical Informatics and Decision Making*, 22(1), Article ID: 36. <https://doi.org/10.1186/s12911-022-01775-z>
- Sisodia, D., & Sisodia, D. S. (2018). Prediction of Diabetes Using Classification Algorithms. *Procedia Computer Science*, 132, 1578–1585. <https://doi.org/10.1016/j.procs.2018.05.122>
- Torkey, H., Ibrahim, E., Hemdan, E. E.-D., El-Sayed, A., & Shouman, M. A. (2022). Diabetes Classification Application with Efficient Missing and Outliers Data Handling Algorithms. *Complex & Intelligent Systems*, 8(1), 237–253. <https://doi.org/10.1007/s40747-021-00349-2>
- Wang, X., Zhai, M., Ren, Z., Ren, H., Li, M., Quan, D., Chen, L., & Qiu, L. (2021). Exploratory Study on Classification of Diabetes Mellitus Through a Combined Random Forest Classifier. *BMC Medical Informatics and Decision Making*, 21(1), Article ID: 105. <https://doi.org/10.1186/s12911-021-01471-4>
- Weerahandi, H. M., Horwitz, L. I., & Blecker, S. B. (2020). Diabetes Phenotyping Using the Electronic Health Record. *Journal of General Internal Medicine*, 35(12), 3716–3718. <https://doi.org/10.1007/s11606-020-06231-0>
- Wichmann, R. M., Fernandes, F. T., Chiavegatto Filho, A. D. P., Ciconelle, A. C. M., de Brito, A. M. E. S., Nunes, B. P., Silva, D. L. e, Anschau, F., de Castro Rodrigues, H., Rocha, H. A. L., dos Reis, J. C. B., de Oliveira Cavalcante, L., de Oliveira, L. P., dos Santos Andrade, L. S., Nasi, L. A., de Maria Felix, M., Mimica, M. J., de Almeida Araujo, M. E., Arnoni, M. V., ... Nuno, V. L. eSant'ana. (2023). Improving the Performance of Machine Learning Algorithms for Health Outcomes Predictions in Multicentric Cohorts. *Scientific Reports*, 13(1), Article ID: 1022. <https://doi.org/10.1038/s41598-022-26467-6>
- Zee, B., Lee, J., Lai, M., Chee, P., Rafferty, J., Thomas, R., & Owens, D. (2022). Digital Solution for Detection of Undiagnosed Diabetes Using Machine Learning-Based Retinal Image Analysis. *BMJ Open Diabetes Research & Care*, 10(6), Article ID: e002914. <https://doi.org/10.1136/bmjdr-2022-002914>
- Zheng, Y., Ley, S. H., & Hu, F. B. (2018). Global Aetiology and Epidemiology of Type 2 Diabetes Mellitus and Its Complications. *Nature Reviews Endocrinology*, 14(2), 88–98. <https://doi.org/10.1038/nrendo.2017.151>
- Zhou, H., Myrzashova, R., & Zheng, R. (2020). Diabetes Prediction Model Based on an Enhanced Deep Neural Network. *EURASIP Journal on Wireless Communications and Networking*, 2020(1), Article ID: 148. <https://doi.org/10.1186/s13638-020-01765-7>
- Zhou, L., Zheng, X., Yang, D., Wang, Y., Bai, X., & Ye, X. (2021). Application of Multi-Label Classification Models for the Diagnosis of Diabetic Complications. *BMC Medical Informatics and Decision Making*, 21(1), Article ID: 182. <https://doi.org/10.1186/s12911-021-01525-7>



Analisis Ketertarikan Pengguna Microsoft Excel Online untuk Pengolahan Data Silsilah Keluarga Menggunakan TAM dan TPB

Fathur Rachman Nufaily ^{(1)*}, Maria Ulfah Siregar ⁽²⁾

¹ Departemen Informatika, UIN Sunan Kalijaga, Yogyakarta, Indonesia

² Departemen Magister Informatika, UIN Sunan Kalijaga, Yogyakarta, Indonesia

e-mail : fr.nufaily@gmail.com, maria.siregar@uin-suka.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 3 Juli 2024, direvisi 29 November 2024, diterima 30 November 2024, dan dipublikasikan 30 September 2025.

Abstract

The use of web-based applications such as Microsoft Excel Online has increased, including for recording family genealogy data. This study aims to analyze the factors influencing the intention and behavior of using this application based on the Technology Acceptance Model (TAM), Theory of Planned Behavior (TPB), and their combined framework. The constructs examined include perceived ease of use, perceived usefulness, attitude, subjective norm, perceived behavioral control, intention, and behavior. This quantitative study collected primary data through questionnaires distributed to family members using Microsoft Excel Online. Data analysis was conducted using SEM-PLS (Structural Equation Modeling-Partial Least Squares) with the assistance of SmartPLS version 4.1.0.2. The results indicate that perceived ease of use and perceived usefulness positively and significantly affect attitude, while attitude, subjective norm, and perceived behavioral control positively influence behavioral intention. Furthermore, behavioral intention has a positive effect on actual usage behavior. These findings suggest that Microsoft Excel Online is reliable for recording family genealogy data and supports technology acceptance among users.

Keywords: Technology Acceptance Model, Theory of Planned Behavior, Microsoft Excel Online, Genealogy, Family

Abstrak

Penggunaan aplikasi berbasis web seperti Microsoft Excel Online semakin meningkat, termasuk untuk pencatatan data silsilah keluarga. Penelitian ini bertujuan untuk menganalisis faktor yang memengaruhi niat dan perilaku penggunaan aplikasi tersebut berdasarkan Technology Acceptance Model (TAM), Theory of Planned Behavior (TPB), serta gabungan keduanya. Konstruk yang diuji meliputi persepsi kemudahan, persepsi kegunaan, sikap, norma subjektif, kontrol perilaku, niat, dan perilaku. Metode penelitian yang digunakan adalah kuantitatif dengan pengumpulan data primer melalui kuesioner yang dibagikan kepada anggota keluarga pengguna Microsoft Excel Online. Analisis dilakukan menggunakan SEM-PLS (Structural Equation Model-Partial Least Square) dengan bantuan SmartPLS versi 4.1.0.2. Hasil penelitian menunjukkan bahwa persepsi kemudahan dan persepsi kegunaan berpengaruh positif dan signifikan terhadap sikap, sedangkan sikap, norma subjektif, dan persepsi kontrol perilaku berpengaruh positif terhadap niat penggunaan. Selanjutnya, niat penggunaan berpengaruh positif terhadap perilaku penggunaan. Temuan ini menunjukkan bahwa Microsoft Excel Online cukup andal untuk pencatatan data silsilah keluarga dan mendukung penerimaan teknologi oleh penggunanya.

Kata Kunci: Technology Acceptance Model, Theory of Planned Behavior, Microsoft Excel Online, Silsilah, Keluarga

1. PENDAHULUAN

Topik hubungan kekeluargaan selalu hangat diperbincangkan dari masa ke masa. Seperti penelitian yang dilakukan oleh (Almuashi et al., 2022; Chergui et al., 2020; Hormann et al., 2020; Liu et al., 2022; Mukherjee et al., 2022; Robinson et al., 2021; Serraoi et al., 2022; Shadrikov, 2020; Wahlström Henriksson & Goedecke, 2021; M. Wang et al., 2020; W. Wang et al., 2023; Yan & Song, 2021; J. Yu et al., 2020). Penelitian-penelitian tersebut bertujuan untuk



Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

memverifikasi hubungan kekerabatan menggunakan survei, *software*, dan lain-lain. Penelitian yang dilakukan oleh Dudgeon et al. (2021) menyimpulkan bahwa memperkuat hubungan kekeluargaan merupakan salah satu upaya pencegahan dari bunuh diri. Dalam era digital saat ini, penguasaan teknologi adalah sesuatu yang menentukan kemajuan suatu negara (Iqbal et al., 2022). Salah satu bentuk penguasaan adalah kemampuan untuk menggunakan teknologi tersebut. Misalnya penggunaan aplikasi berbasis web seperti Microsoft Excel Online yang semakin meningkat. Aplikasi ini memungkinkan pengguna untuk mengakses, mengedit, dan berbagi dokumen dari mana saja dan kapan saja melalui browser, selagi memiliki perangkat komunikasi seperti *smartphone* dan memiliki koneksi internet yang stabil. Salah satu jenis data yang dapat ditulis dengan Microsoft Excel Online adalah data silsilah keluarga. Data silsilah keluarga adalah data yang mencatat hubungan kekerabatan antar anggota keluarga dari generasi ke generasi. Data ini dapat digunakan untuk berbagai tujuan, seperti mengetahui asal-usul keluarga, menjaga hubungan kekeluargaan, menemukan kerabat yang hilang atau sebagai alat pendukung silaturahmi. Pada penelitian yang dilakukan oleh Mardiono et al. (2024) dan W et al. (2021), dibangun suatu aplikasi pencatatan data silsilah keluarga agar anggota keluarga bisa mengetahui asal-usul garis keturunan mereka sehingga bisa memperkuat hubungan kekerabatan. Selain aplikasi tersebut, ada aplikasi FamilySearch Tree yang merupakan aplikasi buatan organisasi genealogi terbesar di dunia, *Genealogical Society of Utah* (Dharma et al., 2020).

Menulis data silsilah keluarga dengan Microsoft Excel Online memiliki tantangan tersendiri. Pertama, pengguna harus memiliki keterampilan dan pengetahuan yang cukup untuk menggunakan aplikasi ini secara optimal. Kedua, pengguna harus memiliki motivasi dan minat yang tinggi untuk mengisi, memperbarui, dan membagikan data silsilah keluarga secara *online*. Ketiga, pengguna harus memiliki akses dan dukungan yang memadai dari sumber-sumber eksternal, seperti internet, perangkat, dan kerabat. Oleh karena itu, penting untuk mengetahui faktor-faktor apa saja yang mempengaruhi niat penggunaan aplikasi Microsoft Excel Online untuk data silsilah keluarga.

Untuk menjawab pertanyaan tersebut, penelitian ini menggunakan teori-teori yang terkait dengan perilaku manusia, yaitu *Technology Acceptance Model* (TAM), *Theory of Planned Behavior* (TPB) serta gabungan dari keduanya. TAM pertama kali diusulkan oleh Fred Davis pada tahun 1986. Model penerimaan teknologi adalah model yang mengukur persepsi individu tentang kegunaan dan kemudahan penggunaan teknologi, yang mempengaruhi niat dan perilaku penggunaan sistem informasi. Model TPB dikembangkan oleh Icek Ajzen yang dikembangkan dari model *Theory of Reasoned Action* (TRA) dengan menambahkan konstruk baru berupa persepsi kontrol perilaku. Model TRA sendiri hanya mengukur sikap dan norma subjektif untuk menganalisis niat individu dalam berperilaku tertentu (Nurfauzan & Priyono, 2022).

Penelitian terdahulu menunjukkan bahwa Microsoft Excel digunakan untuk menganalisis data eksperimen dalam pembelajaran fisika (Nagoro & Wathon, 2023), mengolah nilai rapor siswa (Nursita et al., 2021) dan untuk keperluan ekonomi pada UKM (Nurleli et al., 2023). Adapun dalam penelitian ini, Microsoft Excel digunakan untuk mencatat data silsilah keluarga. Tentunya penggunaan ini memiliki keunikan tersendiri karena sebagian besar data yang ditulis bukan berupa angka. Sedangkan Microsoft Excel biasanya digunakan untuk menghitung dan mengolah data yang bersifat numerik (angka). Alasan teori TAM dan TPB dipilih adalah karena teori tersebut dianggap sebagai model dasar yang ideal untuk mempelajari niat individu dalam menggunakan suatu teknologi, karena TAM berfokus pada faktor-faktor dari perspektif teknologi, sedangkan TPB berfokus pada faktor-faktor yang terkait dengan karakteristik pengguna (Sun et al., 2022).

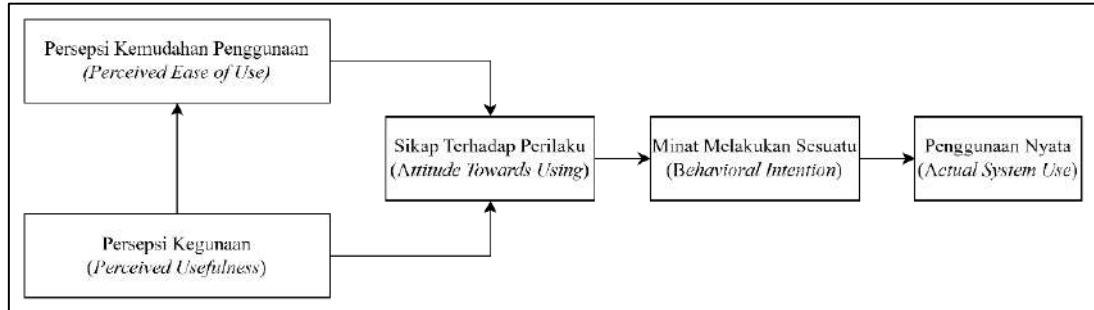
2. METODE PENELITIAN

2.1 Jenis Penelitian

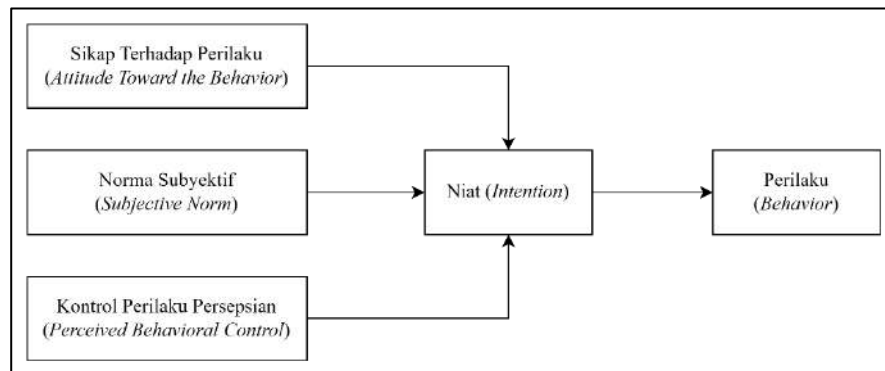
Jenis penelitian dalam penelitian ini menggunakan pendekatan kuantitatif sebagai metodologi utama. Metode penelitian kuantitatif dipilih karena memiliki proses yang sistematis, terstruktur, dan terukur dari awal hingga akhir. Pendekatan ini mencakup berbagai jenis penelitian seperti



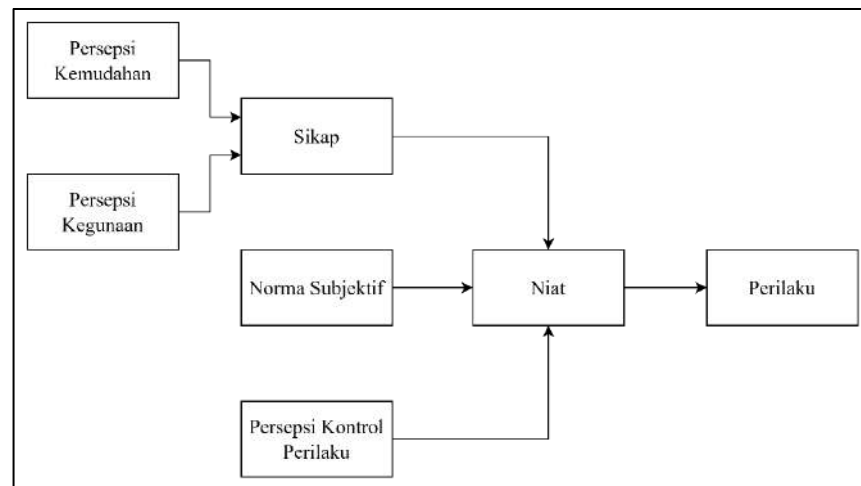
korelasi, deskriptif, kausal komparatif, eksperimen, survei, dan inferensial, yang pada prinsipnya bertujuan untuk menguji atau mengonfirmasi teori maupun asumsi melalui studi mendalam (Siroj et al., 2024).



Gambar 1 Model TAM (Technology Acceptance Model)



Gambar 2 Model TPB (Theory of Planned Behavior)



Gambar 3 Model Penelitian

Model penelitian yang digunakan dalam studi ini meliputi Model TAM (*Technology Acceptance Model*) sebagaimana ditunjukkan pada Gambar 1, Model TPB (*Theory of Planned Behavior*) pada Gambar 2, serta gabungan dari keduanya yang ditampilkan dalam Gambar 3 Model Penelitian. Model TAM digunakan untuk mengukur persepsi individu mengenai kegunaan dan kemudahan penggunaan teknologi yang berpengaruh pada niat serta perilaku penggunaan sistem informasi. Sementara itu, Model TPB digunakan untuk menilai persepsi individu terkait sikap terhadap



perilaku, norma subjektif, dan persepsi kontrol perilaku yang memengaruhi niat (Nurfauzan & Priyono, 2022).

Adapun hipotesis yang akan diuji dalam penelitian ini adalah sebagai berikut:

H1: Persepsi kemudahan penggunaan (*perceived ease of use*) berpengaruh positif terhadap sikap penggunaan (*attitude towards using*) aplikasi Microsoft Excel Online untuk data silsilah keluarga.

H2: Persepsi kegunaan (*perceived usefulness*) berpengaruh positif terhadap sikap penggunaan (*attitude towards using*) aplikasi Microsoft Excel Online untuk data silsilah keluarga.

H3: Sikap terhadap perilaku (*attitude toward the behavior*) berpengaruh positif terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga.

H4: Norma subjektif (*subjective norm*) berpengaruh positif terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga.

H5: Kontrol perilaku persepsian (*perceived behavioral control*) berpengaruh positif terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga.

H6: Niat (*intention*) menggunakan berpengaruh positif terhadap perilaku (*behavior*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga.

2.2 Jenis dan Sumber Data

Jenis data yang digunakan dalam penelitian ini merupakan data primer. Data primer adalah data yang diperoleh secara langsung dari sumber asli, yang kemudian diserahkan kepada pengumpul data atau peneliti untuk dianalisis. Data ini biasanya didapatkan melalui observasi, wawancara, survei, atau eksperimen yang dilakukan oleh peneliti sendiri, memastikan bahwa informasi yang diperoleh adalah akurat dan relevan untuk tujuan penelitian yang spesifik (Maria & Sutabri, 2023). Informasi atau data primer dari studi ini didapat dari tanggapan/jawaban atas pertanyaan kuesioner yang sesuai dengan situasi/keadaan/sentimen yang dirasakan oleh responden secara jujur. Data yang terkumpul berjumlah 98 responden yang diambil secara berkala dari bulan Februari hingga April 2024.

2.3 Teknik Pengumpulan Data

Penelitian kuantitatif dalam penelitian ini menggunakan metode pengumpulan data angket atau kuesioner. Dengan teknik ini, peneliti dapat mengumpulkan data dalam bentuk angka atau statistik yang dapat dianalisis secara kuantitatif. Skala Likert, yang diciptakan oleh Rensis Likert pada tahun 1932, merupakan metode pengukuran yang menggunakan serangkaian pernyataan atau item. Skala ini biasanya terdiri dari empat atau lebih item yang, ketika digabungkan, menghasilkan skor yang mencerminkan karakteristik tertentu dari individu, seperti pengetahuan, sikap, atau perilaku (Darmayanti & Sunarianingsih, 2024). Dalam analisis data, skor gabungan yang merupakan total atau rata-rata dari semua item dapat diaplikasikan. Pendekatan ini dianggap valid karena setiap item dianggap sebagai indikator yang mewakili variabel yang diteliti. Dalam penelitian ini, skala Likert yang digunakan adalah skor 1-5 sebagaimana pada Tabel 1. Skala penilaian 1-5 dipilih karena dapat membuat indikator lebih rinci dalam mengukur variabelnya (Abidin & Nuryana, 2023). Sistem “sangat setuju” hingga “sangat tidak setuju” merupakan pemeringkatan yang digunakan untuk mengukur tingkat persetujuan responden terhadap suatu pernyataan yang ada di dalam kuesioner. Selanjutnya, variabel penelitian beserta indikator konstruksi yang diukur dijabarkan secara rinci pada Tabel 2 Variabel Indikator Konstruksi, yang menjadi dasar dalam penyusunan butir-butir pernyataan pada kuesioner.

Tabel 1 Skala Likert

Penilaian	Skor
Sangat Setuju (SS)	5
Setuju (S)	4
Kurang Setuju (KS)	3
Tidak Setuju (TS)	2
Sangat Tidak Setuju (STS)	1



Tabel 2 Variabel Indikator Konstruksi

Variabel	Indikator	Kode
Persepsi Kemudahan (X1)	Mudah menggunakan	MUDAH1
	Mudah memahami	MUDAH2
	Tidak perlu banyak usaha	MUDAH3
Persepsi Kegunaan (X2)	Meningkatkan efektivitas	GUNA1
	Mempermudah pekerjaan	GUNA2
	Mempercepat pekerjaan	GUNA3
Sikap (X3)	Merasa nyaman	SIKAP1
	Menyenangkan	SIKAP2
	Memberi nilai bagus	SIKAP3
Norma Subjektif (X4)	Dukungan dari orang dekat	NORMS1
	Adanya tekanan sosial	NORMS2
	Adanya partisipasi orang lain	NORMS3
Persepsi Kontrol Perilaku (X5)	Adanya kendali diri	KOPRI1
	Bisa mengontrol penuh	KOPRI2
	Keyakinan diri dalam situasi sulit	KOPRI3
Niat (Y1)	Niat ketertarikan	NIAT1
	Pilihan utama dalam menggunakan	NIAT2
	Keinginan untuk tetap menggunakan	NIAT3
Perilaku (Y2)	Kebiasaan	PRI1
	Durasi penggunaan	PRI2
	Kecanduan menggunakan	PRI3

2.4 Teknik Analisis Data

Teknik analisis data dalam penelitian ini menggunakan teknik analisis deskriptif dan analisis inferensial berupa SEM-PLS (Structural Equation Modeling-Partial Least Squares). Statistik deskriptif adalah jenis analisis yang menunjukkan persentase dan frekuensi serta fenomena yang terdapat dalam data, terutama untuk menampilkan profil dan demografi responden. Analisis deskriptif adalah analisis yang mengandalkan statistik yang mampu mendeskripsikan dan menggambarkan subjek penelitian melalui sampel atau populasi apa adanya tanpa melakukan analisis dan membuat kesimpulan yang berlaku secara umum dalam sebuah penelitian. Tujuan analisis deskriptif adalah deskriptif bertujuan untuk membuat deskripsi, gambar, atau lukisan yang sistematis, menjelaskan hasil deskripsi dan memvalidasi kebenaran dan keakuratan hasil data tersebut (Samputra, 2022).

Statistik deskriptif, juga dikenal sebagai statistik deduktif, merupakan suatu proses pengumpulan, penyusunan, pengaturan, pengolahan, penyajian, dan analisis data numerik untuk memberikan gambaran yang jelas, ringkas, dan teratur tentang suatu fenomena, peristiwa, atau kondisi. Statistik deskriptif berfokus pada tugas mengorganisir dan menganalisis data numerik dengan tujuan memberikan deskripsi yang jelas, ringkas, dan teratur tentang suatu fenomena, peristiwa, atau kondisi, sehingga dapat diinterpretasikan dan dimaknai secara spesifik (Sumarni et al., 2023).

Analisis inferensial (induktif) memiliki tujuan untuk penarikan kesimpulan. Sebelum penarikan kesimpulan dilakukan suatu dugaan yang dapat diperoleh dari statistik deskriptif (Kartomo et al., 2024). Analisis ini memungkinkan diambilnya simpulan secara umum terhadap data yang diolah. Pengolahan data sampel dapat memprediksi dan mengendalikan seluruh populasi berdasarkan data, gejala, dan peristiwa yang ada pada proses penelitian. Fungsi ini dimulai dengan membuat suatu hipotesis atau prediksi tentang populasi (Siregar, 2021). Salah satu metode yang menggunakan statistik inferensial adalah SEM-PLS (Structural Equation Modeling-Partial Least Squares).

Analisis PLS adalah metode statistik multivariat yang membandingkan antara variabel dependen berganda dan variabel independen berganda pula. PLS merupakan salah satu teknik statistik



SEM berbasis varian yang dirancang untuk menangani regresi berganda ketika data mengalami masalah tertentu, seperti sampel penelitian yang kecil, data yang tidak lengkap (*missing value*) dan multikolinieritas. PLS kadang-kadang juga disebut *soft modeling* karena mengendurkan asumsi-asumsi regresi OLS (Ordinary Least Square) yang sangat ketat, seperti tidak terdapat multikolinieritas antara variabel bebas (Sholihin & Ratmono, 2021).

Salah satu perangkat lunak yang praktis untuk melakukan analisis SEM-PLS adalah SmartPLS. SmartPLS adalah sebuah perangkat lunak yang digunakan untuk mengolah data SEM (Structural Equation Modeling), yang popularitasnya meningkat pesat dalam penelitian pendidikan tinggi. Aplikasi ini direkomendasikan oleh Prof. Dae Shik Park dari Universitas Nasional Chungnam, Korea, khususnya untuk digunakan oleh mahasiswa. Selain itu, SmartPLS merupakan aplikasi statistik yang bisa diandalkan, mudah digunakan, dan jika diterapkan dengan baik, dapat sangat membantu peneliti dan dosen dalam proses penelitian (Tabelessy & Pattiruhu, 2022).

3. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan pendekatan kuantitatif untuk menganalisis data yang diperoleh. Analisis dilakukan dengan menggunakan statistik deskriptif serta pengujian hipotesis melalui Partial Least Square (PLS). Variabel yang terlibat dalam penelitian yaitu persepsi kemudahan, persepsi kegunaan, sikap, norma subjektif, persepsi kontrol perilaku, niat, dan perilaku.

3.1 Analisis deskriptif

Dalam penelitian ini, penyajian data deskriptif mencakup rata-rata, standar deviasi, median, nilai minimum, nilai maksimum, kurtosis, dan *skewness* (derajat ketidaksimetrisan suatu distribusi). Analisis deskriptif ini bertujuan untuk merangkum data sehingga lebih mudah dipahami dan memberikan gambaran awal mengenai karakteristik data yang digunakan. Uji kurtosis dan *skewness* digunakan sebagai indikator normalitas data, di mana data dikatakan normal apabila nilai kurtosis dan *skewness* berada dalam rentang -2 hingga 2 (Aji et al., 2020). Berdasarkan Tabel 3, beberapa item seperti GUNA2, GUNA3, NORMS1, NIAT1, dan NIAT2 memiliki nilai kurtosis melebihi 2, yang menunjukkan bahwa data item tersebut tidak normal. Namun demikian, analisis SEM-PLS tidak mensyaratkan normalitas data sehingga kondisi ini tidak menjadi kendala dalam proses analisis (Yamin, 2023).

Tabel 3 Analisis Deskriptif

Item Pengukuran	Rata-Rata	Median	Minimum	Maksimum	Standar Deviasi	Excess Kurtosis	Skewness
MUDAH1	4,01	4	1	5	1,025	0,871	-1,117
MUDAH2	3,949	4	1	5	1,146	-0,011	-0,931
MUDAH3	3,735	4	1	5	1,121	-0,306	-0,693
GUNA1	4,061	4	1	5	1,058	1,664	-1,385
GUNA2	4,337	5	1	5	1,078	2,822	-1,851
GUNA3	4,255	4	1	5	0,951	2,806	-1,62
SIKAP1	3,99	4	1	5	1,015	0,677	-0,99
SIKAP2	3,776	4	1	5	1,005	0,192	-0,695
SIKAP3	4,143	4	1	5	1,05	1,368	-1,313
NORMS1	3,867	4	1	5	0,888	2,28	-1,332
NORMS2	3,173	3	1	5	1,278	-1,021	-0,064
NORMS3	3,735	4	1	5	1,112	0,015	-0,813
KOPRI1	3,939	4	1	5	0,998	0,919	-1,064
KOPRI2	4,01	4	1	5	0,964	1,531	-1,202
KOPRI3	3,643	4	1	5	1,109	-0,232	-0,619
NIAT1	4,194	4	1	5	0,976	2,824	-1,606
NIAT2	4,204	4	1	5	1	2,403	-1,544
NIAT3	4,082	4	1	5	1,007	1,053	-1,141
PRI1	3,776	4	1	5	1,074	-0,211	-0,692
PRI2	3,735	4	1	5	1,208	-0,064	-0,883
PRI3	3,582	4	1	5	1,087	-0,251	-0,455



3.2 Analisis inferensial

Analisis inferensial terdiri dari dua tahapan analisis, yaitu analisis *outer model* dan analisis *inner model*. *Outer model* atau model pengukuran dilakukan dua pengujian, yaitu uji validitas dan uji reliabilitas. Validitas sendiri merujuk pada seberapa akurat instrumen tersebut dalam mengukur apa yang seharusnya diukur. Adapun reliabilitas menunjukkan bahwa hasil yang diperoleh dari instrumen tersebut akan konsisten dan stabil ketika pengukuran diulang berkali-kali dalam situasi yang serupa. Analisis *inner model* dilakukan untuk menjawab hipotesis dalam penelitian ini.

Validitas konvergen dapat ditentukan dengan cara melihat nilai *outer loadings* atau *loading factor* dari masing-masing indikator. Selain itu, perlu dilihat juga nilai AVE. Indikator yang memiliki nilai *loading factor* lebih dari 0,50 dan nilai AVE lebih dari 0,50 dianggap memiliki validitas konvergen yang tinggi (Putra et al., 2022). Formula untuk mencari nilai AVE ditunjukkan pada Pers. (1) (Yamin, 2023).

$$AVE = \frac{\sum \lambda_i^2}{\sum \lambda_i^2 + \sum \theta_i} \quad (1)$$

Tabel 4 Hasil Outer Loadings dan Nilai AVE

Variabel	Item Pengukuran	Indikator	Outer Loadings	Nilai AVE	Ket.
Persepsi Kemudahan	MUDAH1	Mudah menggunakan	0,853	0,766	Valid
	MUDAH2	Mudah memahami	0,906		Valid
	MUDAH3	Tidak perlu banyak usaha	0,866		Valid
Persepsi Kegunaan	GUNA1	Meningkatkan efektivitas	0,893	0,823	Valid
	GUNA2	Mempermudah pekerjaan	0,919		Valid
	GUNA3	Mempercepat pekerjaan	0,910		Valid
Sikap	SIKAP1	Merasa nyaman	0,897	0,733	Valid
	SIKAP2	Menyenangkan	0,860		Valid
	SIKAP3	Meberi nilai bagus	0,880		Valid
Norma Subjektif	NORMS1	Dukungan dari orang dekat	0,890	0,661	Valid
	NORMS2	Adanya tekanan sosial	0,667		Valid
	NORMS3	Adanya partisipasi orang lain	0,864		Valid
Persepsi Kontrol Perilaku	KOPRI1	Adanya kendali diri	0,881	0,776	Valid
	KOPRI2	Bisa mengontrol penuh	0,917		Valid
	KOPRI3	Keyakinan diri dalam situasi sulit	0,844		Valid
Niat	NIAT1	Niat ketertarikan	0,930	0,859	Valid
	NIAT2	Pilihan utama dalam menggunakan	0,921		Valid
	NIAT3	Keinginan untuk tetap menggunakan	0,929		Valid
Perilaku	PRI1	Kebiasaan	0,867	0,772	Valid
	PRI2	Durasi penggunaan	0,888		Valid
	PRI3	Kecanduan menggunakan	0,881		Valid

Tabel 4 adalah hasil dari uji validitas konvergen menggunakan *software* SmartPLS. Pada Tabel 4 semua *item* pengukuran memiliki nilai *outer loadings* lebih dari 0,5 serta keseluruhan nilai AVE melebihi 0,5. Itu artinya, *item* pengukuran dalam penelitian ini sudah memenuhi syarat validitas konvergen. Namun, di antara beberapa *item* pengukuran tersebut, terdapat satu *item* pengukuran yang memiliki nilai *outer loadings* lebih rendah dari pada yang lain, yaitu NORMS2 (0,667).

Analisis validitas diskriminan perlu dilakukan dengan mempertimbangkan kriteria *cross-loadings* dan Fornell-Larcker. Validitas diskriminan merupakan bentuk evaluasi yang bertujuan untuk memastikan bahwa variabel secara teoritis berbeda dan terbukti secara empiris melalui pengujian statistik. *Loading* sebuah indikator pada variabel yang diukur harus lebih besar dari pada *loading*



terhadap variabel lainnya. Pengukuran ini disebut sebagai *cross-loadings*. Sedangkan kriteria Fornell-Larcker menyatakan bahwa akar dari Average Variance Extracted (AVE) pada suatu variabel harus lebih besar daripada korelasi antara variabel tersebut dengan variabel laten lainnya (Sholihin & Ratmono, 2021).

Tabel 5 Hasil Cross Loadings

Item Pengukuran	Kegunaan	Kemudahan	Kontrol Perilaku	Niat	Norma Subjektif	Perilaku	Sikap
GUNA1	0,893	0,766	0,745	0,796	0,666	0,661	0,759
GUNA2	0,919	0,695	0,702	0,846	0,547	0,489	0,753
GUNA3	0,91	0,62	0,688	0,811	0,571	0,479	0,707
KOPRI1	0,697	0,718	0,881	0,732	0,628	0,599	0,725
KOPRI2	0,769	0,741	0,917	0,768	0,621	0,593	0,739
KOPRI3	0,597	0,642	0,844	0,639	0,57	0,65	0,75
MUDAH1	0,666	0,853	0,621	0,621	0,53	0,473	0,563
MUDAH2	0,744	0,906	0,756	0,755	0,551	0,605	0,706
MUDAH3	0,601	0,866	0,704	0,627	0,512	0,58	0,656
NIAT1	0,861	0,718	0,76	0,93	0,653	0,609	0,734
NIAT2	0,86	0,735	0,768	0,921	0,69	0,596	0,785
NIAT3	0,788	0,68	0,731	0,929	0,639	0,695	0,722
NORMS1	0,757	0,646	0,685	0,766	0,89	0,503	0,737
NORMS2	0,171	0,328	0,373	0,306	0,667	0,197	0,277
NORMS3	0,469	0,413	0,541	0,523	0,864	0,371	0,428
PRI1	0,554	0,604	0,668	0,617	0,389	0,867	0,597
PRI2	0,546	0,592	0,586	0,625	0,407	0,888	0,523
PRI3	0,479	0,472	0,569	0,556	0,468	0,881	0,593
SIKAP1	0,777	0,701	0,78	0,762	0,643	0,636	0,897
SIKAP2	0,667	0,612	0,702	0,641	0,542	0,501	0,86
SIKAP3	0,701	0,627	0,717	0,715	0,53	0,565	0,88

Tabel 6 Hasil Fornell-Larcker

Variabel	Persepsi Kegunaan	Persepsi Kemudahan	Persepsi Kontrol Perilaku	Niat	Norma Subjektif	Perilaku	Sikap
Persepsi Kegunaan	0,907						
Persepsi Kemudahan	0,767	0,875					
Persepsi Kontrol Perilaku	0,785	0,797	0,881				
Niat	0,902	0,767	0,812	0,927			
Norma Subjektif	0,656	0,606	0,689	0,713	0,813		
Perilaku	0,601	0,636	0,693	0,684	0,478	0,879	
Sikap	0,816	0,738	0,835	0,806	0,653	0,649	0,879

Pada Tabel 5, dapat dilihat bahwa nilai *loading* dari indikator GUNA1, GUNA2, dan GUNA3 dalam mengukur variabel “Kegunaan” memiliki nilai yang lebih besar dari pada *loading* indikator tersebut dalam mengukur variabel yang lainnya. Begitu pula dengan indikator-indikator yang lain dalam mengukur variabelnya sendiri memiliki nilai *loading* yang lebih tinggi dari pada nilai *loading* dalam mengukur variabel yang lainnya. Pada Tabel 6, dapat dilihat bahwa akar dari Average Variance Extracted (AVE) pada variabel “Kegunaan” lebih besar daripada korelasi antara variabel tersebut dengan variabel laten lainnya. Begitu juga dengan variabel-variabel yang lain. Hal ini mengindikasikan bahwa validitas diskriminan telah terpenuhi.



Reliabilitas menunjukkan akurasi, ketepatan, dan konsistensi instrumen dalam pengukuran variabel. *Internal consistency reliability* menguji kemampuan indikator dalam mengukur variabel latennya. *Internal consistency reliability* dapat dilihat dari nilai *composite reliability* dan *cronbach alpha*. Nilai *composite reliability* dan *cronbach alpha* adalah > 0,60 (Sholihin & Ratmono, 2021). Formula untuk mencari *composite reliability* ditunjukkan pada Pers. (2) (Yamin, 2023). Tabel 7 merupakan hasil dari pengujian Composite Reliability dan Cronbach's Alpha menggunakan SmartPLS. Tabel tersebut menunjukkan nilai *Composite Reliability* dan *Cronbach's Alpha* dari masing-masing variabel. Nilainya menunjukkan > 0,60 yang berarti telah memenuhi kriteria reliabilitas.

$$CR = \frac{(\sum \lambda_i)^2}{(\sum \lambda_i)^2 + \sum \theta_i} \quad (2)$$

Tabel 7 Hasil Composite Reliability dan Cronbach's Alpha

Variabel	Composite Reliability	Cronbach's Alpha	Keterangan
Persepsi Kegunaan	0,933	0,892	Reliabel
Persepsi Kemudahan	0,907	0,848	Reliabel
Persepsi Kontrol Perilaku	0,912	0,856	Reliabel
Niat	0,948	0,918	Reliabel
Norma Subjektif	0,852	0,761	Reliabel
Perilaku	0,910	0,853	Reliabel
Sikap	0,911	0,853	Reliabel

Sebelum memasuki analisis inner model, perlu dilakukan pemeriksaan multikolinier antara variabel dengan *inner VIF* (*Variance Inflated Factor*). Hasil yang menunjukan nilai *inner VIF* kurang dari 5 maka tidak ada multikolinier antara variabel (Yamin, 2023). Adapun formula untuk *inner VIF* ditunjukkan pada Pers. (3) (Yamin, 2021). Hasil uji kolinieritas pada Tabel 8 menunjukkan tidak adanya multikolinier pada hubungan antara variabel. Hal ini dibuktikan dengan hasil estimasi *inner VIF* kurang dari 5.

$$VIF_i = \frac{1}{1 - R_i^2} \quad (3)$$

Tabel 8 Hasil Uji Kolinieritas

Hubungan Variabel	Inner VIF	Keterangan
Kegunaan > Sikap	2,427	Tidak ada multikolinier
Kemudahan > Sikap	2,427	Tidak ada multikolinier
Kontrol Perilaku > Niat	3,759	Tidak ada multikolinier
Niat > Perilaku	1,000	Tidak ada multikolinier
Norma Subjektif > Niat	1,978	Tidak ada multikolinier
Sikap > Niat	3,439	Tidak ada multikolinier

Tabel 9 Hasil Uji Signifikansi Path Coefficients dan P-Value

Hipotesis	Koefisien	P-Value	Keterangan
Kegunaan > Sikap	0,608	0,000	Tidak ada multikolinier
Kemudahan > Sikap	0,272	0,003	Tidak ada multikolinier
Kontrol Perilaku > Niat	0,347	0,003	Tidak ada multikolinier
Niat > Perilaku	0,684	0,000	Tidak ada multikolinier
Norma Subjektif > Niat	0,239	0,006	Tidak ada multikolinier
Sikap > Niat	0,360	0,003	Tidak ada multikolinier

Pengambilan keputusan dalam pengujian hipotesis dapat dilihat dari nilai Path Coefficients dan tingkat signifikansi (p-value). Jika nilai Path Coefficients berada di antara 0 sampai dengan 1



maka dapat dikatakan berpengaruh positif, sedangkan jika nilai berada di antara -1 sampai dengan 0 maka dapat dikatakan negatif (tidak berpengaruh). Adapun p-value menyatakan tingkat pengaruh signifikansi antara variabel. Jika p-value < 0,05 maka berpengaruh secara signifikan. Jika p-value > 0,05 maka berpengaruh tidak signifikan (Yamin, 2023).

Berdasarkan hasil pengujian pada Tabel 9, maka diketahui sebagai berikut:

- 1) Persepsi kemudahan penggunaan (*perceived ease of use*) berpengaruh terhadap sikap penggunaan (*attitude towards using*) aplikasi Microsoft Excel Online untuk data silsilah keluarga. Hasil pengukuran hubungan antara kemudahan terhadap sikap menunjukkan nilai koefisien sebesar 0,272 dan p-value 0,003. Hal ini menunjukkan bahwa persepsi kemudahan penggunaan berpengaruh positif terhadap sikap penggunaan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Artinya hipotesis 1 (H1) diterima.
- 2) Persepsi kegunaan (*perceived usefulness*) berpengaruh terhadap sikap penggunaan (*attitude towards using*) aplikasi Microsoft Excel Online untuk data silsilah keluarga. Hasil estimasi hubungan antara kegunaan terhadap sikap menunjukkan nilai koefisien sebesar 0,608 dan p-value 0,000. Hal ini menunjukkan bahwa persepsi kegunaan berpengaruh positif terhadap sikap penggunaan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Dengan kata lain, hipotesis 2 (H2) diterima.
- 3) Sikap terhadap perilaku (*attitude toward the behavior*) berpengaruh terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga. Hubungan antara sikap terhadap niat memiliki nilai koefisien sebesar 0,360 dengan p-value 0,003. Hal ini menunjukkan bahwa sikap terhadap perilaku penggunaan aplikasi Microsoft Excel Online berpengaruh positif terhadap niat menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Dengan kata lain, hipotesis 3 (H3) diterima.
- 4) Norma subjektif (*subjective norm*) berpengaruh terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga. Hasil pengujian hubungan antara norma subjektif terhadap niat menggunakan menunjukkan nilai koefisien sebesar 0,239 dan p-value 0,006. Hal ini menunjukkan bahwa norma subjektif berpengaruh positif terhadap niat menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Dengan kata lain, hipotesis 4 (H4) diterima.
- 5) Kontrol perilaku persepsian (*perceived behavioral control*) berpengaruh terhadap niat (*intention*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga. Nilai koefisien dari persepsi kontrol perilaku terhadap niat menggunakan adalah 0,347 dengan p-value senilai 0,003. Hal ini menunjukkan bahwa persepsi kontrol perilaku berpengaruh positif terhadap niat menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Dengan kata lain, hipotesis 5 (H5) diterima.
- 6) Niat (*intention*) menggunakan berpengaruh terhadap perilaku (*behavior*) menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga. Hasil pengukuran hubungan antara niat menggunakan terhadap perilaku menggunakan menunjukkan nilai koefisien sebesar 0,684 dan p-value 0,000. Hal ini menunjukkan bahwa niat menggunakan berpengaruh positif terhadap perilaku menggunakan aplikasi Microsoft Excel Online untuk data silsilah keluarga dan berpengaruh signifikan. Dengan kata lain, hipotesis 6 (H6) diterima.

Tabel 10 Hasil Uji R-Square

Variabel Dependen	R-Square	R-Square Adjusted	Keterangan
Sikap	0,697	0,690	Pengaruh kuat
Niat	0,742	0,734	Pengaruh kuat
Perilaku	0,468	0,462	Pengaruh sedang

R-Square adalah ukuran yang menunjukkan seberapa kuat pengaruh variabel independen (eksogen) terhadap variabel dependen (endogen). R-Square adalah nilai yang berada di antara 0 dan 1 yang menunjukkan seberapa besar gabungan variabel independen memengaruhi nilai variabel dependen secara bersamaan. Nilai R-Square memiliki tiga kelompok kategori yaitu kategori tinggi, kategori sedang, dan kategori rendah. Nilai R-Square 0,75 termasuk kategori kuat, nilai R-Square 0,50 termasuk kategori sedang dan nilai R-Square 0,25 termasuk kategori lemah (Hair et al., 2019). Tabel 10 adalah hasil dari pengukuran R-Square. Semakin mendekati



angka 1, semakin tinggi nilai R-Square, maka semakin besar pengaruh variabel independen terhadap variabel dependen. Sisanya, yang tidak dapat dijelaskan oleh variabel independen, merupakan komponen error.

F-Square mengukur seberapa besar perubahan R-Square ketika variabel eksogen tertentu yang berpengaruh terhadap variabel endogen dihilangkan dari model. Nilai F-Square juga memiliki tiga kelompok kategori yaitu kategori tinggi, kategori sedang, dan kategori rendah. Nilai F-Square 0,35 termasuk kategori kuat, nilai F-Square 0,15 termasuk kategori sedang dan nilai F-Square 0,02 termasuk kategori lemah (Hair et al., 2019). Formula untuk mencari nilai F-Square ditunjukkan pada Pers. (4) (Yamin, 2023). Berdasarkan hasil perhitungan pada Tabel 11, dapat diketahui bahwa nilai F-Square kategori tinggi adalah pengaruh kegunaan terhadap sikap (0,502) dan pengaruh niat terhadap kontrol perilaku (0,880). Sedangkan nilai F-Square yang lainnya termasuk dalam kategori rendah, yaitu di bawah 0,15.

$$f^2 = \frac{R_{included}^2 - R_{excluded}^2}{1 - R_{included}^2} \quad (4)$$

Tabel 11 Hasil Uji F-Square

Hipotesis	F-Square	Keterangan
Kegunaan > Sikap	0,502	Tinggi
Kemudahan > Sikap	0,100	Rendah
Kontrol Perilaku > Niat	0,124	Rendah
Niat > Perilaku	0,880	Tinggi
Norma Subjektif > Niat	0,112	Rendah
Sikap > Niat	0,147	Rendah

Goodness of Fit (GoF) digunakan untuk memvalidasi model struktural secara menyeluruh. GoF *index* adalah ukuran tunggal yang memvalidasi performa gabungan dari model pengukuran dan model struktural. Nilai GoF didapat dari akar kuadrat dari hasil kali antara *average communalities index* dan nilai rata-rata R-Square. Nilai GoF berkisar antara 0 sampai 1 dengan kriteria: 0,1 (GoF rendah), 0,25 (GoF sedang), dan 0,36 (GoF tinggi). Formula untuk mencari GoF index ditunjukkan pada Pers. (5) (Yamin, 2023). Berdasarkan hasil perhitungan pada Tabel 12, rata-rata *communality* bernilai 0,776 yang merupakan nilai rata-rata dari semua nilai *outer loadings* yang dikuadratkan. Sedangkan nilai GoF index adalah 0,702 yang merupakan akar dari hasil perkalian antara rata-rata *communality* dan rata-rata R-Square. Nilai tersebut mengindikasikan bahwa GoF *index* dalam penelitian ini termasuk dalam kategori tinggi.

$$GoF = \sqrt{AVE \times R^2} \quad (5)$$

Tabel 12 Hasil Uji GoF Index

Hipotesis	F-Square	Keterangan
Kegunaan > Sikap	0,502	Tinggi

Tabel 13 Hasil Uji F-Square

Variabel Endogen	F-Square	Keterangan
Niat	0,502	Memiliki Predictive Relevance
Perilaku	0,100	Memiliki Predictive Relevance
Sikap	0,124	Memiliki Predictive Relevance

F-Square atau *predictive relevance* berperan untuk memvalidasi model. Pengukuran ini dilakukan jika variabel laten endogen menggunakan model pengukuran reflektif. Nilai *predictive relevance* yang baik adalah jika nilainya > 0 yang menandakan variabel laten eksogen cocok (sesuai) sebagai variabel penjelas yang dapat memprediksi variabel endogennya (Salsabila et al., 2024). Dapat dilihat pada Tabel 13 bahwa F-Square dari semua variabel endogen memiliki



nilai yang lebih besar dari 0, yang berarti sudah memiliki *predictive relevance*. Variabel “niat” memiliki *predictive relevance* yang paling tinggi dari pada variabel “perilaku” dan “sikap”.

3.3 Model Pembandingan

Peneliti akan membandingkan uji signifikansi Path Coefficients dari model gabungan TAM dan TPB dengan model TAM tanpa digabung dan model TPB tanpa digabung. Hasil komparasi tidak ada perubahan nilai pada Tabel 14 dan Tabel 15 kecuali pengaruh variabel sikap terhadap niat dalam model TAM. Pada model gabungan TAM dan TPB, nilai koefisien menunjukkan angka 0,360. Pada model TAM, nilai koefisien menunjukkan angka yang lebih besar, yaitu 0,806. Hal ini menunjukkan bahwa dalam model TAM, variabel sikap memiliki pengaruh yang lebih besar terhadap niat dari pada pengaruh sikap terhadap niat dalam model gabungan. Hal ini wajar karena dalam model gabungan, niat dipengaruhi oleh tiga variabel, sedangkan dalam model TAM niat hanya dipengaruhi oleh satu variabel saja.

Tabel 14 Komparasi Hasil Uji Signifikansi Path Coefficients TAM dan Gabungan

Hipotesis	Koefisien	Koefisien TAM
Kegunaan > Sikap	0,608	0,608
Kemudahan > Sikap	0,272	0,272
Kontrol Perilaku > Niat	0,347	
Niat > Perilaku	0,684	0,684
Norma Subjektif > Niat	0,239	
Sikap > Niat	0,360	0,806

Tabel 15 Komparasi Hasil Uji Signifikansi Path Coefficients TPB dan Gabungan

Hipotesis	Koefisien	Koefisien TPB
Kegunaan > Sikap	0,608	
Kemudahan > Sikap	0,272	
Kontrol Perilaku > Niat	0,347	0,346
Niat > Perilaku	0,684	0,684
Norma Subjektif > Niat	0,239	0,239
Sikap > Niat	0,360	0,361

4. KESIMPULAN

Berdasarkan hasil analisis dan pembahasan sebelumnya, dapat disimpulkan bahwa variabel persepsi kemudahan dan persepsi kegunaan memiliki pengaruh signifikan terhadap sikap penggunaan aplikasi Microsoft Excel Online dalam pencatatan data silsilah keluarga. Hal ini ditunjukkan oleh nilai koefisien sebesar 0,272 dari variabel kemudahan dan 0,608 dari variabel kegunaan dalam memengaruhi sikap. Selanjutnya, variabel sikap, norma subjektif, dan kontrol perilaku juga terbukti berpengaruh signifikan terhadap niat penggunaan aplikasi tersebut, dengan masing-masing nilai koefisien sebesar 0,360, 0,239, dan 0,347. Niat penggunaan aplikasi Microsoft Excel Online pun berpengaruh signifikan terhadap perilaku penggunaan aplikasi tersebut, sebagaimana dibuktikan oleh nilai koefisien sebesar 0,684. Adapun perbandingan antara metode TAM dan gabungan serta TPB dan gabungan menunjukkan tidak adanya perbedaan signifikan, kecuali pada variabel sikap dalam memengaruhi niat antara metode TAM dan gabungan, sebagaimana tercantum dalam tabel 14.

Penelitian ini memiliki beberapa kekurangan dan keterbatasan. Pertama, sampel yang digunakan hanya berasal dari satu komunitas keluarga di daerah Sumenep yang menggunakan Microsoft Excel sebagai alat pencatatan silsilah keluarga, padahal komunitas keluarga lain juga mengenal dan mungkin menggunakan aplikasi tersebut untuk tujuan serupa. Kedua, penelitian ini tergolong dalam studi keperilakuan yang masih berfokus pada niat perilaku, sehingga belum sepenuhnya melibatkan individu yang benar-benar telah berperilaku menggunakan Microsoft Excel untuk pencatatan silsilah keluarga mereka.



Sebagai saran, peneliti selanjutnya diharapkan dapat mengumpulkan lebih banyak komunitas keluarga yang menggunakan Microsoft Excel sebagai aplikasi pencatatan silsilah, guna meningkatkan kualitas penelitian. Selain itu, penggunaan metode kualitatif lapangan juga disarankan agar dapat menggali lebih dalam kelebihan dan kekurangan aplikasi Microsoft Excel dalam pencatatan silsilah keluarga, sehingga dapat menjadi dasar dalam pengembangan perangkat lunak khusus untuk keperluan tersebut.

DAFTAR PUSTAKA

- Abidin, A. R., & Nuryana, I. K. D. (2023). Perbandingan Metode Klasifikasi Data Mining untuk Mengukur Tingkat Kepuasan Mahasiswa Terhadap Sistem Informasi Penilaian Nonakademik UNESA (SIPENA). *Journal of Emerging Information Systems and Business Intelligence*, 4(4), 129–138. <https://doi.org/10.26740/jeisbi.v4i4.56966>
- Aji, H., Irmansyah, D., Wicaksono, A., & Gibran, A. F. (2020). *Ukuran Penyebaran Data (Kemiringan & Keruncingan) "Angka Terjangkit DBD Tahun 2008-2017."* <https://doi.org/10.31219/osf.io/tym4s>
- Almuashi, M., Hashim, S. Z. M., Yusoff, N., Syazwan, K. N., & Ghabban, F. (2022). Siamese Convolutional Neural Network and Fusion of the Best Overlapping Blocks for Kinship Verification. *Multimedia Tools and Applications*, 81(27), 39311–39342. <https://doi.org/10.1007/s11042-022-12735-0>
- Chergui, A., Ouchtati, S., Mavromatis, S., Bekhouche, S., Lashab, M., & Sequeira, J. (2020). Kinship Verification Through Facial Images Using CNN-Based Features. *Traitement Du Signal*, 37(1), 1–8. <https://doi.org/10.18280/ts.370101>
- Darmayanti, N. W. S., & Sunarianingsih, N. L. P. (2024). Portrait of Students' Science Process Skills in Science Learning at Elementary School 1 Susut. *Dharmas Education Journal (DE_Journal)*, 5(2), 1183–1193. <https://doi.org/10.56667/dejournal.v5i2.1626>
- Dharma, A. C. L. H., Aknuranda, I., & Rokhmawati, R. I. (2020). Pengujian Usability untuk Aplikasi Silsilah Keluarga FamilySearch Tree dengan Pendekatan Kuantitatif. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 4, 2226–2235. <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/7572>
- Dudgeon, P., Blustein, S., Bray, A., Calma, T., McPhee, R., & Ring, I. (2021). *Connection Between Family, Kinship and Social and Emotional Wellbeing*. <https://doi.org/10.25816/jx22-vq08>
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate Data Analysis* (8th ed.). Cengage Learning.
- Hormann, S., Knoche, M., & Rigoll, G. (2020). A Multi-Task Comparator Framework for Kinship Verification. *2020 15th IEEE International Conference on Automatic Face and Gesture Recognition (FG 2020)*, 863–867. <https://doi.org/10.1109/FG47880.2020.00106>
- Iqbal, D. M., Arista, D. N., Salsabila, H., & Nur, S. Q. (2022). Penyuluhan Penggunaan Microsoft Excel dalam Mengolah Laporan Kegiatan Pemberdayaan Kesejahteraan Keluarga di Desa Kertamulya. *Proceedings UIN Sunan Gunung Djati Bandung*, 2(4), 13–28. <https://proceedings.uinsgd.ac.id/index.php/proceedings/article/view/1320>
- Kartomo, T., Taufik, D. A., Kartutu, S. J., Tenu, M. W., & Sudana, I. W. (2024). Analisis Peran Statistika Terapan Dalam Bidang Bisnis, Kesehatan, dan Lingkungan. *Journal of Comprehensive Science*, 3(2), 394–402. <https://doi.org/10.59188/JCS.V3i2.626>
- Liu, F., Li, Z., Yang, W., & Xu, F. (2022). Age-Invariant Adversarial Feature Learning for Kinship Verification. *Mathematics*, 10(3), Article ID: 480. <https://doi.org/10.3390/math10030480>
- Mardiono, A. P., Rosid, M. A., & Busono, S. (2024). Rancang Bangun Sistem Informasi Silsilah Keluarga Berbasis Website Menggunakan Metode Rapid Application Development (RAD). *JUPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, 9(4), 1869–1880. <https://doi.org/10.29100/jupi.v9i4.5516>
- Maria, F., & Sutabri, T. (2023). Pengukuran Kualitas Website E-Learning di SMA Muhammadiyah 1 Palembang dengan Metode WebQual. *Indonesian Journal of Multidisciplinary on Social and Technology*, 1(2), 121–127. <https://doi.org/10.31004/ijmst.v1i2.134>
- Mukherjee, M., Meenpal, T., & Goyal, A. (2022). FuseKin: Weighted Image Fusion Based Kinship Verification Under Unconstrained Age Group. *Journal of Visual Communication and Image Representation*, 84, 103470. <https://doi.org/10.1016/j.jvcir.2022.103470>



- Nagoro, N., & Wathon, A. (2023). Penggunaan Microsoft Excel dalam Analisis Data Eksperimen pada Pembelajaran Fisika. *Sistim Informasi Manajemen*, 6(1), 30–46. <https://oj.lapamu.com/index.php/sim/article/view/33>
- Nurfauzan, J. A., & Priyono, A. (2022). Analisis TAM dan TPB dalam Penerimaan Aplikasi Perdagangan Saham Seluler (Mobile) di Kalangan Investor di Indonesia. *Selekta Manajemen: Jurnal Mahasiswa Bisnis & Manajemen*, 1(4), 79–96. <https://journal.uui.ac.id/selma/article/view/24883>
- Nurleli, S. B., Faturkhman, A., & Ulfah, P. (2023). The Effect of Usefulness and Ease of Use on Behavioral Interest in Using the Microsoft Excel Application with Attitude of Use as a Mediation Variable in SMEs in Banyumas District. *Jurnal Riset Akuntansi Soedirman*, 2(1), 1–15. <https://doi.org/10.32424/1.jras.2023.2.1.7852>
- Nursita, L., Astina, A., Isakasari, I., & Amiruddin, I. (2021). Efektivitas Penggunaan Microsoft Excel dalam Pengolahan Nilai Rapor Siswa SMA Negeri 11 Bone. *Educational Leadership: Jurnal Manajemen Pendidikan*, 1(1), 1–9. <https://doi.org/10.24252/edu.v1i1.21994>
- Putra, I. G. C., Manuari, I. A. R., & Puspayanti, N. K. D. (2022). Pengaruh Corporate Governance Terhadap Profitabilitas dan Nilai Perusahaan Manufaktur yang Terdaftar di Bursa Efek Indonesia (BEI). *WACANA EKONOMI (Jurnal Ekonomi, Bisnis dan Akuntansi)*, 21(1), 105–118. <https://doi.org/10.22225/we.21.1.2022.105-118>
- Robinson, J. P., Qin, C., Shao, M., Turk, M. A., Chellappa, R., & Fu, Y. (2021). The 5th Recognizing Families in the Wild Data Challenge: Predicting Kinship from Faces. *2021 16th IEEE International Conference on Automatic Face and Gesture Recognition (FG 2021)*, 01–07. <https://doi.org/10.1109/FG52635.2021.9666990>
- Salsabila, K. L. N., Handayani, J., & Kusuma, S. Y. (2024). Pengaruh Customer Relationship Management dan Kualitas Pelayanan Terhadap Loyalitas Nasabah dengan Kepuasan Sebagai Variabel Mediasi. *JIBEMA: Jurnal Ilmu Bisnis, Ekonomi, Manajemen, dan Akuntansi*, 2(2), 83–108. <https://doi.org/10.62421/jibema.v2i2.86>
- Samputra, H. (2022). Analisis Deskriptif Statistik Faktor-Faktor Penyebab Kematian Ibu di Provinsi Riau Tahun 2019. *Indonesian Council of Premier Statistical Science*, 1(1), 26–33. <https://doi.org/10.24014/icopss.v1i1.18933>
- Serraoui, I., Laiadi, O., Ouamane, A., Dornaika, F., & Taleb-Ahmed, A. (2022). Knowledge-Based Tensor Subspace Analysis System for Kinship Verification. *Neural Networks*, 151, 222–237. <https://doi.org/10.1016/j.neunet.2022.03.020>
- Shadrikov, A. (2020). Achieving Better Kinship Recognition Through Better Baseline. *2020 15th IEEE International Conference on Automatic Face and Gesture Recognition (FG 2020)*, 872–876. <https://doi.org/10.1109/FG47880.2020.00137>
- Sholihin, M., & Ratmono, D. (2021). *Analisis SEM-PLS dengan WarpPLS 7.0 untuk Hubungan Nonlinier dalam Penelitian Sosial dan Bisnis* (1st ed., Vol. 1). Penerbit Andi.
- Siregar, I. A. (2021). Analisis dan Interpretasi Data Kuantitatif. *ALACRITY: Journal of Education*, 1(2), 39–48. <https://doi.org/10.52121/alacrity.v1i2.25>
- Siroj, R. A., Afgani, W., Fatimah, F., Septaria, D., & Salsabila, G. Z. (2024). Metode Penelitian Kuantitatif Pendekatan Ilmiah untuk Analisis Data. *Jurnal Review Pendidikan dan Pengajaran*, 7(3), 11279–11289. <https://doi.org/10.31004/jrpp.v7i3.32467>
- Sumarni, E., Adawiah, E. R., & Yurna, Y. (2023). Sarana Berpikir Ilmiah (Bahasa, Logika, Matematika dan Statistika). *Jurnal Pendidikan Berkarakter*, 1(4), 106–122. <https://doi.org/10.51903/pendekar.v1i4.299>
- Sun, xianglong, Ning, X., & Junman, D. (2022). Using Combined-TAM-TPB Model to Understand Passenger Acceptance of Driverless Bus—A Case from Suzhou, China. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4098799>
- Tabelessy, W., & Pattiruhu, J. R. (2022). Pengenalan Aplikasi SmartPLS Bagi Mahasiswa Baru Program Studi Magister Manajemen Universitas Pattimura. *COMMUNIO: Jurnal Pengabdian Kepada Masyarakat*, 1(2), 82–88. <https://jurnal.litnuspublisher.com/index.php/jpkm/article/view/25>
- W, H. D., Nugroho, I., & Siahaan, R. (2021). Aplikasi Silsilah Marga Siahaan (Somba Debata) Berbasis Android. *Go Infotech: Jurnal Ilmiah STMIK AUB*, 27(1), 85–96. <https://doi.org/10.36309/goi.v27i1.147>



- Henriksson, H. W., & Goedecke, K. (Eds.). (2021). *Close Relations*. Springer Nature Singapore. <https://doi.org/10.1007/978-981-16-0792-9>
- Wang, M., Shu, X., Feng, J., Wang, X., & Tang, J. (2020). Deep Multi-Person Kinship Matching and Recognition for Family Photos. *Pattern Recognition*, 105, Article ID: 107342. <https://doi.org/10.1016/j.patcog.2020.107342>
- Wang, W., You, S., Karaoglu, S., & Gevers, T. (2023). A Survey on Kinship Verification. *Neurocomputing*, 525, 1–28. <https://doi.org/10.1016/j.neucom.2022.12.031>
- Yamin, S. (2021). *Tutorial Statistik: SPSS, LISREL, WARPPLS & JASP*. Dewangga Energi Internasional.
- Yamin, S. (2023). *Olah Data Statistik: SMARTPLS 3 SMARTPLS 4 AMOS & STATA* (3rd ed.). Dewangga Energi Internasional.
- Yan, H., & Song, C. (2021). Multi-Scale Deep Relational Reasoning for Facial Kinship Verification. *Pattern Recognition*, 110, Article ID: 107541. <https://doi.org/10.1016/j.patcog.2020.107541>
- Yu, J., Li, M., Hao, X., & Xie, G. (2020). Deep Fusion Siamese Network for Automatic Kinship Verification. *2020 15th IEEE International Conference on Automatic Face and Gesture Recognition (FG 2020)*, 892–899. <https://doi.org/10.1109/FG47880.2020.00127>



Algoritma K-Means dan Analisis Komponen Utama untuk Mengatasi Multikolinearitas pada Pengelompokan Kabupaten Tertinggal

Firna Aviliana ^{(1)*}, Putriaji Hendikawati ⁽²⁾

¹ Departemen Matematika, Universitas Negeri Semarang, Semarang, Indonesia

² Departemen Statistika Terapan dan Komputasi, Universitas Negeri Semarang, Semarang, Indonesia

e-mail : aviliana17@students.unnes.ac.id, putriaji.mat@mail.unnes.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 28 Juli 2024, direvisi 7 Oktober 2024, diterima 8 Oktober 2024, dan dipublikasikan 30 September 2025.

Abstract

Underdeveloped areas are regions that frequently face developmental challenges in various aspects such as infrastructure, education, and healthcare. Presidential Regulation Number 63 of 2020 designates 62 regencies in Indonesia as underdeveloped areas. This study categorizes the 62 underdeveloped regencies based on education and health indicators. The methods used are the k-means algorithm and principal component analysis due to multicollinearity in the data. MANOVA is conducted to determine the influence of the cluster results on the Human Development Index (HDI), Average Years of Schooling (AYS), Expected Years of Schooling (EYS), and Life Expectancy (LE). Due to multicollinearity in the education indicator data, principal component analysis was performed, resulting in three main components. The k-means analysis groups the 62 regencies into three clusters based on education indicators and two clusters based on health indicators. Further analysis using MANOVA shows the influence of the education and health clusters on HDI, AYS, EYS, and LE, indicated by statistical test results showing $p\text{-value} < \alpha(0.05)$. Thus, education and health indicators influence the categorization of underdeveloped areas.

Keywords: K-Means Algorithm, Principal Component Analysis, Multicollinearity, MANOVA, Underdeveloped Regencies

Abstrak

Daerah tertinggal merupakan daerah yang sering menghadapi tantangan pembangunan dalam berbagai hal, seperti infrastruktur, pendidikan, dan kesehatan. Peraturan Presiden Nomor 63 tahun 2020 menetapkan 62 kabupaten di Indonesia ke dalam daerah tertinggal. Penelitian ini mengelompokkan 62 kabupaten tertinggal berdasarkan indikator pendidikan dan kesehatan. Metode yang digunakan adalah algoritma k-means dan analisis komponen utama karena terjadi multikolinearitas dalam data. MANOVA dilakukan untuk mengetahui kesesuaian pengaruh hasil *cluster* terhadap nilai Indeks Pembangunan Manusia (IPM), Rata-rata Lama Sekolah (RLS), Harapan Lama Sekolah (HLS), dan Umur Harapan Hidup (UHH). Adanya multikolinearitas dalam data indikator pendidikan sehingga dilakukan analisis komponen utama dan terbentuk tiga komponen utama. Analisis k-means mengelompokkan 62 kabupaten ke dalam tiga *cluster* berdasarkan indikator pendidikan dan dua *cluster* berdasarkan indikator kesehatan. Hasil analisis lanjutan menggunakan MANOVA menunjukkan adanya kesesuaian pengaruh *cluster* pendidikan dan kesehatan terhadap nilai IPM, RLS, HLS, dan UHH ditandai dari hasil uji statistik yang menunjukkan nilai $p\text{-value} < \alpha(0,05)$. Dengan demikian, indikator pendidikan dan kesehatan memberikan pengaruh pada penentuan kelompok pada daerah tertinggal.

Kata Kunci: Algoritma K-Means, Analisis Komponen Utama, Multikolinearitas, MANOVA, Kabupaten Tertinggal

1. PENDAHULUAN

Peningkatan kesejahteraan masyarakat di berbagai wilayah adalah tujuan utama pembangunan nasional. Pembangunan nasional melibatkan banyak pihak, termasuk pemerintah, sektor swasta,



Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

masyarakat sipil, dan organisasi internasional. Telah banyak program dan kebijakan yang diambil untuk meningkatkan kesejahteraan masyarakat di berbagai bidang, seperti ekonomi, pendidikan, kesehatan, dan infrastruktur. Namun, pada kenyataannya masih terdapat ketidaksetaraan atau kesenjangan pembangunan antarwilayah sehingga masih ditemukan kabupaten tertinggal di berbagai provinsi di Indonesia.

Pendidikan dan kesehatan merupakan komponen yang ikut menentukan kesejahteraan masyarakat. Amrullah et al. (2022) menyatakan bahwa pendidikan merupakan salah satu cara untuk meningkatkan kualitas sumber daya manusia, yang diharapkan dapat meningkatkan kesejahteraan manusia. Badan Pusat Statistik (BPS) menyatakan bahwa Angka Partisipasi Murni (APM) dan Angka Partisipasi Kasar (APK) menjadi indikator dalam keberhasilan pendidikan (Prasetyawati, 2023). Sitinjak et al., (2022) mengungkapkan bahwa pembangunan kesehatan merupakan komponen penting dari pembangunan nasional, yang berarti melakukan upaya kesehatan untuk memastikan bahwa setiap orang memiliki kesempatan untuk hidup dengan kesehatan yang optimal. Adanya sarana dan prasarana kesehatan yang baik, seperti puskesmas, rumah sakit, dan fasilitas kesehatan lainnya, serta tenaga kesehatan yang memadai mampu meningkatkan kesehatan masyarakat secara keseluruhan.

Peraturan Presiden (Perpres) Nomor 63 Tahun 2020 Tentang Penetapan Daerah Tertinggal Tahun 2020-2024 menyatakan bahwa daerah tertinggal adalah daerah kabupaten yang wilayah serta masyarakatnya kurang berkembang dibandingkan dengan daerah lain dalam skala nasional. Selanjutnya dalam Pasal 2 menyebutkan kriteria yang menjadi penetapan daerah tertinggal di antaranya adalah perekonomian masyarakat, sumber daya manusia, sarana dan prasarana, kemampuan keuangan daerah, aksesibilitas, dan karakteristik daerah. Peraturan Presiden tersebut menetapkan 62 kabupaten di Indonesia yang termasuk ke dalam daerah tertinggal. Ditinjau dari data BPS tahun 2022, yang tergolong ke dalam daerah tertinggal memiliki Indeks Pembangunan Manusia (IPM) sedang, bahkan beberapa kabupaten tergolong rendah. IPM didasarkan pada tiga dimensi utama, yaitu umur panjang dan hidup sehat, pengetahuan, dan standar hidup yang layak (Setiawan et al., 2022). Umur panjang dan hidup sehat sangat dipengaruhi oleh kualitas kesehatan yang dicerminkan oleh Umur Harapan Hidup (UHH), pengetahuan dipengaruhi oleh kualitas pendidikan, yang dicerminkan oleh Rata-rata Lama Sekolah (RLS) dan Harapan Lama Sekolah (HLS), sedangkan standar hidup layak dipengaruhi oleh kualitas ketenagakerjaan pada daerah tersebut yang dicerminkan oleh Pendapatan Nasional Bruto (PNB) per kapita (PPP).

Pendidikan dan kesehatan memainkan peran penting dalam peningkatan IPM serta mewujudkan tercapainya kesejahteraan masyarakat. Indikator-indikator seperti fasilitas, akses, serta pelayanan dalam bidang pendidikan dan kesehatan menjadi penentu untuk mengukur kualitas hidup masyarakat. Kabupaten yang masih dalam kategori daerah tertinggal perlu dikelompokkan berdasarkan indikator pendidikan dan kesehatan agar proses pembangunan terlaksana tepat sasaran. Hasil pengelompokan ini berguna untuk mengidentifikasi kebutuhan khusus dari setiap kelompok kabupaten, terutama dalam bidang pendidikan dan kesehatan. Hasil pengelompokan menjadi alat yang penting dalam upaya meningkatkan kesejahteraan masyarakat di kabupaten tertinggal agar tercapai pemerataan pembangunan nasional.

Analisis *cluster* merupakan teknik mengelompokkan objek atau data ke dalam *cluster* berdasarkan kesamaan dan kemiripan karakteristiknya. Dalam satu *cluster* memiliki karakteristik yang sama dan memiliki perbedaan karakteristik pada objek *cluster* lain (Amrullah et al., 2022). Salah satu metode untuk mengelompokkan data yaitu menggunakan algoritma k-means. Algoritma k-means bekerja dengan mengelompokkan data ke dalam *k* kelompok (*cluster*) yang ditentukan sebelumnya, dengan meminimalkan variasi atau jarak antara titik data dalam satu *cluster* (Aryanto et al., 2024). K-means merupakan metode non-hirarki yang paling banyak digunakan untuk pengelompokan karena kemudahannya dan kemampuannya untuk mengelompokkan dalam jumlah data yang besar dengan waktu komputasi yang cepat dan efisien (Lenama et al., 2023). Algoritma k-means merupakan metode yang mudah dipahami dan diimplementasikan, serta mampu menghasilkan hasil yang cukup akurat.



Beberapa penelitian terdahulu yang melakukan pengelompokan menggunakan algoritma k-means, di antaranya Amrullah et al. (2022) mengenai pengelompokan faktor penunjang pendidikan di Kabupaten Karawang, diperoleh hasil dua *cluster* optimal dengan nilai Davies-Bouldin Index (DBI) 0,408. Penelitian lain dilakukan oleh Octaviyani et al. (2022) mengenai pengelompokan gizi balita, diperoleh lima *cluster* optimal dengan nilai DBI 0,5224. Yudhistira & Andika (2023) juga melakukan penelitian pengelompokan data nilai siswa menggunakan k-means, diperoleh tiga *cluster* terbaik dengan nilai *silhouette coefficient* 0,489.

Setelah diketahui kategori pengelompokan daerah tertinggal tersebut, selanjutnya dilakukan uji kesesuaian pengaruh pengelompokan tersebut terhadap variabel utama pembangunan nasional, yaitu IPM, RLS, HLS, dan UHH menggunakan MANOVA. MANOVA adalah uji statistik untuk mengukur pengaruh variabel independen (variabel bebas) pada skala kategori terhadap variabel dependen (variabel terikat) pada skala data kuantitatif (Pursitasari et al., 2024). Tujuan penelitian ini adalah menguji hubungan hasil *cluster* terhadap nilai IPM, RLS, HLS, dan UHH. Apakah pengelompokan berdasarkan indikator pendidikan dan kesehatan mempengaruhi berbagai aspek pembangunan manusia di daerah tertinggal sehingga hasil *cluster* yang diberikan dapat mendukung pembuatan kebijakan yang lebih efektif dan tepat sasaran.

2. METODE PENELITIAN

Penelitian ini menggunakan analisis *cluster* dengan algoritma k-means untuk memperoleh hasil *cluster* kabupaten tertinggal berdasarkan indikator pendidikan dan kesehatan. Selanjutnya hasil *cluster* tersebut dianalisis lanjut menggunakan MANOVA untuk menguji apakah terdapat pengaruh yang signifikan dari hasil *cluster* yang terbentuk terhadap Indeks Pembangunan Manusia (IPM), Rata-rata Lama Sekolah (RLS), Harapan Lama Sekolah, dan Umur Harapan Hidup (UHH). Analisis MANOVA dilakukan untuk mengetahui kategori hasil *cluster* yang diperoleh memberikan kesesuaian pengaruh pada tingkat nilai IPM, RLS, HLS, dan UHH.

2.1 Data dan Sumber Data

Data yang digunakan dalam penelitian ini adalah data sekunder yang diambil dari *website* Badan Pusat Statistik (BPS) melalui laman <https://www.bps.go.id>, meliputi data indikator pendidikan dan kesehatan di 62 kabupaten tertinggal di Indonesia pada tahun 2022. Pada analisis *cluster*, data indikator pendidikan memuat 14 variabel, yaitu jumlah TK/RA (X_1), jumlah SD/MI (X_2), jumlah SMP/MTs (X_3), jumlah SMA/SMK/MA (X_4), jumlah guru TK/RA (X_5), jumlah guru SD/MI (X_6), jumlah guru SMP/MTs (X_7), jumlah guru SMA/SMK/MA (X_8), APM SD/MI (X_9), APM SMP/MTs (X_{10}), APM SMA/SMK/MA (X_{11}), APK SD/MI (X_{12}), APK SMP/MTs (X_{13}), dan APK SMA/SMK/MA (X_{14}), sedangkan indikator kesehatan memuat 5 variabel, yaitu jumlah dokter (X_{15}), jumlah perawat (X_{16}), jumlah bidan (X_{17}), jumlah rumah sakit (X_{18}), jumlah puskesmas (X_{19}). Analisis MANOVA dilakukan setelah diperoleh hasil analisis *cluster* dua indikator yang masing-masing memuat variabel bebas dan terikat. Pada indikator pendidikan, variabel bebas yaitu hasil *cluster* pendidikan (X_{20}), sedangkan variabel terikat meliputi IPM (Y_1), RLS (Y_2), dan HLS (Y_3). Pada indikator kesehatan, variabel bebas yaitu hasil *cluster* kesehatan (X_{21}), sedangkan variabel terikat meliputi IPM (Y_4) dan UHH (Y_5).

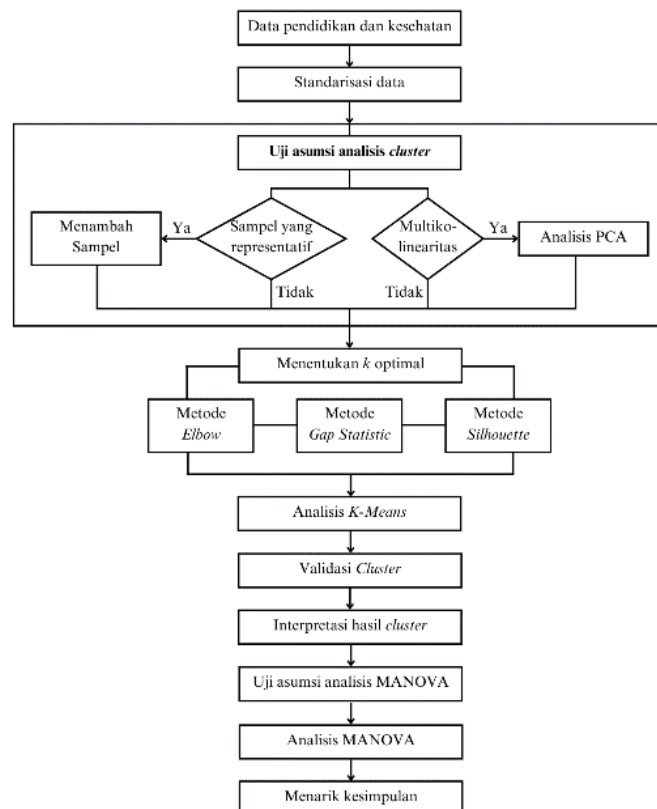
2.2 Tahapan Analisis Data

Analisis data dalam penelitian ini dilakukan dengan bantuan software RStudio. Tahapan analisis dapat dilihat pada Gambar 1 yang dimulai dengan pengumpulan data yang bersumber dari publikasi BPS. Setelah itu, dilakukan standarisasi data agar semua variabel memiliki skala yang sama dan tidak menimbulkan dominansi antarvariabel, mengingat data pendidikan dan kesehatan memiliki satuan yang berbeda. Standarisasi ini bertujuan menyamakan distribusi variabel sehingga lebih seragam (Wijaya et al., 2024). Selanjutnya, dilakukan pengujian asumsi *cluster* yang mencakup representativitas sampel dan uji nonmultikolinearitas. Apabila ditemukan multikolinearitas, maka dilakukan analisis komponen utama atau Principal Component Analysis (PCA).



Tahap berikutnya adalah menentukan jumlah cluster optimal (k) menggunakan tiga metode, yaitu metode elbow, metode gap statistic, dan metode silhouette. Setelah jumlah cluster optimal diperoleh, analisis dilanjutkan menggunakan algoritma k-means. Validasi hasil clustering kemudian dilakukan dengan menggunakan Silhouette Index (SI) untuk memastikan kualitas pemisahan antarcluster. Selanjutnya, hasil cluster diinterpretasikan secara menyeluruh.

Tahapan berikutnya melibatkan pengujian asumsi MANOVA, termasuk uji normalitas multivariat dan uji homoskedastisitas. Setelah asumsi terpenuhi, analisis dilakukan menggunakan MANOVA untuk menguji perbedaan antarcluster berdasarkan variabel penelitian. Tahapan terakhir adalah menarik kesimpulan dari seluruh rangkaian analisis yang telah dilakukan.



Gambar 1 Tahapan Penelitian

3. HASIL DAN PEMBAHASAN

Penerapan algoritma k-means dengan analisis komponen utama dan dilanjut analisis MANOVA dalam pengelompokan kabupaten tertinggal berdasarkan indikator pendidikan dan kesehatan dapat membantu meningkatkan efektivitas upaya peningkatan pendidikan dan kesehatan di kabupaten tertinggal. Tahapan dalam analisis dapat diuraikan sebagai berikut.

3.1 Analisis K-Means

Sebelum melakukan analisis *cluster*, uji asumsi *cluster* dilakukan pada masing-masing data indikator agar mendapatkan hasil yang lebih akurat. Hair et al. (2013, dikutip dalam Hamidah et al., 2022) menyebutkan uji asumsi analisis *cluster* meliputi sampel yang representatif dan tidak terjadi multikolinearitas.

- 1) Pengecekan sampel yang representatif dilakukan menggunakan uji KMO (*Kaiser Meyer Olkin*). Menurut Kaiser (1974, dikutip dalam Hamidah et al., 2022), sampel dikatakan mewakili populasi jika nilai KMO lebih dari 0,5. Dalam pengujian ini menunjukkan hasil KMO data

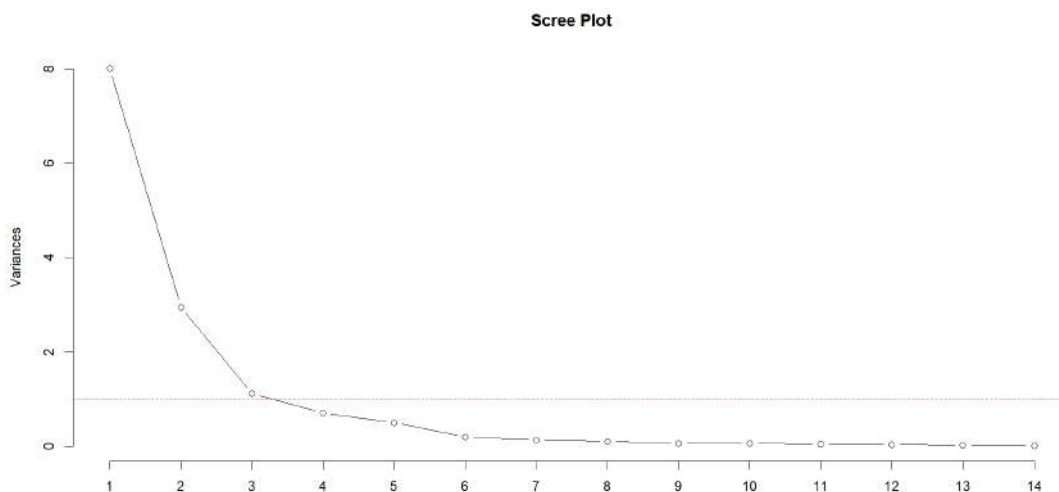


indikator pendidikan sebesar 0,82, sedangkan data indikator kesehatan memiliki nilai KMO sebesar 0,8. Dengan demikian kedua data indikator tersebut memenuhi syarat sampel yang representatif.

- 2) Uji multikolinearitas dilakukan dengan penghitungan nilai VIF antar variabel. Dikatakan memenuhi syarat multikolinearitas apabila nilai VIF lebih kecil dari 10 atau dinotasikan sebagai $VIF < 10$ (Fajar & Indrawati, 2020). Hasil dari penghitungan nilai VIF pada data indikator pendidikan dapat dilihat dalam Tabel 7 di Lampiran A dan data indikator kesehatan dalam Tabel 1. Pada Tabel 7 di Lampiran A diketahui terdapat nilai $VIF > 10$ maka data indikator pendidikan terindikasi multikolinearitas. Jika suatu data terjadi multikolinearitas menandakan adanya korelasi yang kuat antar variabel sehingga perlu dilakukan analisis komponen utama atau *Principal Component Analysis* (PCA) (Thamrin & Murni, 2022). Pada Tabel 1 memiliki nilai $VIF < 10$, artinya data indikator kesehatan tidak terjadi multikolinearitas.

Tabel 1 Nilai VIF Antar Variabel Kesehatan

Nilai VIF	X_{15}	X_{16}	X_{17}	X_{18}	X_{19}
X_{15}		2,735	2,153	2,939	2,941
X_{16}	2,992		2,586	3,152	2,917
X_{17}	2,552	2,801		3,490	3,495
X_{18}	1,129	1,107	1,132		1,115
X_{19}	1,433	1,298	1,436	1,414	

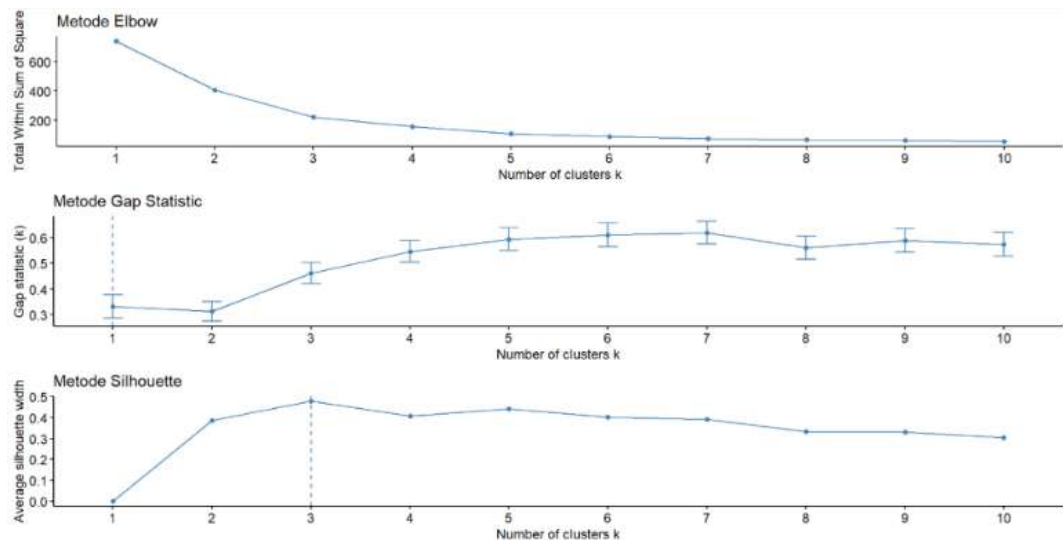


Gambar 2 Scree Plot

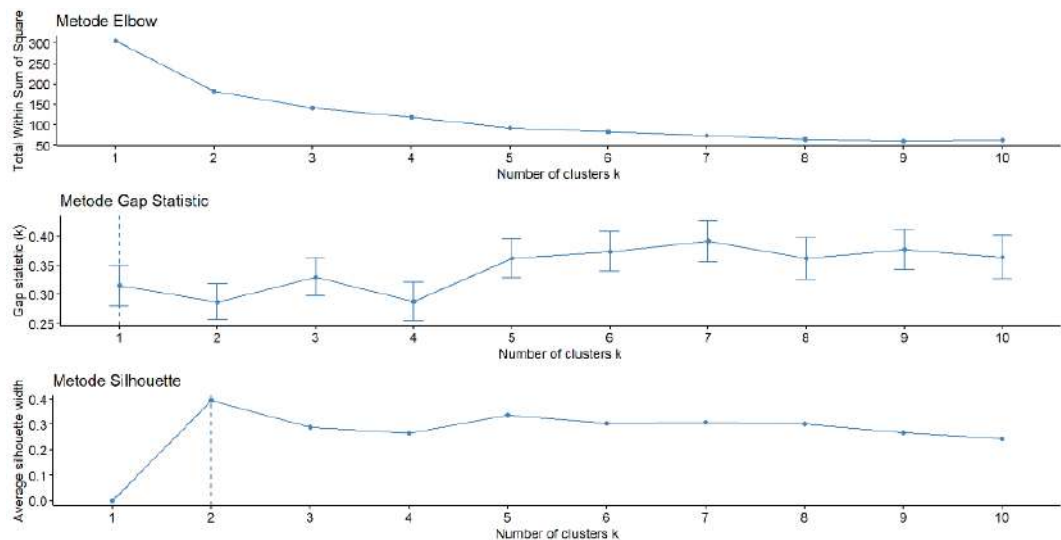
Analisis komponen utama atau *Principal Component Analysis* (PCA) adalah suatu cara yang digunakan untuk mereduksi variabel yang banyak dan mengatasi data yang terdapat multikolinearitas (Amelia & Kholijah, 2023). Menurut Mariana (2013) ada tiga kriteria pemilihan k komponen utama yaitu memilih akar ciri (nilai eigen) yang lebih dari satu, nilai proporsi varians kumulatif minimal 80%, dan menggunakan *scree plot*. Data indikator pendidikan mengalami multikolinearitas sehingga dilakukan analisis komponen utama terlebih dahulu. Pencarian nilai eigen dan proporsi varians kumulatif data indikator pendidikan dapat dilihat pada Tabel 8 di Lampiran A yang menunjukkan terdapat tiga komponen utama yang memiliki nilai eigen lebih dari satu serta PC 3 menunjukkan proporsi kumulatif sebesar 0,8634 artinya data sudah mampu menjelaskan sebesar 86,33% dari keragaman data aslinya sehingga menurut kriteria ini dipilih tiga komponen utama karena telah melebihi 80%. Berdasarkan *scree plot* pada Gambar 2 menunjukkan titik tiga garis *scree plot* mulai melandai atau bisa dilihat dari garis horizontal putus-putus, yakni terdapat tiga titik yang berada di atas garis sehingga dipilih komponen utama sebanyak tiga. Ketiga kriteria pemilihan menunjukkan hasil yang sama, yaitu tiga komponen



utama sehingga tiga komponen utama tersebut yang akan digunakan dalam analisis *cluster* menggunakan k-means.



Gambar 3 Plot Penentuan k Optimal Indikator Pendidikan

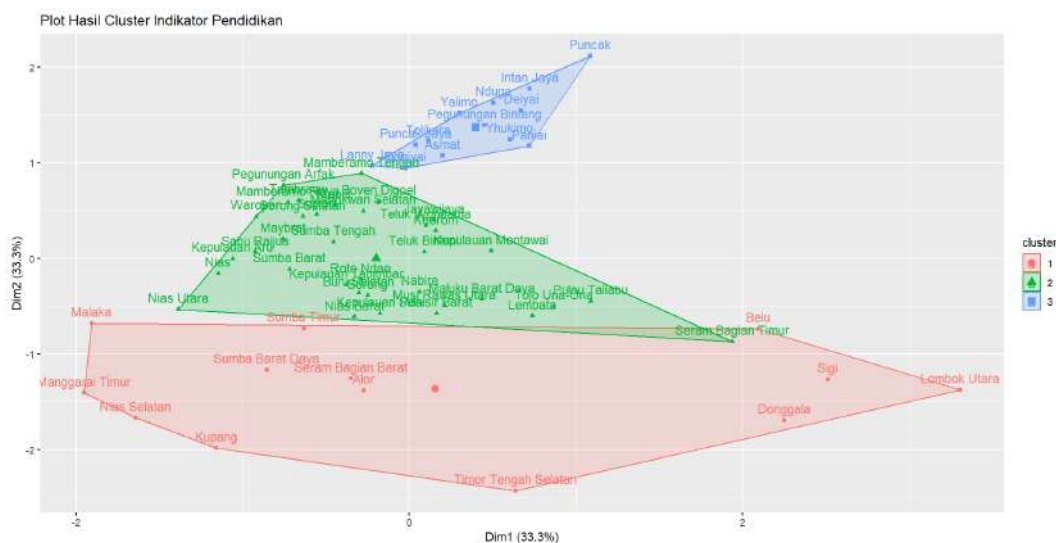


Gambar 4 Plot Penentuan k Optimal Indikator Kesehatan

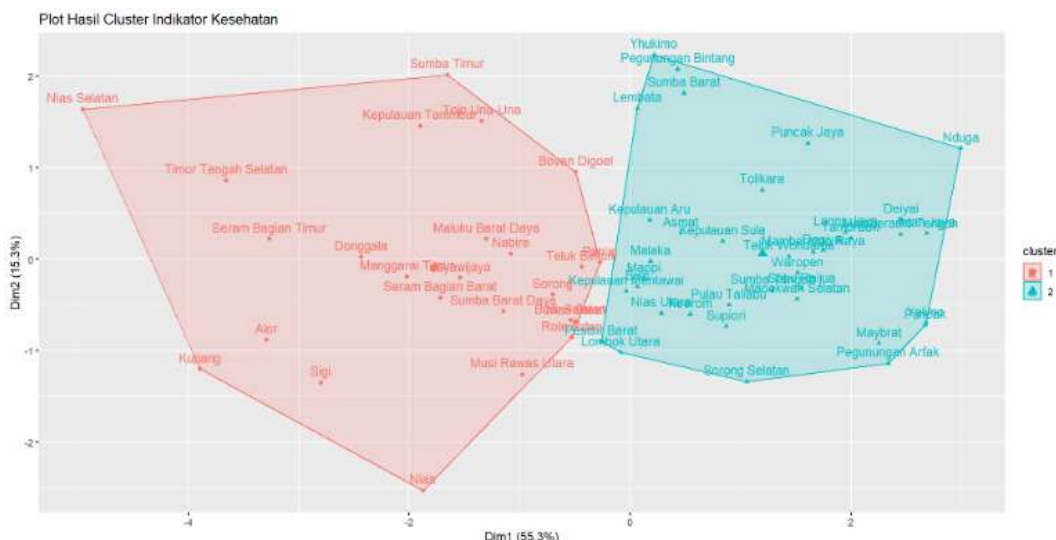
Penentuan k optimal dapat dilakukan dengan metode *elbow*, metode *gap statistic*, dan metode *silhouette* (Yuan & Yang, 2019). Metode *elbow* dijelaskan bahwa untuk menentukan k optimal dilihat pada persentase grafik yang membentuk sudut siku (Lashiyanti et al., 2023). Metode *gap statistic* dilihat dari garis yang dihubungkan dengan garis putus-putus (Zulyani et al., 2023). Metode *silhouette* menunjukkan nilai yang paling baik dilihat pada presentasi garis yang paling tinggi (Lashiyanti et al., 2023). Gambar 3 menunjukkan plot penentuan k optimal pada data indikator pendidikan, yaitu nilai k optimal pada metode *elbow* menunjukkan tiga *cluster*, metode *gap statistic* menunjukkan satu *cluster*, sedangkan metode *silhouette* menunjukkan tiga *cluster*. Dua metode menunjukkan hasil yang sama sehingga tiga dipilih sebagai banyaknya *cluster* optimal. Gambar 4 menunjukkan plot penentuan k optimal pada data indikator kesehatan, yaitu pada metode *elbow* menunjukkan dua *cluster*, metode *gap statistic* menunjukkan satu *cluster*, sedangkan metode *silhouette* menunjukkan dua *cluster*. Dua metode menunjukkan hasil yang sama sehingga dua dipilih sebagai banyaknya *cluster* optimal.



Bholowalia et al. (Faujia & Subarkah, 2022) menjelaskan bahwa algoritma k-means diawali dengan penentuan jumlah *cluster* yang akan terbentuk dan mendefinisikan nilai inisiasi *centroid* (pusat *cluster*). Dengan menggunakan jarak terdekat dari *centroid*, k-means mengelompokkan data hanya pada satu *cluster* sehingga satu data tidak akan menempati dua *cluster* atau lebih. Hasil pembentukan k-means divisualisasikan menggunakan Rstudio yang dapat dilihat pada Gambar 5 dan 6. Gambar 5 menunjukkan hasil visualisasi *cluster* pendidikan yang mengelompokkan 62 kabupaten tertinggal ke dalam tiga *cluster*, yakni *cluster 1* terdiri dari 13 kabupaten, *cluster 2* terdiri dari 36 kabupaten, dan *cluster 3* terdiri dari 13 kabupaten. Gambar 6 menunjukkan hasil visualisasi *cluster* kesehatan yang mengelompokkan 62 kabupaten tertinggal ke dalam dua *cluster*, yakni *cluster 1* terdiri dari 25 kabupaten dan *cluster 2* terdiri dari 37 kabupaten.



Gambar 5 Visualisasi *Cluster* Pendidikan



Gambar 6 Visualisasi *Cluster* Kesehatan

Tahap validasi *cluster* yang telah terbentuk dilakukan menggunakan Silhouette Index (SI). Silhouette Index akan mengevaluasi penempatan setiap objek dalam setiap *cluster* dengan membandingkan jarak rata-rata objek dalam satu *cluster* dan jarak antara objek dengan *cluster* berbeda. Semakin besar nilai koefisien *silhouette* akan semakin baik kualitas suatu kelompok



(Nahdliyah et al., 2019). Hasil validasi *cluster* menggunakan *silhouette index* dapat dilihat pada Tabel 2. Hasil pengelompokan terbaik data pendidikan menunjukkan 3 *cluster* dengan nilai *silhouette index* sebesar 0,3696. Hal ini sesuai dengan penentuan banyaknya *k* optimal menggunakan metode *elbow* dan metode *silhouette* pada tahap sebelumnya sehingga pengelompokan indikator pendidikan ke dalam 3 *cluster* memiliki kualitas yang baik dan kuat. Namun, terdapat perbedaan pada indikator kesehatan yang mana pada validasi *silhouette index* menunjukkan banyaknya *cluster* terbaik 4, sedangkan dengan menggunakan metode *elbow* dan metode *silhouette* pada tahap sebelumnya menunjukkan 2 *cluster*. Namun demikian, pengelompokan 2 *cluster* memiliki nilai *silhouette index* sebesar 0,2351 menduduki urutan ke dua setelah pengelompokan 4 *cluster* yang memiliki nilai *silhouette index* sebesar 0,2491. Dengan demikian, pengelompokan 2 *cluster* pada indikator kesehatan juga memiliki kualitas pengelompokan yang baik dan kuat.

Tabel 2 Validasi *Cluster* Pendidikan dan Kesehatan

Banyak Cluster	Indikator Pendidikan	Indikator Kesehatan
2	0,3578	0,2351
3	0,3696	0,2009
4	0,3293	0,2491
5	0,3437	0,1733

Karakteristik masing-masing *cluster* ditandai dengan besaran rata-rata tiap *cluster* sehingga diketahui kategori tiap *cluster*-nya. Hasil rata-rata tiap *cluster* ditunjukkan pada Tabel 3 untuk *cluster* pendidikan dan Tabel 4 untuk *cluster* kesehatan. Dengan demikian diperoleh rincian sebagai berikut.

- 1) Hasil analisis *cluster* pendidikan menggunakan metode k-means dengan tiga *cluster* diperoleh hasil sebagai berikut:
 - a) *Cluster* 1 terdiri dari 13 kabupaten, yaitu Nias Selatan, Lombok Utara, Sumba Timur, Kupang, Timor Tengah Selatan, Belu, Alor, Sumba Barat Daya, Manggarai Timur, Malaka, Donggala, Sigi, dan Seram Bagian Barat. Anggota *cluster* 1 memiliki indikator pendidikan yang lebih tinggi atau memadai.
 - b) *Cluster* 2 terdiri dari 36 kabupaten, yaitu Nias, Nias Utara, Nias Barat, Kepulauan Mentawai, Musi Rawas Utara, Pesisir Barat, Sumba Barat, Lembata, Rote Ndao, Sumba Tengah, Sabu Raijua, Tojo Una-Una, Kepulauan Tanimbar, Kepulauan Aru, Seram Bagian Timur, Maluku Barat Daya, Buru Selatan, Kepulauan Sula, Pulau Taliabu, Teluk Wondama, Teluk Bintuni, Sorong Selatan, Sorong, Tambrauw, Maybrat, Manokwari Selatan, Pegunungan Arfak, Jayawijaya, Nabire, Boven Digoel, Mappi, Keerom, Waropen, Supiori, Mamberamo Raya, dan Mamberamo Tengah. Anggota *cluster* 2 memiliki indikator pendidikan yang sedang atau cukup memadai.
 - c) *Cluster* 3 terdiri dari 13 kabupaten, yaitu Paniai, Puncak Jaya, Asmat, Yahukimo, Pegunungan Bintang, Tolikara, Nduga, Lanny Jaya, Yalimo, Puncak, Dogiyai, Intan Jaya, dan Deiyai. Anggota *cluster* 3 memiliki indikator pendidikan yang rendah atau terbatas sehingga perlu perhatian khusus dari pemerintah.
- 2) Hasil analisis *cluster* kesehatan menggunakan metode k-means dengan dua *cluster* diperoleh hasil sebagai berikut:
 - a) *Cluster* 1 terdiri dari 25 kabupaten, yaitu Nias, Nias Selatan, Nias Barat, Musi Rawas Utara, Sumba Timur, Kupang, Timor Tengah Selatan, Alor, Rote Ndao, Sumba Barat Daya, Manggarai Timur, Donggala, Tojo Una-Una, Sigi, Kepulauan Tanimbar, Seram Bagian Barat, Seram Bagian Timur, Maluku Barat Daya, Buru Selatan, Teluk Bintuni, Sorong, Jayawijaya, Nabire, Paniai, Boven Digoel. Anggota *cluster* 1 memiliki jumlah sarana prasarana dan tenaga medis yang memadai sehingga memiliki indikator kesehatan yang tinggi.
 - b) *Cluster* 2 terdiri dari 37 kabupaten, yaitu Nias Utara, Kepulauan Mentawai, Pesisir Barat, Lombok Utara, Sumba Barat, Belu, Lembata, Sumba Tengah, Sabu Raijua, Malaka, Kepulauan Aru, Kepulauan Sula, Pulau Taliabu, Teluk Wondama, Sorong Selatan, Tambrauw, Maybrat, Manokwari Selatan, Pegunungan Arfak, Puncak Jaya, Mappi, Asmat, Yahukimo, Pegunungan Bintang, Tolikara, Keerom, Waropen,



Supiori, Mamberamo Raya, Nduga, Lannya Jaya, Mamberamo Tengah, Yalimo, Puncak, Dogiyai, Intan Jaya, Deiyai. Anggota *cluster* 2 memiliki jumlah sarana prasarana dan tenaga medis yang terbatas sehingga memiliki indikator kesehatan yang rendah.

Tabel 3 Karakteristik *Cluster* Pendidikan

Variabel	Cluster 1	Cluster 2	Cluster 3
Jumlah Sekolah TK/RA	126,615385	52,444444	8,923077
Jumlah Sekolah SD/MI	297,8462	110,3056	71,0000
Jumlah Sekolah SMP/MTs	117,53846	40,52778	17,15385
Jumlah Sekolah SMA/SMK/MA	62,15385	19,86111	6,00000
Jumlah Guru TK/RA	385,53846	162,83333	28,69231
Jumlah Guru SD/MI	1949,9231	1001,3611	394,3846
Jumlah Guru SMP/MTs	1778,6923	584,7778	181,4615
Jumlah Guru SMA/SMK/MA	1414,69231	391,50000	93,23077
APM SD/MI (%)	94,54077	94,21778	68,21308
APM SMP/MTs (%)	70,65769	70,94861	48,34615
APM SMA/SMK/MA (%)	56,00000	56,92278	28,39769
APK SD/MI (%)	111,07769	109,57000	81,60308
APK SMP/MTs (%)	89,78077	93,45306	69,65769
APK SMA/SMK/MA (%)	81,35000	87,29694	46,66231

Tabel 4 Karakteristik *Cluster* Kesehatan

Variabel	Cluster 1	Cluster 2
Jumlah Dokter	48,88000	19,72973
Jumlah Perawat	475,880	174,973
Jumlah Bidan	370,5600	119,1351
Jumlah Rumah Sakit Umum	1,440000	1,054054
Jumlah Puskesmas	21,52000	13,43243

3.2 Analisis MANOVA

Setelah dilakukan analisis *cluster*, selanjutnya akan dilakukan analisis MANOVA dengan tujuan menguji apakah terdapat kesesuaian pengaruh dari hasil *cluster* yang terbentuk terhadap nilai IPM, RLS, HLS, dan UHH. Namun, sebelum melakukan analisis MANOVA harus memenuhi syarat uji asumsi MANOVA. Terdapat dua asumsi yang harus dipenuhi, yaitu normalitas multivariat dan homoskedastisitas (Hamidah et al., 2022). Pengujian asumsi ini dilakukan untuk memastikan bahwa MANOVA dapat digunakan dan menghasilkan keputusan yang tepat (Sutrisno & Wulandari, 2018).

a) Uji normalitas multivariat dengan hipotesis sebagai berikut:

H_0 : Data berdistribusi normal multivariat

H_1 : Data tidak berdistribusi normal multivariat

Pengujian normalitas multivariat dilakukan menggunakan uji Mardia menunjukkan $p - value > \alpha(0,05)$ pada pengujian hasil *cluster* pendidikan terhadap IPM, RLS, dan HLS, serta pengujian hasil *cluster* kesehatan terhadap IPM dan UHH. Dengan demikian, keduanya memiliki keputusan terima H_0 sehingga data berdistribusi normal multivariat. Oleh karena itu, uji asumsi normalitas multivariat terpenuhi.

b) Uji homoskedastisitas dengan hipotesis sebagai berikut:

H_0 : Data bersifat homoskedastisitas

H_1 : Data tidak bersifat homoskedastisitas

Pengujian homoskedastisitas dilakukan menggunakan uji Levene menunjukkan $p - value(0,3332) > \alpha(0,05)$ pada pengujian hasil *cluster* pendidikan terhadap IPM, RLS, dan HLS, sedangkan pada pengujian hasil *cluster* kesehatan terhadap IPM dan UHH memiliki $p - value(0,03473) < \alpha(0,05)$. Dengan demikian, data pendidikan memiliki keputusan terima H_0 sehingga data bersifat homoskedastisitas, tetapi data kesehatan memiliki keputusan tolak H_0 .



sehingga data tidak bersifat homoskedastisitas. Menurut Hamidah et al. (2022), jika terdapat uji asumsi yang tidak terpenuhi, maka dapat menggunakan uji statistik Wilk's lambda karena kriteria pada uji tersebut dianggap lebih *robust* dibanding kriteria pada uji lainnya. Dengan demikian, analisis MANOVA dapat dilanjutkan.

Analisis MANOVA pada data pendidikan menggunakan empat uji statistik, yaitu Pillai's trace, Wilk's lambda, Hotelling trace, dan Roy's largest root, sedangkan data kesehatan menggunakan uji statistik Wilk's lambda karena terdapat salah satu uji asumsi yang tidak terpenuhi. Hipotesis penelitian MANOVA yang digunakan sebagai berikut.

H_0 : Hasil *cluster* tidak berpengaruh signifikan secara multivariat

H_1 : Hasil *cluster* berpengaruh signifikan secara multivariat

Tabel 5 menunjukkan uji statistik pengaruh *cluster* pendidikan terhadap nilai IPM, RLS, dan HLS. Hasil uji statistik menyatakan nilai $p - value < \alpha(0,05)$ yang berarti tolak H_0 . Dengan demikian, dapat disimpulkan bahwa terdapat kesesuaian pengaruh yang signifikan antara hasil *cluster* (*cluster* 1, *cluster* 2, dan *cluster* 3) terhadap nilai IPM, RLS, dan HLS. *Cluster* dengan kategori tinggi memiliki nilai IPM, RLS, dan HLS yang tinggi juga, begitupun sebaliknya.

Tabel 6 menunjukkan uji statistik pengaruh *cluster* kesehatan terhadap nilai IPM dan UHH. Hasil uji *Wilk's lambda* menunjukkan $p - value < \alpha(0,05)$ yang berarti tolak H_0 . Dengan demikian, dapat disimpulkan bahwa terdapat pengaruh yang signifikan antara hasil *cluster* (*cluster* 1 dan *cluster* 2) terhadap nilai IPM dan UHH. *Cluster* dengan kategori tinggi memiliki nilai IPM dan UHH yang tinggi juga, begitupun sebaliknya.

Hasil *cluster* pendidikan dan kesehatan memang secara nyata berpengaruh pada IPM, RLS, HLS, dan UHH. Antara lain ditunjukkan pada Kabupaten Nias Selatan yang tergolong kategori tinggi pada indikator pendidikan dan kesehatan. Kabupaten Nias Selatan memiliki nilai IPM yaitu 63,17, nilai RLS 6,23, nilai HLS 12,48, dan nilai UHH 69,21. Dibandingkan dengan Kabupaten Puncak Jaya yang tergolong dalam kategori rendah di kedua indikator tersebut memiliki nilai IPM 49,84, nilai RLS 4,03, nilai HLS 7,50, dan nilai UHH 65,66. Terbukti bahwa kabupaten dengan *cluster* tinggi akan memiliki nilai IPM, RLS, HLS, dan UHH yang tinggi dibandingkan dengan *cluster* rendah. Dengan demikian kabupaten-kabupaten yang termasuk dalam kategori rendah, baik indikator pendidikan maupun kesehatan harus mendapatkan perhatian lebih dari pemerintah agar tercipta peningkatan kesejahteraan masyarakat.

Tabel 5 Uji Statistik MANOVA *Cluster* Pendidikan

Uji Statistik	p-value	Keputusan	Keterangan
Pillai's trace	1,367e-09	Tolak H_0	Berbeda signifikan
Wilk's lambda	1,367e-09	Tolak H_0	Berbeda signifikan
Hotelling trace	1,367e-09	Tolak H_0	Berbeda signifikan
Roy	1,367e-09	Tolak H_0	Berbeda signifikan

Tabel 6 Uji Statistik MANOVA *Cluster* Kesehatan

Uji Statistik	p-value	Keputusan	Keterangan
Pillai's trace	0,000282	Tolak H_0	Berbeda signifikan

4. KESIMPULAN

Berdasarkan hasil penelitian dan pembahasan dapat disimpulkan bahwa algoritma k-means dapat dilakukan untuk mengelompokkan kabupaten tertinggal berdasarkan data indikator pendidikan dan kesehatan. Pengelompokan 62 kabupaten tertinggal berdasarkan indikator pendidikan terbagi menjadi tiga *cluster* dengan nilai *silhouette index* sebesar 0,3696. *Cluster* 1 mencakup 13 kabupaten dengan sarana prasarana dan tenaga pendidik yang memadai sehingga memiliki indikator pendidikan yang tinggi. *Cluster* 2 mencakup 36 kabupaten dengan jumlah sarana prasarana dan tenaga pendidik yang cukup memadai sehingga memiliki indikator



pendidikan yang sedang. *Cluster* 3 mencakup 13 kabupaten dengan sarana prasarana dan tenaga pendidik yang terbatas sehingga memiliki indikator pendidikan yang rendah. Hasil analisis MANOVA menunjukkan adanya pengaruh hasil *cluster* pendidikan terhadap IPM, RLS, dan HLS.

Sementara indikator kesehatan terbagi ke dalam dua *cluster dengan nilai silhouette index sebesar 0,2351*. *Cluster* 1 mencakup 25 kabupaten dengan jumlah sarana prasarana dan tenaga medis yang memadai sehingga memiliki indikator kesehatan yang tinggi. *Cluster* 2 mencakup 37 kabupaten dengan jumlah sarana prasarana kesehatan dan tenaga medis yang terbatas sehingga memiliki indikator kesehatan yang rendah. Hasil analisis MANOVA menunjukkan adanya pengaruh hasil *cluster* kesehatan terhadap IPM dan UHH. Dengan kata lain masing-masing pertumbuhan bidang pendidikan dan kesehatan memiliki pengaruh yang signifikan terhadap Indeks Pembangunan Manusia, Rata-rata Lama Sekolah (RLS), Harapan Lama Sekolah (HLS), dan Umur Harapan Hidup (UHH). Hasil penelitian ini sejalan dengan penelitian yang dilakukan oleh Hasibuan et al. (2023).

DAFTAR PUSTAKA

- Amelia, D., & Kholijah, G. (2023). Analisis Cluster Pengelompokan Provinsi di Indonesia Berdasarkan Sub Sektor Nilai Tukar Petani. *DEMOS: Journal of Demography, Ethnography and Social Transformation*, 3(1), 1–12. <https://doi.org/10.30631/demos.v3i1.1812>
- Amrullah, A., Purnamasari, I., Sari, B. N., Garno, G., & Voutama, A. (2022). Analisis Cluster Faktor Penunjang Pendidikan Menggunakan Algoritma K-Means (Studi Kasus: Kabupaten Karawang). *Jurnal Informatika dan Rekayasa Elektronik*, 5(2), 244–252. <https://doi.org/10.36595/jire.v5i2.701>
- Aryanto, R. P., Nilogiri, A., & Wardoyo, A. E. (2024). Klasterisasi Jumlah Penduduk Provinsi Jawa Timur Tahun 2021-2023 Menggunakan Algoritma K-Means. *JISKA (Jurnal Informatika Sunan Kalijaga)*, 9(2), 134–146. <https://doi.org/10.14421/jiska.2024.9.2.134-146>
- Fajar, M. A., & Indrawati, L. (2020). Pengaruh Belanja Pendidikan, Belanja Kesehatan dan Belanja Perumahan dan Fasilitas Umum Terhadap Indeks Pembangunan Manusia (Studi Kasus pada Pemerintah Daerah Kabupaten Cianjur). *Indonesian Accounting Research Journal*, 1(1), 108–118. <https://jurnal.polban.ac.id/iarj/article/view/2366>
- Faujia, R. A., & Subarkah, M. Z. (2022). Analisis Klaster K-Means dan Visualisasi Data Spasial Berdasarkan Karakteristik Persebaran Covid-19 dan Pelanggaran Protokol Kesehatan di Jawa Tengah. *Seminar Nasional Official Statistics*, 2022(1), 813–822. <https://doi.org/10.34123/semnasoffstat.v2022i1.1222>
- Hamidah, N., Santoso, R., & Rusgiyono, A. (2022). Klasterisasi Provinsi di Indonesia Berdasarkan Faktor Penyebaran Covid-19 Menggunakan Model-Based Clustering t-Multivariat. *Jurnal Gaussian*, 11(1), 56–66. <https://doi.org/10.14710/j.gauss.v11i1.33999>
- Hasibuan, S. R., Harahap, I., & Tambunan, K. (2023). Pengaruh Pertumbuhan Ekonomi, Pendidikan dan Kesehatan Terhadap Indeks Pembangunan Manusia di Provinsi Sumatera Utara. *Jurnal Manajemen Akuntansi (JUMSI)*, 3(1), 272–285. <https://doi.org/10.36987/jumsi.v3i1.4023>
- Lashiyanti, A. R., Munthe, I. R., & Nasution, F. A. (2023). Optimisasi Klasterisasi Nilai Ujian Nasional dengan Pendekatan Algoritma. *Jurnal Ilmu Komputer dan Sistem Informasi (JIKOMSI)*, 6(1), 14–20. <https://ejournal.sisfokomtek.org/index.php/jikom/article/view/1550>
- Lenama, N., Kleden, M. A., & Pasangka, I. G. (2023). K-Means Clustering Analysis pada Pengelompokan Kabupaten/Kota di Provinsi Nusa Tenggara Timur Berdasarkan Indikator Pendidikan. *Jurnal Cakrawala Ilmiah*, 2(9), 3365–3376. <https://doi.org/10.53625/jcijurnalcakrawalailmiah.v2i9.5653>
- Mariana, M. (2013). Analisis Komponen Utama. *Jurnal Matematika dan Pembelajarannya*, 1(2), 99–114. <https://doi.org/10.33477/mp.v1i2.304>
- Nahdliyah, M. A., Widiari, T., & Prahutama, A. (2019). Metode K-Medoids Clustering dengan Validasi Silhouette Index dan C-Index (Studi Kasus Jumlah Kriminalitas Kabupaten/Kota di Jawa Tengah Tahun 2018). *Jurnal Gaussian*, 8(2), 161–170. <https://doi.org/10.14710/j.gauss.8.2.161-170>



- Octaviyani, N. R., Mayasari, R., & Susilawati, S. (2022). Implementasi Algoritma K-Means Clustering Status Gizi Balita. *Jurnal Ilmiah Wahana Pendidikan*, 2022(13), 370–381. <https://doi.org/10.5281/zenodo.6962588>
- Peraturan Presiden (Perpres) Nomor 63 Tahun 2020 Tentang Penetapan Daerah Tertinggal Tahun 2020-2024, Pub. L. No. 63, Pemerintah Pusat Indonesia (2020).
- Prasetyawati, M. D. (2023). *Statistik Daerah Provinsi Jawa Tengah 2023* (D. Nursetyohadi & M. Samuharwadi, Eds.). Badan Pusat Statistik Provinsi Jawa Tengah.
- Pursitasari, I. D., Harianto, B., Wulan, S. S., Hermanto, D., & Ardianto, D. (2024). Multivariate Analysis of Variance (MANOVA) di Bidang Kesehatan dan Pendidikan MIPA. *Jurnal Ilmiah Kanderang Tingang*, 15(1), 117–126. <https://doi.org/10.37304/jikt.v15i1.307>
- Setiawan, I. N., Kristiani, S. N., Nurarifin, N., Delyana, S., Setyawati, N., & Arsyi, F. A. (2022). *Indeks Pembangunan Manusia 2022* (W. Winardi, Y. Karyono, M. Mutijo, & D. H. Santoso, Eds.). Badan Pusat Statistik.
- Sitinjak, D. K., Pangestu, B. A., & Sari, B. N. (2022). Clustering Tenaga Kesehatan Berdasarkan Kecamatan di Kabupaten Karawang Menggunakan Algoritma K-Means. *Journal of Applied Informatics and Computing*, 6(1), 47–54. <https://doi.org/10.30871/jaic.v6i1.3855>
- Sutrisno, S., & Wulandari, D. (2018). Multivariate Analysis of Variance (MANOVA) untuk Memperkaya Hasil Penelitian Pendidikan. *AKSIOMA: Jurnal Matematika dan Pendidikan Matematika*, 9(1), 37. <https://doi.org/10.26877/aks.v9i1.2472>
- Thamrin, D. R., & Murni, D. (2022). Analisis Cluster Hierarki Metode Single Linkage pada Kabupaten/Kota di Provinsi Sumatera Barat Berdasarkan Indikator Kesehatan. *Journal of Mathematics UNP*, 7(3), 45–51. <https://doi.org/10.24036/unpjomath.v7i3.12988>
- Wijaya, T. A., Utami, E., & Fatta, H. (2024). Perbandingan Algoritma DBSCAN dan K-Means Clustering untuk Pengelompokan Data Gangguan PT. PLN UID Kalselteng. *Innovative: Journal of Social Science Research*, 4(1), 8846–8854. <https://doi.org/10.31004/innovative.v4i1.8920>
- Yuan, C., & Yang, H. (2019). Research on K-Value Selection Method of K-Means Clustering Algorithm. *J—Multidisciplinary Scientific Journal*, 2(2), 226–235. <https://doi.org/10.3390/j2020016>
- Yudhistira, A., & Andika, R. (2023). Pengelompokan Data Nilai Siswa Menggunakan Metode K-Means Clustering. *Journal of Artificial Intelligence and Technology Information (JAITI)*, 1(1), 20–28. <https://doi.org/10.58602/jaiti.v1i1.22>
- Zulyani, A. A., Irawan, A. S. Y., & Jamaludin, A. (2023). Penerapan Data Mining Menggunakan Algoritma K-Means untuk Menentukan Tingkat Vaksinasi pada Kecamatan Tambun Selatan. *Innovative: Journal of Social Science Research*, 3(3), 7037–7050. <https://j-innovative.org/index.php/Innovative/article/view/2946>



LAMPIRAN A

Tabel 7 Nilai VIF Antar Variabel Pendidikan

Nilai VIF	X ₁	X ₂	X ₃	X ₄	X ₅	X ₆	X ₇	X ₈	X ₉	X ₁₀	X ₁₁	X ₁₂	X ₁₃	X ₁₄
X ₁		12,430	12,895	12,990	2,106	12,645	12,983	13,152	12,705	13,157	13,136	13,112	12,981	13,168
X ₂	24,883		16,750	25,037	24,894	9,430	21,500	25,642	26,376	25,017	26,341	26,367	26,210	26,325
X ₃	30,893	20,046		21,368	30,613	27,883	20,616	31,219	31,282	30,662	29,718	31,550	31,308	31,149
X ₄	21,170	20,384	14,536		21,418	20,903	20,108	15,436	20,969	21,106	20,446	21,368	21,349	21,307
X ₅	2,086	12,318	12,657	13,017		12,241	12,933	12,605	12,684	12,886	12,981	12,982	13,004	13,002
X ₆	43,926	16,364	40,430	44,544	42,928		36,050	39,950	45,725	43,871	45,735	45,715	45,347	45,634
X ₇	20,886	17,278	13,844	19,848	21,005	16,695		19,389	21,046	21,174	21,197	21,087	20,832	20,982
X ₈	21,112	20,562	20,917	15,203	20,427	18,461	19,346		21,146	21,019	21,107	20,762	21,070	21,124
X ₉	9,781	10,144	10,052	9,905	9,859	10,134	10,071	10,142		8,557	10,064	3,884	9,781	10,008
X ₁₀	7,106	6,750	6,913	6,995	7,027	6,821	7,109	7,073	6,004		5,943	6,424	4,597	7,069
X ₁₁	7,441	7,072	7,027	7,106	7,424	7,458	7,464	7,449	7,405	6,233		7,381	7,299	3,968
X ₁₂	6,672	6,702	6,701	6,671	6,669	6,696	6,669	6,581	2,567	6,052	6,629		5,807	6,084
X ₁₃	3,026	3,052	3,046	3,054	3,060	3,043	3,019	3,060	2,962	1,984	3,004	2,660		3,066
X ₁₄	5,577	5,577	5,507	5,537	5,560	5,564	5,524	5,574	5,506	5,543	2,967	5,501	5,572	

Tabel 8 Nilai Eigen dan Proporsi Kumulatif Data Pendidikan

Komponen Utama	PC1	PC2	PC3	PC4	PC5	PC6	PC7	PC8	PC9	PC10	PC11	PC12	PC13	PC14
Nilai Eigen	8,0097	2,9519	1,1257	0,7090	0,5077	0,1934	0,1373	0,1040	0,0670	0,0637	0,0564	0,0399	0,0226	0,1116
Proporsi Kumulatif	0.5721	0.7830	0.8634	0.9140	0.9503	0.9641	0.9739	0.9814	0.9862	0.9907	0.9947	0.9976	0.9992	1.0000



Spatial Decision Support System to Determine the Feasibility of Evacuation Posts in Natural Disasters

Nuril Afni Alviola ^{(1)*}, Agung Teguh Wibowo Almais ⁽²⁾, A'la Syauqi ⁽³⁾, Totok Chamidy ⁽⁴⁾,
Puspa Miladin Nuraida Safitri A Basid ⁽⁵⁾, Anisa Anisa ⁽⁶⁾, M. Dafa Wardana ⁽⁷⁾

Department of Informatic Engineering, UIN Maulana Malik Ibrahim, Malang, Indonesia

e-mail : {nurilafnia, anisanisag73, mdafawardana}@gmail.com,

{agung.twa, syuqi, fachrulk}@ti.uin-malang.ac.id, puspa.miladin@uin-malang.ac.id.

* Corresponding author.

This article was submitted on 5 August 2024, revised on 26 February 2025, accepted on 28 February 2025, and published on 30 September 2025.

Abstract

This study aimed to improve the accuracy of determining the feasibility of evacuation posts after natural disasters using the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) within a Spatial Decision Support System (SDSS). A dataset of 50 evacuation posts from the 2021 Mount Semeru eruption was analyzed. The Rank Order Centroid (ROC) method was applied for criteria weighting, and TOPSIS was used to process the data. Results showed 72% accuracy, confirming that TOPSIS is a passable method for assessing post-feasibility based on accessibility, sanitation, and refugee facilities. Although the focus is on evaluating post-disaster evacuation posts, the system can be adapted for use in various other types of disasters. However, it is still dependent on historical data and lacks real-time adaptability. Future research can integrate Artificial Intelligence (AI) and Machine Learning (ML) with real-time data to improve decision-making in disaster management.

Keywords: After Natural Disasters, Evacuation Posts, Management Disasters, Spatial Decision Support System, Technique For Order Preference By Similarity to Ideal Solution (TOPSIS)

Abstrak

Penelitian ini bertujuan meningkatkan akurasi penentuan lokasi pos pengungsian bencana alam menggunakan metode Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) dalam Sistem Pendukung Keputusan Spasial (SDSS). Dataset yang digunakan terdiri dari 50 pos pengungsian erupsi Gunung Semeru tahun 2021. Metode Rank Order Centroid (ROC) digunakan untuk memberikan bobot pada kriteria, dan TOPSIS untuk pemrosesan data. Hasil penelitian menunjukkan akurasi 72%, membuktikan bahwa TOPSIS cukup baik dalam mengevaluasi kelayakan pos berdasarkan aksesibilitas, sanitasi, dan fasilitas pengungsi. Meskipun berfokus pada evaluasi pos pengungsian pasca bencana alam, sistem ini dapat beradaptasi ke berbagai jenis bencana lainnya, seperti gempa bumi, banjir, badai, dan kebakaran hutan. Namun, sistem ini masih bergantung pada data historis dan belum dapat menganalisis secara real-time, yang membatasi respons cepat dalam situasi darurat. Penelitian di masa depan dapat mengintegrasikan Kecerdasan Buatan (AI) dan Machine Learning (ML) dengan data real-time untuk meningkatkan akurasi dan pengambilan keputusan dalam manajemen bencana.

Kata Kunci: Pasca Bencana Alam, Pos Pengungsian, Manajemen Bencana, Sistem Pendukung Keputusan Berbasis Spasial, Technique For Order Preference By Similarity to Ideal Solution (TOPSIS)

1. INTRODUCTION

Natural disasters are events that occur as a result of natural phenomena, such as earthquakes, tsunamis, floods, landslides, volcanic eruptions, and extreme weather events like storms and forest fires (Meilano et al., 2020). Natural disasters, as defined by Law No. 24 of 2007, are events resulting from natural phenomena. These disasters cause significant damage and losses across various sectors (Fathiyaturahma, 2023). Indonesia experienced 2.594 natural disasters from January to August 2023, including floods, extreme weather, landslides, forest fires, tidal waves,



This article is distributed following Attribution-NonCommercial CC BY-NC as stated on <https://creativecommons.org/licenses/by-nc/4.0/>.

earthquakes, droughts, and volcanic eruptions (Badan Nasional Penanggulangan Bencana, 2023). This highlights Indonesia's high vulnerability to natural disasters (Ozbay et al., 2019). Different regions in Indonesia face different types of natural disasters due to their diverse geographical conditions (Setiawan & Fitriani, 2024). Lumajang Regency in East Java is particularly prone to disasters, with Mount Semeru, the highest mountain in Java, located there. In 2021, a volcanic eruption in Lumajang caused significant damage, disrupting access to bridges, village roads, and homes (Perwita et al., 2023).

In response to disasters, the government, with various stakeholders, establishes evacuation and protection posts to provide critical assistance (Ismeti et al., 2023). These posts are essential in disaster management and are divided into eight sub-clusters, including shelter, water and sanitation, coordination and management, security, child protection, and psychosocial support (Kementerian Sosial Republik Indonesia, 2019). Evacuation posts must meet specific criteria to minimize disaster risks, including adequate land, water supply, sanitation facilities, and access to public services. The Regional Disaster Management Agency (BPBD) determines evacuation posts by referring to the contingency plan (Renkon). Renkon is a joint agreement in disaster management. However, not all evacuees are in the evacuation posts established by BPBD, so some evacuees occupy places that are not in the designated Renkon area. A proper evacuation post is in accordance with the agreed-upon Renkon; therefore, the locations of evacuation posts not included in the Renkon have subjective feasibility. In an emergency, people tend to make decisions under emotional pressure, which can make decision-making ineffective. Determining the feasibility of an evacuation can utilize information technology. In decision-making, surveyors require technology to make structured and objective decisions. One of the information technologies that can be used to determine the feasibility of evacuation posts is the Decision Support System (DSS).

A DSS is a system that can be used for objective decision-making in the face of information uncertainty (Ojo et al., 2024). A DSS uses a Computer-Based Information System (CBIS) developed to support flexible and interactive decision-making systems (Trieu et al., 2024). A DSS will help make decisions based on appropriate criteria to get an existing alternative (Xie et al., 2024). One type of DSS that can be used to make decisions is the Spatial Decision Support System (SDSS). An SDSS is a decision support system that uses a spatial or geographic approach. Spatial is an approach that utilizes geographical location (Keenan & Jankowski, 2019). Web-based SDSS are designed to support all types of data, from vector to raster. The location approach in this decision support system can use latitude and longitude coordinates. The spatial approach in the research to be carried out is in the form of coordinates for the locations of established evacuation posts. Surveyors will input data related to the location of the evacuation post, and the system will analyze and process the information to determine the suitability of the input data against predetermined criteria. After the data is processed, the system selects the best alternative, considering it the most suitable based on the input data.

A DSS in decision-making employs a specialized method known as Multiple Criteria Decision Making (MCDM). MCDM is an approach to decision-making that involves many criteria or decision support factors (Shao et al., 2024). MCDM employs several approaches, one of which is the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS). TOPSIS is one of the multi-criteria decision support system methods used to solve the problem of selecting the best alternative (Dewi et al., 2021). Decision support systems can be developed using the TOPSIS approach in decision-making. The TOPSIS approach is carried out by determining the alternative that is closest to the positive and negative ideal solution values based on existing criteria. This approach can help overcome the multiple criteria that exist within the system. The TOPSIS approach helps overcome the complexity of criteria and provides priority rankings on alternatives, enabling the system to make decisions quickly.

The TOPSIS approach was selected as the main decision-making instrument in this study to assess the viability of evacuation posts following Mount Semeru's 2021 eruption. Because it can handle multiple criteria simultaneously and rank options according to how closely they approach



an ideal solution, the TOPSIS technique was chosen. The approach is especially helpful for complex decision-making situations that involve multiple variables, such as safety, accessibility, and available facilities. The study aims to provide a methodical and impartial approach for determining the best evacuation points using TOPSIS.

The main reference for this research is the research conducted by Hadiguna et al., titled “Implementing a Web-Based Decision Support System for Disaster Logistics: A Case Study of an Evacuation Location Assessment for Indonesia.” The research developed a web-based decision support system used to evaluate the feasibility of public facilities as post-disaster evacuation sites (Hadiguna et al., 2014). Unlike previous research that employed the TOPSIS method to assess the level of building damage following a natural disaster, this research focuses on the subject of evacuation post-eligibility. The calculation of the weight of each criterion uses the Rank Order Centroid (ROC) method. This research also applies the SDSS by obtaining latitude and longitude points on Google Maps. This research uses seven criteria and three alternatives. The alternatives used are categorized as feasible, somewhat feasible, and not feasible. At the same time, the criteria considered include health facilities, drinking and cooking water, bathing, washing, and toilet facilities, closed bathrooms and toilets, waste disposal, and refugee privacy areas.

2. METHODS

The system programming in this study was carried out using the PHP (Hypertext Preprocessor) programming language. The method implementation utilized the Laragon database to display views from the queries of each stage. Figure 1 shows that the developed system includes input, process, and output stages. There are three stages in the research data processing cycle: input, process, and output. The input stage is the stage for entering data into the Spatial Decision Support System (SDSS). The data that needs to be input into the system consists of location, criteria, alternatives, weights, and a decision matrix. The process of inputting the longitude and latitude points into the SDSS uses the coordinate points on Google Maps. The process stage is the stage of data processing to achieve the desired results. This stage involves processes such as determining the divider value, data normalization, weighted normalization, identifying the maximum and minimum values, determining the positive and negative ideal solution values, determining the size of the separator, and assigning alternative preference values. The output stage is the result of the data processing that has been performed. The output in question is an alternative result that matches the input data. In the system for determining the feasibility of natural disaster refugee posts, the alternatives that emerge as output are categorized as feasible, quite feasible, and not feasible. The system design based on these three stages is structured as follows.

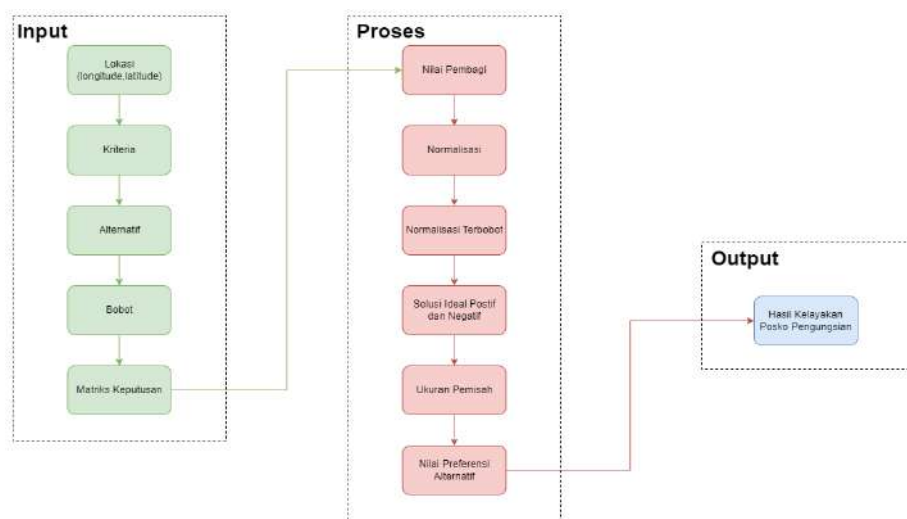


Figure 1 System Design



This research utilizes data collected and compiled by other parties or sources, known as secondary data. The secondary data for this study were obtained from data released by the National Disaster Management Agency (BNPB) and the Regional Disaster Management Agency (BPBD) of Lumajang District. The dataset consisted of 50 records on evacuation posts for the 2021 eruption of Mount Semeru. Several criteria were obtained from this previously collected data. Table 1 presents the criteria that will be used as supporting factors for determining suitable alternatives. These criteria are useful to assist surveyors in determining the feasibility of post-natural disaster refugee camps in Lumajang Regency.

Table 1 Research Criteria

Criteria Code	Criteria Name
K1	Drinking and cooking water
K2	Bathing, washing, and toilet water
K3	Closed Bathroom
K4	Toilet
K5	Waste Disposal
K6	Refugee Privacy Temp
K7	Drinking and cooking water

The original data obtained from BNPB and BPBD will be processed by providing a rating scale on predetermined criteria. Table 2 presents the criteria assessment provisions used in the research. The assessment scale for each criterion uses a statistical scale, namely an ordinal scale. An ordinal scale provides ranked data or sequential values, where each value is assigned a rank. Values on an ordinal scale have a lower or smaller value than other values. The ordinal scale also lacks a definite distance between each category. The value on the ordinal scale does not represent an exact value, but rather indicates a value that is considered higher or lower than the comparison value. Based on a research journal entitled "Application of the ROC-TOPSIS Method in the Decision on Recipients of the Family Hope Program," sequential category data values can be provided from the range of 1-3 for each data rank according to the sub-criteria used. In addition to the criteria obtained from the description above, there are alternatives, which are the conclusions drawn from comparing the criteria. Three alternatives will be used in this study as described in Table 3.

Table 2 Criteria Value Conditions

Criteria	Sub- Criteria	Value
Health Facilities	None	1
	> 1.6 km	2
	<= 1.6 km	3
Drinking and cooking water	Not available	1
	Existing and insufficient	2
	Existing and sufficient	3
Bathing, washing, and toilet water	None	1
	Existent and insufficient	2
	Present and sufficient	3
Closed Bathroom	None	1
	Present and insufficient	2
	Present and sufficient	3
Toilet	None	1
	Present and insufficient	2
	Present and sufficient	3
Waste Disposal	None	1
	Present and insufficient	2
Refugee Privacy Temp	None	1
	Existing and insufficient	2



Table 3 Research Alternative

Alternative Code	Alternative Name
A1	Not Feasible
A2	Decent Enough
A3	Feasible

2.1 Rank Order Centroid (ROC)

The determination of criteria weights is performed using the ROC method. Before determining the weight, the criteria will be sorted based on priority. Criteria will also be determined, including the type of benefit or cost associated with each option. The type of criteria is obtained after the rating scale for each criterion is determined. Criteria that have an assessment scale where the lower the assessment provision, the better the criterion will be, are a type of cost criterion. If, on the contrary, the type of criteria is a benefit. The determination of the type of criteria is described in Table 4.

Table 4 Research Criteria

Criteria Code	Criteria Name	Type
K1	Drinking and cooking water	Cost
K2	Bathing, washing, and toilet water	Benefit
K3	Closed Bathroom	Benefit
K4	Toilet	Benefit
K5	Waste Disposal	Benefit
K6	Refugee Privacy Temp	Benefit
K7	Drinking and cooking water	Benefit

In determining weights, the ROC (Rank Order Centroid) method prioritizes the most important criteria. This prioritization is expressed by ordering the criteria as shown in Eq. (1), with the corresponding ROC weights following the same order, as presented in Eq. (2). The ROC method calculates the weight of each criterion to reflect its relative importance in a decision support system (DSS), as illustrated in Eq. (3) (Damanik & Utomo, 2020). Using an average approximation, the weights for each criterion can be determined as shown in Eq. (4)–(7) (Simanjuntak et al., 2022), ensuring that higher-priority criteria receive greater weights.

$$K_1 \geq K_2 \geq K_3 \geq \dots \geq K_n \quad (1)$$

$$W_1 \geq W_2 \geq W_3 \geq \dots \geq W_n \quad (2)$$

$$W_k = \frac{1}{k} \sum_{i=k}^n \left(1 + \frac{1}{k}\right) \quad (3)$$

$$W_1 = \frac{1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{k}}{k} \quad (4)$$

$$W_2 = \frac{0 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{k}}{k} \quad (5)$$

$$W_3 = \frac{0 + 0 + \frac{1}{3} + \dots + \frac{1}{k}}{k} \quad (6)$$

$$W_k = \frac{0 + \dots + 0 + \frac{1}{k}}{k} \quad (7)$$



2.2 Technique for Order Preference by Similarity to Ideal Solution (TOPSIS)

TOPSIS, or Technique for Order Preference by Similarity to Ideal Solution, is a multi-criteria decision-making (MCDM) method used to identify the best alternative based on its proximity to the positive ideal solution (PIS) and distance from the negative ideal solution (NIS). This approach is fundamental in MCDM, where decision-making involves multiple criteria that influence the selection process (Baydaş et al., 2024). TOPSIS calculates how far each alternative deviates from the PIS, which represents the highest value for each criterion, and the NIS, which represents the lowest value for each criterion. The distances between alternatives and ideal solutions are computed using the Euclidean method (Divayana & Suyasa, 2021). The TOPSIS procedure begins with constructing the decision matrix as shown in Eq. (8), followed by calculating the divider value (Eq. 9) to normalize the data (Eq. 10). Each criterion is then weighted to obtain the weighted normalized values (Eq. 11). Positive and negative ideal solutions are determined for benefit and cost criteria as illustrated in Eq. (12) and Eq. (13). Subsequently, the separation distances from the PIS and NIS are calculated using Eq. (14) and Eq. (15). Finally, the preference value for each alternative is derived as presented in Eq. (16), which enables ranking the alternatives according to their closeness to the ideal solution (Dewi et al., 2021).

$$D = (x_{ij})_{m \times n} \quad (8)$$

$$N = \sqrt{\sum_{i=1}^m x_{ij}^2} \quad (9)$$

$$r_{ij} = \left(\frac{x_{ij}}{N}\right)_{m \times n} \quad (10)$$

$$y_{ij} = (w_j \cdot r_{ij})_{m \times n} \quad (11)$$

$$A^+ = \{(^{max}_i y_{ij} | j \in J_b), (^{min}_i y_{ij} | j \in J_c)\} \quad (12)$$

$$A^- = \{(^{min}_i y_{ij} | j \in J_b), (^{max}_i y_{ij} | j \in J_c)\} \quad (13)$$

$$S_i^+ = \sqrt{\sum_{j=1}^n (y_{ij} - A^+)^2} \quad (14)$$

$$S_i^- = \sqrt{\sum_{j=1}^n (y_{ij} - A^-)^2} \quad (15)$$

$$v_i = \frac{S_i^-}{S_i^+ + S_i^-} \quad (16)$$

3. RESULTS AND DISCUSSION

The TOPSIS method is implemented using PHP and MYSQL programming. The stages of the TOPSIS method calculation process are carried out by querying the views contained in the database. The TOPSIS calculation process at each stage is carried out with interconnected views. The PHP programming language is used to implement the TOPSIS method in the form of a website. The following are the results of implementing the TOPSIS method.



3.1 Data Processing

The data processed is information on the evacuation post for the 2021 eruption of Mount Semeru. The total dataset used consists of 50 data points on evacuation posts for the 2021 eruption of Mount Semeru. The data obtained are data on the name of the refugee post, refugee address, latitude point, longitude point, distance to health facilities, availability and adequacy of bathing, washing and toilet water, availability and adequacy of bathing and drinking water, availability and adequacy of bathrooms, availability, adequacy and number of toilets, amount of drinking water, amount of clean water, availability of waste disposal, and availability of refugee privacy places. The data obtained is in the form of category data on each criterion. The data will be converted into quantitative data so that it can be processed in the calculation of the TOPSIS method. Table 5 presents the data preprocessing that can be used in research.

Table 5 Preprocessing Data Result

No.	Pos Name	Pos Address	Latitude	Longitude	K1	K2	K3	K4	K5	K6	K7
1	Balai Desa Gesang	Dusun Krajan I, Gesang, Tempeh	-8,17882052032194	113,147427494563	3	3	3	3	3	2	2
2	Balai desa Jambearum	Pakem, Jambearum, Pasrujambe, Lumajang	-8,1683721537445	113,078814435035	3	3	3	3	3	2	2
3	Balai desa Karang Anom	Jl. Raya Semeru, Tambakrejo Wetan, Karang Anom, Pasrujambe	-8,3824920798675	113,145563024991	3	3	3	3	3	2	2
4	Balai desa Klopok Sawit	Krajan, Klopokawit, Kec. Candipuro, Lumajang	-8,14873467600424	113,081888178848	2	3	3	3	3	2	2
5	Balai desa Pasrujambe	Jl. Raya Semeru, Pasrujambe, Pasrujambe	-8,113237391364702	113,044102294561	2	3	3	3	3	2	2

3.2 System Interface Implementation

The decision support system in determining the feasibility of post-disaster refugee camps is implemented using a website-based user interface. The website features several pages tailored to the user's needs. Figure 2 shows the SDSS website for the feasibility of post-disaster refugee camps. The website display consists of a login page, dashboard, evacuation post input, user, criteria, alternatives, criteria weights, rating scale, assessment input, TOPSIS method process results page, and user profile page.

The SDSS implementation on the web comprises several pages that facilitate ease of system interaction with users. The website features maps that display the locations of evacuation posts, along with information related to each post, such as its eligibility as an evacuation site. The input page also provides a map display, allowing users to directly select the location point of the evacuation post without manually entering the longitude and latitude points. The result display of evacuation posts that have been checked with SDSS will display the stages of the calculation process until the results of the TOPSIS calculation process are obtained.



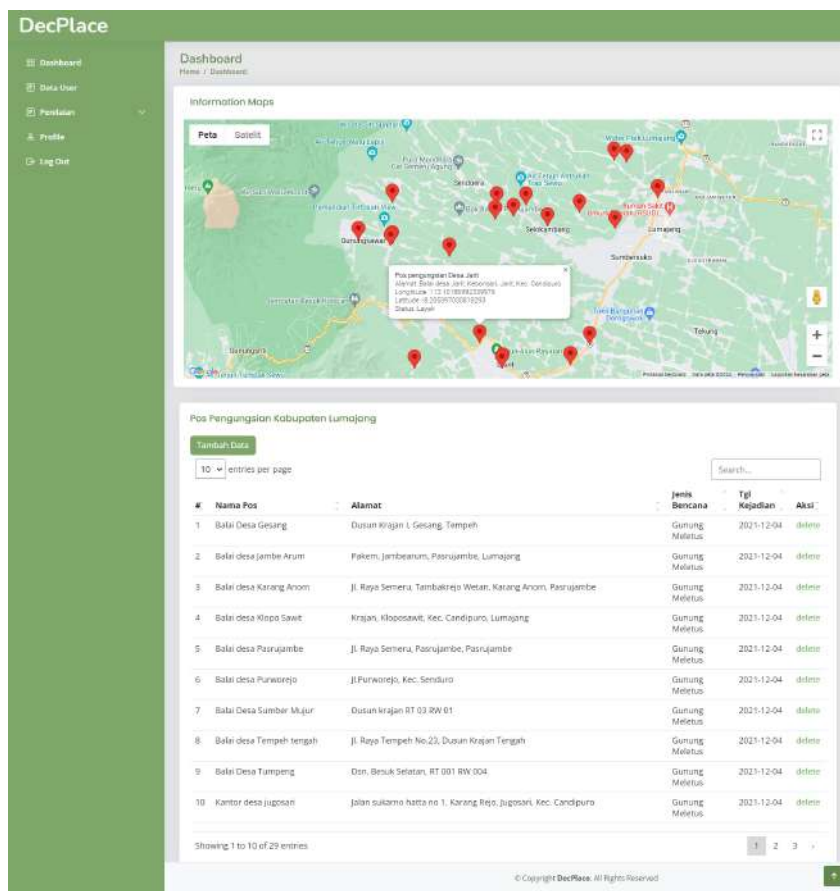


Figure 2 Dashboard Page

3.3 Accuracy Calculation

The accuracy of the TOPSIS method is calculated by comparing the corresponding data between the original data and the system data. Table 6 presents the calculation results obtained using the TOPSIS method and actual data. From the system calculations above, it is evident that 14 data points yield different results from the TOPSIS calculation data. To determine the accuracy of the TOPSIS method, the steps in Eq. (17) can be taken.

Table 6 Comparison Table of Original Data with TOPSIS Data

No.	Evacuation Name	K1	K2	K3	K4	K5	K6	K7	Original Data	TOPSIS Result	Description
1	Balai Desa Gesang	3	3	3	3	3	2	2	Feasible	Feasible	Suitable
2	Balai desa Jambe Arum	3	3	3	3	3	2	2	Feasible	Feasible	Suitable
3	Balai desa Karang Anom	3	3	3	3	3	2	2	Feasible	Feasible	Suitable
...	...	...	...	...	...	...	...	...	...	...	...
50	PPKM MIKRO	2	3	3	3	3	2	2	Feasible	Feasible	Suitable

$$Accuracy = \frac{\text{suitable data}}{\text{total number of data}} \times 100\% = \frac{36}{50} \times 100\% = 0,72 \times 100\% = 72\% \quad (17)$$



3.4 Discussion

The test data consist of records from refugee camps following the eruption of Mount Semeru in Lumajang Regency in 2021. The data used is in the form of data on the name of the refugee post, the address of the refugee post, the supply of facilities located at the refugee post, such as health facilities, bathing, washing, and toilet water supplies, drinking and cooking water supplies, bathrooms, toilets, waste disposal, and refugee privacy areas. Each evacuation post's data includes a description of its eligibility, which is based on the analysis by BPBD Lumajang District. The refugee post eligibility data obtained will be compared with the data generated by the system to assess the accuracy of the method. The data obtained is qualitative, so that it will be processed into quantitative data for further analysis. The qualitative data will be given a rating scale in accordance with Table 2.

Criteria 1 is the top priority in the form of health facility distance. This type of criterion is a cost criterion because the smaller the distance to the health facility, the higher the value obtained. Meanwhile, the farther the distance of health facilities, the smaller the value obtained. The meaning of the unavailability of health facilities is the absence of health posts such as emergency health posts, village health posts (POSKESDES), or the nearest health center from the location of the evacuation post. Criteria 1 in the preprocessed data yields values of "3" for 22 data points, "2" for 28 data points, and "1" for 0 data points. From these results, it is evident that all evacuation posts meet the criteria for health facilities that must be located near the evacuation post and are therefore a priority for consideration.

In criterion 2, namely the supply of drinking and cooking water, which has a ranking of the 2nd level of importance among all criteria. This criterion is a benefit criterion because the more water supplies and the adequacy of drinking and cooking water, the higher the value obtained. In this criterion, refugee posts that have drinking and cooking water data considered "available and sufficient" will be assigned a value of 3, as it is deemed to have a sufficient water supply for all refugees in the refugee post. For data that has the availability of cooking and drinking water, but is insufficient for all refugees, it will be assigned a value of 2. Meanwhile, the refugee post does not have a supply of drinking and cooking water; it will be given a value of 1. From the data that has undergone the data preprocessing stage, it is observed that there are 50 refugee posts with a value of "3". From this amount of data, it can be seen that all refugee posts have sufficient water availability for all refugees.

In criterion 3, namely bathing, washing, and toilet water supply or bathing washing latrines, which have a ranking of the 3rd level of importance from the total of all criteria. This criterion is a benefit criterion because the more water supply and water adequacy, the higher the value obtained. Similar to the 2nd criterion, this criterion also considers the same sub-criteria in assigning value to the data. In this criterion, 46 data points were obtained for a value of "3", 3 data points for a value of "2", and 1 data point with a value of "1". From this data, it is evident that some refugee posts have sufficient water supplies for toilets, while others do not. This can be caused by many factors, one of which may be a small or nonexistent water source.

In criterion 4, namely a closed bathroom, which has a ranking of the 4th level of importance among all criteria. This criterion is a benefit criterion because the more bathroom supplies are available, the higher the value obtained. In this criterion, refugee posts that have "available and sufficient" data will be assigned a value of 3, as they are considered to have sufficient bathrooms for all refugees in the refugee post. For data that has a bathroom but is insufficient for all refugees, it will be assigned a value of 2. Meanwhile, if the refugee post does not have a bathroom, it will be given a value of 1. Criteria 4 in the preprocessed data yields values of "3" for 46 data points, "2" for 3 data points, and "1" for 1 data point. From this amount of data, it can be seen that most refugee camps have sufficient availability and adequacy of bathrooms for refugees.

In criterion 5, namely the number of toilets that have a ranking of the 5th level of importance of the total of all criteria. This criterion is a benefit criterion because the more the number of toilets,



the higher the value obtained. In this criterion, refugee posts that have “there and enough” data will be given a value of 3. For data that has a bathroom but is insufficient for all refugees, a value of 2 will be assigned. If the refugee post does not have a toilet, a value of 1 will be given. Criteria 5 in the preprocessed data yields values of “3” for 44 data points, “2” for 5 data points, and “1” for 1 data point. From this amount of data, it can be seen that most refugee camps have toilets for refugees and a sufficient number of toilets for all refugees. Meanwhile, five refugee camps have toilets, but there are not enough for all refugees. This can happen if the number of toilets available is not proportional to the number of refugees at the refugee camp.

In criterion 6, namely waste disposal, which has a ranking of the 6th level of importance among all criteria. This criterion is a benefit criterion because if there is waste disposal, the value at the evacuation post will be higher, namely “2”. However, refugee camps that do not have waste disposal will receive a lower value, namely “1”. Criterion 6 in the preprocessed data yields a value of “2” for 49 data points and a value of “1” for 1 data point. From this amount of data, it can be seen that almost all refugee camps have proper waste disposal.

In criterion 7, namely the place of refugee privacy, which has a ranking of the 7th level of importance out of all criteria. This criterion is a benefit criterion because having a place of privacy for refugees increases the value of the refugee camp, specifically to “2”. However, if the refugee camp does not have a place of privacy, the value obtained will be lower, namely “1”. A private space for refugees can be in the form of a plywood partition or a room for several family units (KK) who are still part of one family. This privacy space is useful for maintaining the privacy of each family or promoting peace among refugees. Criterion 7 in the preprocessed data contains values of “2” for 39 data points and “1” for 11 data points.

The results of testing the accuracy of the TOPSIS method obtained an accuracy of 72%. According to Pratiwi and Wibowo, the 72% accuracy result suggests that the TOPSIS method can accurately predict the correct data (Niu et al., 2024). The system testing showed an accuracy rate of 72%, with 36 matching results and 14 non-matching results when compared to actual data from BPBD Lumajang. While this accuracy indicates that TOPSIS can effectively assess evacuation posts based on historical data, the method has some limitations, particularly in distinguishing between the “Feasible” and “Decent Enough” categories. This suggests that the weighting process used in TOPSIS may not fully capture the conditions of post-disaster situations, where external factors such as infrastructure damage and population density may affect the actual feasibility of evacuation sites.

Although this study focused on the 2021 Mount Semeru eruption, the proposed SDSS-based TOPSIS system has the potential for wide applications. Many disaster-prone regions worldwide, including earthquake zones (such as some areas in Indonesia), flood-prone areas, hurricane-prone regions, and wildfire-affected areas, could benefit from a structured decision-making tool to evaluate evacuation post-feasibility. By adjusting the criteria weights based on specific geographic and disaster conditions, this system can be customized to suit various disaster scenarios, ensuring more effective preparedness and response strategies.

The existing system’s dependence on historical data and lack of real-time adaptability. Decision-making in a rapidly changing crisis scenario requires dynamic analysis based on real-time data, rather than static records. To improve system performance, future research can combine machine learning (ML) with artificial intelligence (AI). Real-time evaluation of evacuation post-feasibility is made possible by AI-driven models that continually learn from previous evacuation results. Combining TOPSIS with AI/ML would enable the system not only to evaluate previously used evacuation sites but also to predict the best locations for future evacuation posts. This hybrid approach will increase decision-making speed, adaptability, and accuracy, making it more relevant for large-scale disaster management. Furthermore, real-time integration can help disaster response teams quickly identify the safest and most suitable locations for evacuation posts, ensuring disaster preparedness and response.



Overall, while TOPSIS has proven to be a reliable and well-structured method for evaluating evacuation posts, its effectiveness can be further enhanced by integrating AI and real-time data processing. Future research should explore hybrid decision-making models that incorporate machine learning and real-time geographic data analysis, allowing for more precise and adaptive disaster response planning. This research lays the foundation for further advancements in data-driven disaster management systems, helping decision-makers optimize the selection of evacuation posts on both local and global scales.

4. CONCLUSIONS

This research successfully implemented the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) method in a web-based Spatial Decision Support System (SDSS) to determine the feasibility of evacuation posts after a natural disaster. The system was developed using PHP and MySQL, with the Rank Order Centroid (ROC) method applied for weighting the criteria. System testing showed an accuracy rate of 72%, with 36 matching results and 14 non-matching results compared to actual data from BPBD Lumajang. This system is designed to evaluate previously used evacuation posts, ensuring whether a location remains suitable for future disasters. By considering seven key criteria, including accessibility, sanitation, and refugee facilities, the system enables decision-makers to assess the effectiveness of evacuation posts based on historical data. Although this research focuses on the 2021 Mount Semeru eruption, the system has global potential and can be adapted for use in various disaster-prone areas. TOPSIS-based SDSS can be customized to accommodate different geographic conditions and disaster types, including earthquakes, floods, hurricanes, and wildfires. However, the system still relies on historical data and cannot yet determine evacuation post locations in real time. Future research can integrate artificial intelligence (AI) and machine learning (ML) to analyze data and dynamically predict the optimal evacuation post locations. This approach will enhance the system's accuracy and efficiency in making quick decisions during emergencies. In conclusion, this study contributes to the development of a data-driven decision support system for evaluating evacuation posts. It also serves as a foundation for future research in AI/ML-based systems that can determine evacuation post locations in real time, improving disaster response efforts worldwide.

REFERENCES

- Badan Nasional Penanggulangan Bencana. (2023). *Geoportal Data Bencana Indonesia*. Badan Nasional Penanggulangan Bencana. <https://gis.bnpb.go.id>
- Baydaş, M., Yılmaz, M., Jović, Ž., Stević, Ž., Özüyar, S. E. G., & Özçil, A. (2024). A Comprehensive MCDM Assessment for Economic Data: Success Analysis of Maximum Normalization, CODAS, and Fuzzy Approaches. *Financial Innovation*, 10(1), Article ID: 105. <https://doi.org/10.1186/s40854-023-00588-x>
- Damanik, S., & Utomo, D. P. (2020). Implementasi Metode ROC (Rank Order Centroid) dan Waspas dalam Sistem Pendukung Keputusan Pemilihan Kerjasama Vendor. *KOMIK (Konferensi Nasional Teknologi Informasi dan Komputer)*, 4(1), 242–248. <https://doi.org/10.30865/komik.v4i1.2690>
- Dewi, R. K., Jonemaro, E. M. A., Kharisma, A. P., Farah, N. A., & Dewantoro, M. F. (2021). TOPSIS for Mobile Based Group and Personal Decision Support System. *Register: Jurnal Ilmiah Teknologi Sistem Informasi*, 7(1), 43–49. <https://doi.org/10.26594/register.v7i1.2140>
- Divayana, D. G. H., & Suyasa, P. W. A. (2021). Simulation of TOPSIS Calculation in Discrepancy-Tat Twam Asi Evaluation Model. *Register: Jurnal Ilmiah Teknologi Sistem Informasi*, 7(2), 136–148. <https://doi.org/10.26594/register.v7i2.2196>
- Fathiyaturahma, A. (2023). Analysis of Prediction of Economic Losses due to Flood Disaster Hazards in Industrial Estates in Karawang Regency. *International Journal of Disaster Management*, 5(3), 203–212. <https://doi.org/10.24815/ijdm.v5i3.29740>
- Hadiguna, R. A., Kamil, I., Delati, A., & Reed, R. (2014). Implementing a Web-Based Decision Support System for Disaster Logistics: A Case Study of an Evacuation Location Assessment for Indonesia. *International Journal of Disaster Risk Reduction*, 9, 38–47. <https://doi.org/10.1016/j.ijdr.2014.02.004>



- Ismeti, I., Palipadang, L., Tavip, Moh., & Weru, A. (2023). Regional Government Responsibility Related to Disaster Mitigation Through Human Rights-Based Spatial Policies in Palu City. *International Journal of Disaster Management*, 5(3), 213–226. <https://doi.org/10.24815/ijdm.v5i3.30987>
- Keenan, P. B., & Jankowski, P. (2019). Spatial Decision Support Systems: Three Decades On. *Decision Support Systems*, 116, 64–76. <https://doi.org/10.1016/j.dss.2018.10.010>
- Kementerian Sosial Republik Indonesia. (2019). *Panduan Shelter untuk Kemanusiaan*. Kementerian Sosial Republik Indonesia. <https://sheltercluster.org/indonesia-2019-responses/documents/shelter-sub-cluster-panduan-shelter-untuk-kemanusiaan-bahasa>
- Meilano, I., Tiaratama, A. L., Wijaya, D. D., Maulida, P., Susilo, S., & Fitri, I. H. (2020). Analisis Potensi Gempa di Selatan Pulau Jawa Berdasarkan Pengamatan GPS. *Jurnal Lingkungan dan Bencana Geologi*, 11(3), 151–159. <https://doi.org/10.34126/jlbg.v11i3.352>
- Niu, S., Yin, Q., Ma, J., Song, Y., Xu, Y., Bai, L., Pan, W., & Yang, X. (2024). Enhancing Healthcare Decision Support Through Explainable AI Models for Risk Prediction. *Decision Support Systems*, 181, Article ID: 114228. <https://doi.org/10.1016/j.dss.2024.114228>
- Ojo, A., Rizun, N., Walsh, G., Mashinchi, M. I., Venosa, M., & Rao, M. N. (2024). Prioritising National Healthcare Service Issues from Free Text Feedback – A Computational Text Analysis & Predictive Modelling Approach. *Decision Support Systems*, 181(2022), Article ID: 114215. <https://doi.org/10.1016/j.dss.2024.114215>
- Ozbay, E., Çavuş, Ö., & Kara, B. Y. (2019). Shelter Site Location Under Multi-Hazard Scenarios. *Computers & Operations Research*, 106, 102–118. <https://doi.org/10.1016/j.cor.2019.02.008>
- Perwita, C. A., Aprilia, F., Maryanto, S., Arrasyid, H., & Tsabitah, A. F. (2023). Hazards Mitigation of Lahar Flows on Semeru Volcano After the 4 December 2021 Eruption Based on PS-InSAR. *International Journal of Disaster Management*, 5(3), 193–202. <https://doi.org/10.24815/ijdm.v5i3.29098>
- Setiawan, E., & Fitriani, W. A. (2024). Disaster Mitigation Strategies Based on Risk Matrix and House of Risk (HoR) Phase 2. *Jurnal Pengelolaan Sumberdaya Alam dan Lingkungan (Journal of Natural Resources and Environmental Management)*, 14(2), 341–353. <https://doi.org/10.29244/jpsl.14.2.341>
- Shao, J., Zhong, S., Tian, M., & Liu, Y. (2024). Combining Fuzzy MCDM with Kano Model and FMEA: A Novel 3-Phase MCDM Method for Reliable Assessment. *Annals of Operations Research*, 342(1), 725–765. <https://doi.org/10.1007/s10479-024-05878-w>
- Simanjuntak, P., Mesran, M., & Sianturi, R. D. (2022). Sistem Pendukung Keputusan Seleksi Penerima Dokter di Rumah Sakit Umum Bhakti dengan Menerapkan Metode Oreste dan ROC. *Resolusi: Rekayasa Teknik Informatika dan Informasi*, 2(3), 121–127. <https://doi.org/10.30865/resolusi.v2i3.307>
- Trieu, V.-H., Vu, H. Q., Indulska, M., & Li, G. (2024). A Computer Vision-Based Concept Model to Recommend Domestic Overseas-Like Travel Experiences: A Design Science Study. *Decision Support Systems*, 181, Article ID: 114149. <https://doi.org/10.1016/j.dss.2023.114149>
- Xie, Y., Yeoh, W., & Wang, J. (2024). How Self-Selection Bias in Online Reviews Affects Buyer Satisfaction: A Product Type Perspective. *Decision Support Systems*, 181(2023), Article ID: 114199. <https://doi.org/10.1016/j.dss.2024.114199>



Implementasi Metode YOLOv8 Mendeteksi Komputer Aktif dengan Subjek Layar Monitor

Frisky Wijaya ^{(1)*}, Dedy Hermanto ⁽²⁾

Departemen Informatika, Universitas Multi Data, Palembang, Indonesia

e-mail : friskywijaya@mhs.mdp.ac.id, dedy@mdp.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 26 Agustus 2024, direvisi 19 Februari 2025, diterima 20 Februari 2025, dan dipublikasikan 30 September 2025.

Abstract

Computers are one example of technological advances used in education. The use of computers that are not turned off can cause damage to computer components, and the use of electrical energy can increase. Student disobedience in turning off school laboratory computers when finished using them causes teachers to conduct manual checks by visiting each computer laboratory in the school. Deep learning is a machine learning algorithm that uses artificial neural networks. Deep learning is usually used for image recognition, voice identification, and data pattern analysis. Therefore, this study will apply the Deep Learning method, specifically YOLOv8, which aims to detect active computers based on the subject of the monitor screen and is expected to provide information about computers that are still active in the school laboratory. Based on the study's results, which detected 10 active computers, the 200-epoch model was selected with 100% accuracy at a speed of 2 ms. Twenty active computers were selected, with 200 epoch models achieving 95% accuracy at a speed of 6 ms per epoch. Thirty active computers were selected, with 100 epoch models achieving 96.67% accuracy at a speed of 3 ms.

Keywords: Deep Learning, Computers, School Laboratory, YOLOv8, Monitor Screen

Abstrak

Komputer merupakan perangkat yang memiliki perkembangan teknologi yang cukup pesat serta digunakan dalam dunia pendidikan. Penggunaan komputer yang tidak dimatikan dapat menyebabkan kerusakan pada komponen-komponen komputer serta penggunaan energi listrik dapat meningkat. Ketidakpatuhan siswa dalam mematikan komputer laboratorium sekolah ketika selesai digunakan menyebabkan guru melakukan pengecekan manual dengan mengunjungi setiap laboratorium komputer sekolah. Deep learning merupakan algoritma pembelajaran mesin yang menggunakan jaringan saraf tiruan. Deep learning biasanya digunakan untuk pengenalan citra, identifikasi suara, dan analisis pola data. Penelitian kali ini akan menggunakan metode Deep Learning yaitu YOLOv8 bertujuan untuk mendeteksi komputer aktif dengan subjek layar monitor serta diharapkan dapat memberikan informasi terkait komputer yang masih aktif di laboratorium sekolah. Dari hasil penelitian, untuk mendeteksi 10 komputer aktif dipilih model 200 epoch memiliki akurasi 100% dengan kecepatan 2 ms. Pada skenario 20 komputer aktif, model dengan 200 epoch memberikan akurasi 95% dengan kecepatan 6 ms. Sementara itu, untuk 30 komputer aktif, model dengan 100 epoch memiliki akurasi 96,67% dengan kecepatan 3 ms.

Kata Kunci: Deep Learning, Komputer, Laboratorium Sekolah, YOLOv8, Layar Monitor

1. PENDAHULUAN

Dalam era digital saat ini, penggunaan komputer tidak terpisahkan dari berbagai aspek kehidupan seperti di lingkungan perkantoran, pendidikan, dan industri. Komputer digunakan untuk berbagai aktivitas mulai dari pengolahan data, komunikasi serta bisnis. Keberadaan komputer semakin banyak digunakan oleh pengguna karena manfaatnya (Monalia et al., 2022). Pada dasarnya, beberapa bagian komputer memiliki jangka waktu hidup terbatas, yang dapat berkurang jika digunakan secara terus menerus. Salah satu cara untuk mencegah terjadinya kerusakan komponen komputer tersebut dengan mematikan komputer ketika selesai digunakan. Penggunaan komputer di laboratorium sekolah menjadi masalah mengenai bagaimana tingkat kepatuhan siswa dalam mematikan komputer setelah selesai digunakan. Masih terdapat siswa



yang belum patuh dalam mematikan komputer setelah digunakan, sehingga guru melakukan pemantauan secara manual dengan mengunjungi setiap laboratorium komputer untuk mengetahui informasi tersebut.

Berdasarkan permasalahan yang telah dijelaskan, dibutuhkan sebuah perangkat lunak yang mampu memberikan informasi kepada guru terkait komputer yang belum dimatikan. Penelitian saat ini menggunakan Artificial Intelligence, sebuah teknologi yang berkembang sangat pesat dan mampu menyelesaikan masalah yang sulit diselesaikan dengan mudah. Secara umum, Artificial Intelligence dibagi menjadi dua bagian yaitu Deep Learning dan Machine Learning (Raup et al., 2022).

Pada penelitian berjudul "Klasifikasi Jenis Burung Menggunakan Metode CNN dan Arsitektur ResNet-50" memperoleh nilai akurasi 98% dengan *optimizer* SGD (Alberto & Hermanto, 2023). Adapun penelitian lain untuk pengenalan telapak tangan menggunakan Convolutional Neural Network (CNN) menggunakan metode Alexnet memperoleh nilai akurasi 97,01% pada epoch ke-7 (Hardi & Sundari, 2023). Penelitian lain untuk mengenali wajah orang juga telah dilakukan menggunakan algoritma HOG dengan *fitcecoc multiclass* SVM yang memperoleh tingkat akurasi 98,5714% (Laia et al., 2023). Dalam penelitian berjudul "Object Detection untuk Menentukan Citra Buah-Buahan dengan Metode YOLO" memperoleh nilai akurasi mAP sebesar 91% (Saputra et al., 2023).

Pada penelitian kali ini akan menerapkan Deep Learning dengan model Convolutional Neural Network (CNN) yang dikenal dengan "YOLO" (You Only Look Once) untuk deteksi komputer aktif melalui layar monitor di laboratorium sekolah. YOLO memiliki tiga bagian utama yaitu *backbone*, *neck*, dan *head* (Ali & Zhang, 2024; Gupta et al., 2025; Sadrawi et al., 2023). Model YOLO memiliki bobot ringan sehingga dapat berjalan secara *real-time* pada perangkat keras dengan memanfaatkan GPU untuk meningkatkan efisiensi deteksi (Alfarizi et al., 2023; Betti & Tucci, 2023). Deteksi dan pengenalan *real-time* menjadi keunggulan utama YOLO dengan kemampuan pengenalan mencapai 45 *frame* per detik serta memberikan kinerja yang sangat baik dalam situasi waktu nyata (Luthfi, 2021; Yaseen, 2024). Deteksi dan pengenalan objek dilakukan oleh YOLO dengan cepat dan akurat dengan kecepatan mencapai 22 ms per gambar (Amin & Obermaisser, 2025; Mulyana & Rofik, 2022). Pada penelitian ini diharapkan dapat mencegah terjadinya kerusakan komputer dan diharapkan dapat membantu mengurangi biaya pemakaian listrik dari komputer yang belum dimatikan.

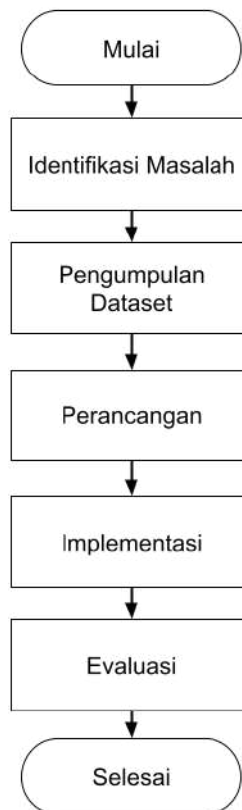
2. METODE PENELITIAN

Tahapan penelitian yang dilakukan untuk mendeteksi keberadaan komputer aktif menggunakan metode YOLOv8 ditunjukkan pada Gambar 1. Tahap awal yang dilakukan yaitu pengumpulan informasi yang berhubungan dengan topik penelitian, mempelajari jurnal yang berkaitan dengan deteksi keberadaan komputer aktif menggunakan metode You Only Look Once (YOLO). Tahap selanjutnya dilakukan pengumpulan *dataset* yang disusun secara *primer*. Data tersebut berupa gambar yang diambil dengan menggunakan *smartphone* Redmi Note 10S dengan spesifikasi kamera yaitu 64 Megapixel (MP). Pengambilan gambar dilakukan di laboratorium komputer, kemudian gambar yang diperoleh dipotong sesuai kebutuhan. Jumlah *dataset* yang terkumpul sebanyak 700 gambar, terdiri dari 490 data *train*, 140 data validasi, dan 70 data *test*. Contoh *dataset* yang diambil dalam penelitian ini ditunjukkan pada Gambar 2.

Pada tahap perancangan, dilakukan perancangan sistem yang diperlukan untuk melakukan penelitian yaitu penggunaan metode YOLO untuk deteksi keberadaan komputer aktif. Proses ini dimulai dengan membuat label dari setiap *dataset* yang telah di kumpulkan, sebagaimana yang ditunjukkan pada Gambar 3. Pengumpulan *dataset* dilakukan secara *primer*, kemudian dilakukan *labeling* dengan bantuan aplikasi Roboflow. Salah satu contoh proses pelabelan dapat dilihat pada Gambar 4. Setelah proses pelabelan selesai, *dataset* diekspor menjadi data *train* terdiri dari 490 gambar, data *test* terdiri dari 70 gambar, dan data validasi terdiri dari 140 gambar. Pemisahan *dataset* tersebut menghasilkan *source code* yang digunakan pada tahap pelatihan model. Proses pemisahan *dataset* ditunjukkan pada Gambar 5.



Pada tahap implementasi dilakukan dengan bantuan *Google Colab* dengan bahasa pemrograman *Python* menggunakan metode YOLO. Di mana akan melatih *dataset* yang telah disiapkan sebelumnya sehingga menghasilkan sebuah model. Model tersebut akan digunakan untuk *testing* deteksi keberadaan komputer aktif dalam pembuatan perangkat lunak berbasis *website* dengan bahasa pemrograman *Python* serta dengan bantuan *framework flask* dan *library OpenCV*. YOLO memproses gambar secara *real-time* pada empat puluh lima (45) *frames per second* (Khairunnas et al., 2021). Arsitektur YOLO terdiri atas dua puluh empat lapisan yang terhubung penuh, diikuti oleh empat lapisan dan dua lapisan yang terhubung penuh. YOLO menggambarkan identifikasi objek sebagai masalah regresi Tunggal, gambar ini akan dihubungkan dengan sebuah *bounding box* spasial yang terpisah dan probabilitas kelas yang terkait (Hutauruk et al., 2020).



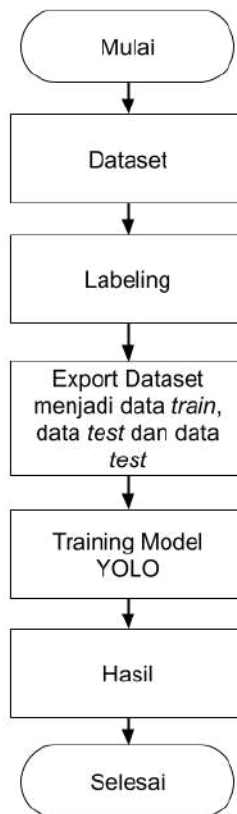
Gambar 1 Tahapan Penelitian



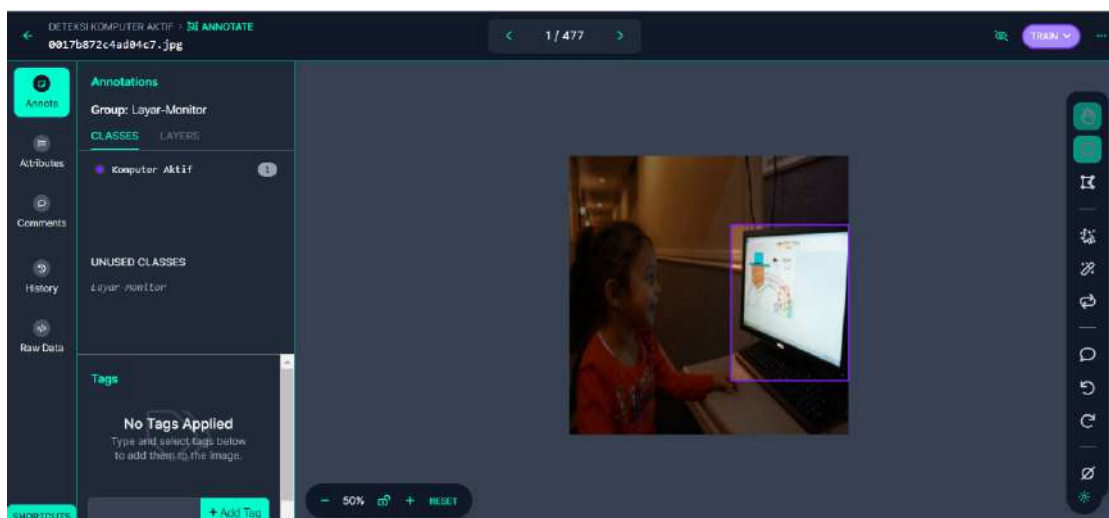
Gambar 2 Contoh Dataset Komputer



Proses dimulai dengan menambahkan dan menginstall *dependency* yang dibutuhkan, sebagaimana ditunjukkan pada potongan kode program yang ditunjukkan pada Gambar 6. Kemudian lakukan proses *training* model dengan epoch yang ditentukan, misalnya 100 epoch, menggunakan file data.yaml yang telah diunduh dan disimpan di lokasi yang ditentukan seperti yang ditunjukkan pada potongan kode program pada Gambar 7. Hasil dari proses *training* tersebut, seperti yang ditunjukkan pada Gambar 8, akan disimpan dan dapat diunduh pada folder *runs/detect/train/weights/last.pt*.

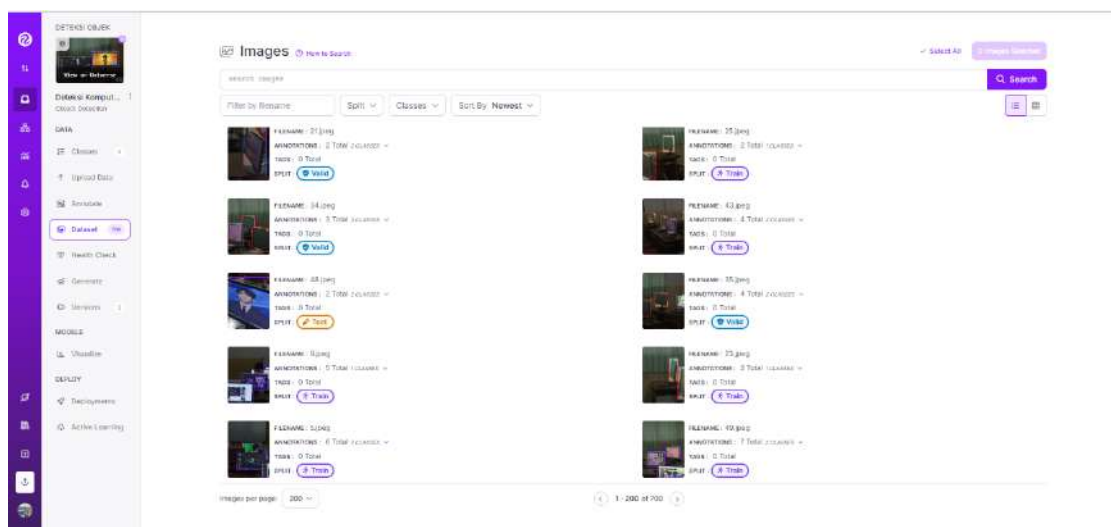


Gambar 3 Perancangan Model YOLO



Gambar 4 Contoh *Labeling* dengan Bantuan Roboflow



Gambar 5 Contoh Hasil *Export Dataset*

```
!pip install ultralytics==8.0.196
from IPython import display
display.clear_output()
import ultralytics ultralytics.checks()
from ultralytics import YOLO
from IPython.display import display, Image
```

Gambar 6 Kode Program Instalasi *Dependency*

```
%cd {HOME}
!yolo task=detect mode=train model=yolov8s.pt
data={dataset.location}/data.yaml epochs=100
```

Gambar 7 Kode Program Training Model

```
100 epochs completed in 0.356 hours.
Optimizer stripped from runs/detect/train/weights/last.pt, 22.5MB
Optimizer stripped from runs/detect/train/weights/best.pt, 22.5MB

Validating runs/detect/train/weights/best.pt...
Ultralytics YOLOv8.0.196 Python-3.10.12 torch-2.2.1+cu121 CUDA:0 (Tesla T4, 15102MiB)
Model summary (fused): 168 layers, 11126358 parameters, 0 gradients, 28.4 GFLOPs
```

Class	Images	Instances	Box(P)	R	mAP50	mAP50-95
all	140	622	0.955	0.957	0.982	0.796
Komputer Aktif	140	344	0.943	0.942	0.975	0.805
Komputer TidakAktif	140	278	0.966	0.971	0.99	0.786

```
Speed: 0.3ms preprocess, 6.5ms inference, 0.0ms loss, 6.6ms postprocess per image
```

Gambar 8 Hasil Model Proses Pelatihan

Tabel 1 Confusion Matrix

		Aktual	
		Positif	Negatif
Prediksi	Positif	True Positive (TP)	False Positive (FP)
	Negatif	False Negative (FN)	True Negative (TN)

Setelah implementasi, tahap evaluasi dilakukan menggunakan Confusion Matrix. Confusion Matrix memberikan pemahaman tentang kesalahan model dengan mengungkapkan saling ketergantungan antara kelas yang diberi label dengan yang diprediksi (Erbani et al., 2024).



Evaluasi ini melibatkan beberapa komponen yaitu akurasi, *precision*, dan *recall*. *Accuracy* digunakan untuk menggambarkan akurasi dari suatu model yang dalam mengklasifikasikan dengan benar (Romadloni et al., 2022). Akurasi merupakan prediksi data benar (positif dan negatif) dengan keseluruhan data (Rahayu et al., 2021), seperti yang ditunjukkan pada Pers. (1). *Precision* digunakan untuk menggambarkan akurasi antara data yang diminta dengan hasil prediksi yang diberikan oleh model (Romadloni et al., 2022). *Precision* merupakan prediksi data benar positif dibandingkan dengan keseluruhan hasil yang diprediksi positif (Rahayu et al., 2021), seperti yang ditunjukkan pada Pers. (2). *Recall* digunakan untuk menggambarkan keberhasilan model dalam menemukan Kembali sebuah informasi (Romadloni et al., 2022). *Recall* merupakan prediksi data benar positif yang berbeda dari keseluruhan data benar positif (Rahayu et al., 2021), seperti yang ditunjukkan pada Pers. (3).

$$Accuracy = \frac{TN + TP}{TN + FN + TP + FP} \quad (1)$$

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

3. HASIL DAN PEMBAHASAN

Pengujian dilakukan menggunakan metode YOLOv8 dengan beberapa epoch yaitu 100 epoch, 150 epoch, dan 200 epoch. Pengujian tersebut dilakukan pada salah satu laboratorium komputer yang memiliki jumlah komputer 36 komputer dengan tiga kondisi yaitu 10, 20, dan 30 komputer aktif. Dari hasil pengujian tersebut dilakukan analisis, kemudian akan dipilih epoch yang memiliki akurasi tertinggi pada setiap kondisi dan waktu tercepat.

3.1 Pengujian 100 Epoch

Pada pengujian dengan 100 epoch, hasil pengujian dari ketiga kondisi dapat dilihat pada Tabel 2 dan Gambar 9. Pada kondisi pertama, jumlah komputer aktif sebanyak 10 unit, seluruhnya berhasil terdeteksi dengan benar. Jumlah komputer tidak aktif sebanyak 26 unit, dengan 24 unit terdeteksi benar dan 2 unit tidak terdeteksi. Pada kondisi kedua, jumlah komputer aktif sebanyak 20 unit, namun hanya 17 unit yang berhasil terdeteksi, sedangkan 3 unit tidak terdeteksi. Jumlah komputer tidak aktif sebanyak 16 unit, dengan 13 unit terdeteksi benar dan 3 unit tidak terdeteksi. Sementara itu, pada kondisi ketiga, jumlah komputer aktif sebanyak 30 unit, dengan 29 unit berhasil terdeteksi dan 1 unit tidak terdeteksi. Jumlah komputer tidak aktif sebanyak 6 unit, dengan 4 unit terdeteksi benar dan 2 unit tidak terdeteksi.

Tabel 2 Pengujian 100 Epoch

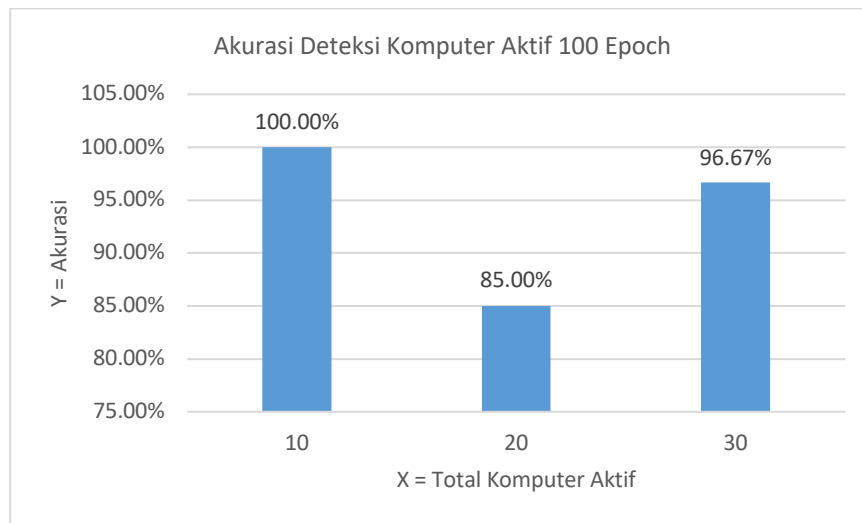
Kondisi	Komputer Aktif	Komputer Tidak Aktif	Terdeteksi Aktif	Terdeteksi Tidak Aktif	Tidak Terdeteksi	Terdeteksi Salah	Lama Deteksi
1	10	26	10	24	2	-	4 ms
2	20	16	17	13	6	-	2 ms
3	30	6	29	4	3	-	3 ms



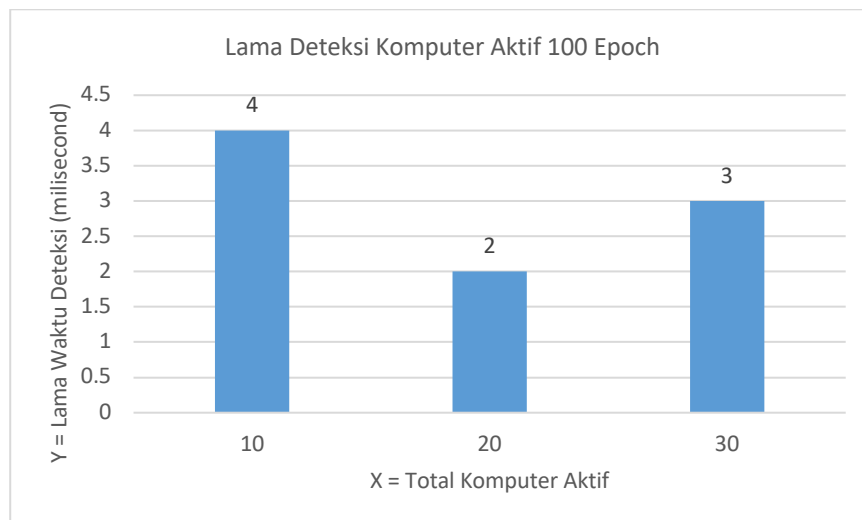
Gambar 9 Pengujian 100 Epoch



Nilai akurasi pada pengujian dengan 100 epoch menunjukkan hasil yang bervariasi pada setiap kondisi. Pada kondisi pertama, dengan total 10 komputer aktif yang seluruhnya berhasil terdeteksi, diperoleh nilai akurasi sebesar 100%. Pada kondisi kedua, dengan total 20 komputer aktif namun hanya 17 komputer yang berhasil terdeteksi, diperoleh nilai akurasi sebesar 85%. Sementara itu, pada kondisi ketiga dengan total 30 komputer aktif dan 29 komputer berhasil terdeteksi, nilai Akurasi yang dicapai adalah 96,67%. Visualisasi nilai akurasi dari ketiga kondisi tersebut dapat dilihat pada Gambar 10. Selanjutnya, waktu deteksi juga berbeda pada setiap kondisi pengujian. Pada kondisi pertama, dengan total 10 komputer aktif, waktu deteksi yang dibutuhkan adalah 4 ms. Pada kondisi kedua, dengan 20 komputer aktif, waktu deteksi tercatat 2 ms. Sedangkan pada kondisi ketiga, dengan 30 komputer aktif, waktu deteksi yang dicatat adalah 3 ms. Perbandingan lama deteksi pada ketiga kondisi tersebut dapat dilihat pada Gambar 11.



Gambar 10 Diagram Akurasi Deteksi Komputer Aktif 100 Epoch



Gambar 11 Diagram Lama Deteksi Komputer Aktif 100 Epoch

3.2 Pengujian 150 Epoch

Pada pengujian dengan 150 epoch, hasil pengujian dari ketiga kondisi dapat dilihat pada Tabel 3 dan Gambar 12. Pada kondisi pertama, jumlah komputer aktif sebanyak 10 unit, seluruhnya berhasil terdeteksi dengan benar. Jumlah komputer tidak aktif sebanyak 26 unit, dengan 25 unit



terdeteksi benar dan 1 unit terdeteksi salah. Pada kondisi kedua, jumlah komputer aktif sebanyak 20 unit, dengan 19 unit berhasil terdeteksi dan 1 unit tidak terdeteksi. Jumlah komputer tidak aktif sebanyak 16 unit, dengan 14 unit terdeteksi benar, 1 unit terdeteksi salah, dan 2 unit tidak terdeteksi. Sementara itu, pada kondisi ketiga, jumlah komputer aktif sebanyak 30 unit, seluruhnya berhasil terdeteksi dengan benar. Jumlah komputer tidak aktif sebanyak 6 unit, dan seluruhnya terdeteksi benar.

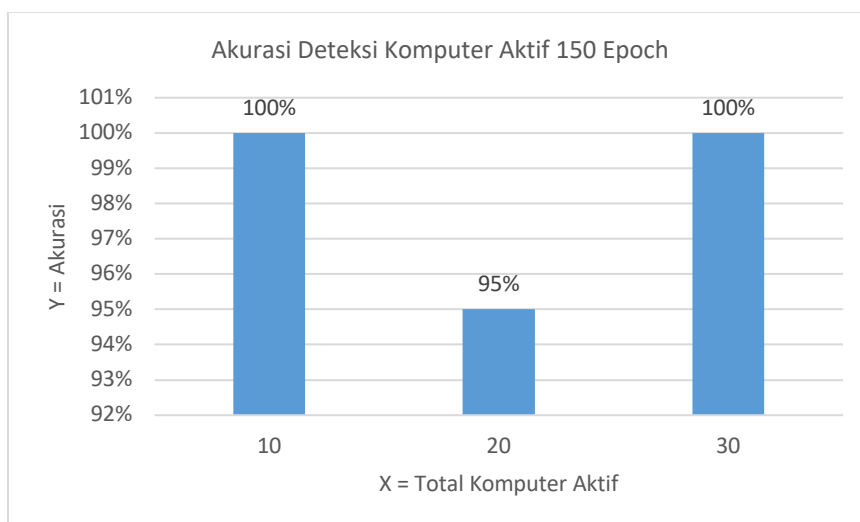
Tabel 3 Pengujian 150 Epoch

Kondisi	Komputer Aktif	Komputer Tidak Aktif	Terdeteksi Aktif	Terdeteksi Tidak Aktif	Tidak Terdeteksi	Terdeteksi Salah	Lama Deteksi
1	10	26	10	25	-	1	5 ms
2	20	16	19	14	2	1	7 ms
3	30	6	30	6	-	-	3 ms



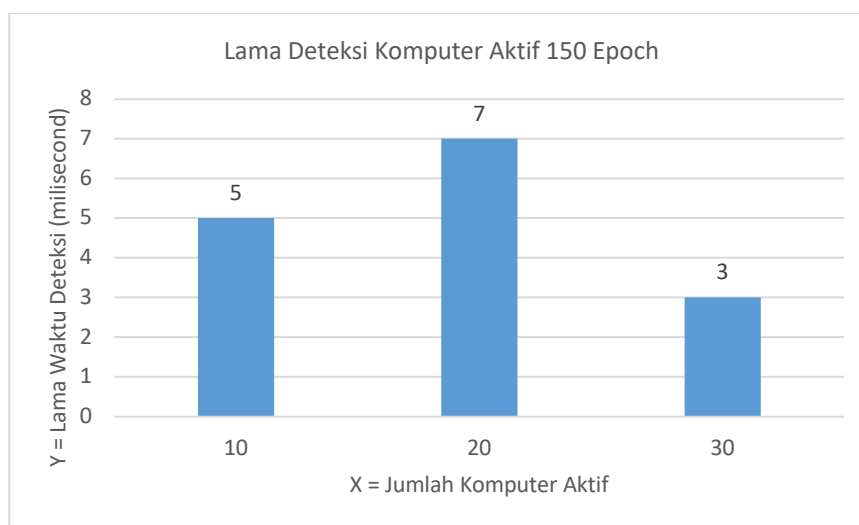
Gambar 12 Pengujian 150 Epoch

Nilai akurasi pada pengujian dengan 150 epoch menunjukkan hasil yang sangat baik. Pada kondisi pertama, dengan total 10 komputer aktif yang seluruhnya berhasil terdeteksi, diperoleh nilai akurasi sebesar 100%. Pada kondisi kedua, dengan total 20 komputer aktif namun hanya 19 komputer yang berhasil terdeteksi, diperoleh nilai akurasi sebesar 95%. Sementara itu, pada kondisi ketiga dengan total 30 komputer aktif yang seluruhnya berhasil terdeteksi, nilai akurasi mencapai 100%. Visualisasi nilai akurasi dari ketiga kondisi tersebut dapat dilihat pada Gambar 13. Selanjutnya, waktu deteksi juga berbeda pada setiap kondisi pengujian. Pada kondisi pertama, dengan total 10 komputer aktif, waktu deteksi yang dibutuhkan adalah 5 ms. Pada kondisi kedua, dengan 20 komputer aktif, waktu deteksi tercatat 7 ms. Sedangkan pada kondisi ketiga, dengan 30 komputer aktif, waktu deteksi yang dicatat adalah 3 ms. Perbandingan lama deteksi pada ketiga kondisi tersebut dapat dilihat pada Gambar 14.



Gambar 13 Diagram Akurasi Deteksi Komputer Aktif 150 Epoch





Gambar 14 Diagram Lama Deteksi Komputer Aktif 150 Epoch

3.3 Pengujian 200 Epoch

Pada pengujian dengan 200 epoch, hasil pengujian dari ketiga kondisi dapat dilihat pada Tabel 4 dan Gambar 15. Pada kondisi pertama, jumlah komputer aktif sebanyak 10 unit, seluruhnya berhasil terdeteksi dengan benar. Jumlah komputer tidak aktif sebanyak 26 unit, dengan 21 unit terdeteksi benar dan 5 unit tidak terdeteksi. Pada kondisi kedua, jumlah komputer aktif sebanyak 20 unit, dengan 19 unit berhasil terdeteksi dan 1 unit tidak terdeteksi. Jumlah komputer tidak aktif sebanyak 16 unit, dengan 15 unit terdeteksi benar dan 1 unit tidak terdeteksi. Sementara itu, pada kondisi ketiga, jumlah komputer aktif sebanyak 30 unit, dengan 27 unit berhasil terdeteksi dan 3 unit tidak terdeteksi. Jumlah komputer tidak aktif sebanyak 6 unit, dengan 5 unit terdeteksi benar dan 1 unit tidak terdeteksi.

Tabel 4 Pengujian 200 Epoch

Kondisi	Komputer Aktif	Komputer Tidak Aktif	Terdeteksi Aktif	Terdeteksi Tidak Aktif	Tidak Terdeteksi	Terdeteksi Salah	Lama Deteksi
1	10	26	10	21	5	-	2 ms
2	20	16	19	15	2	-	6 ms
3	30	6	27	5	4	-	4 ms

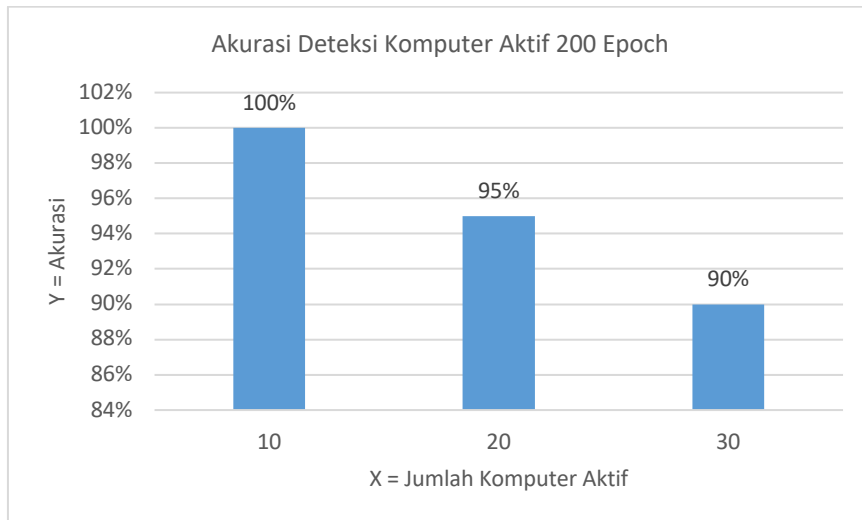


Gambar 15 Pengujian 200 Epoch

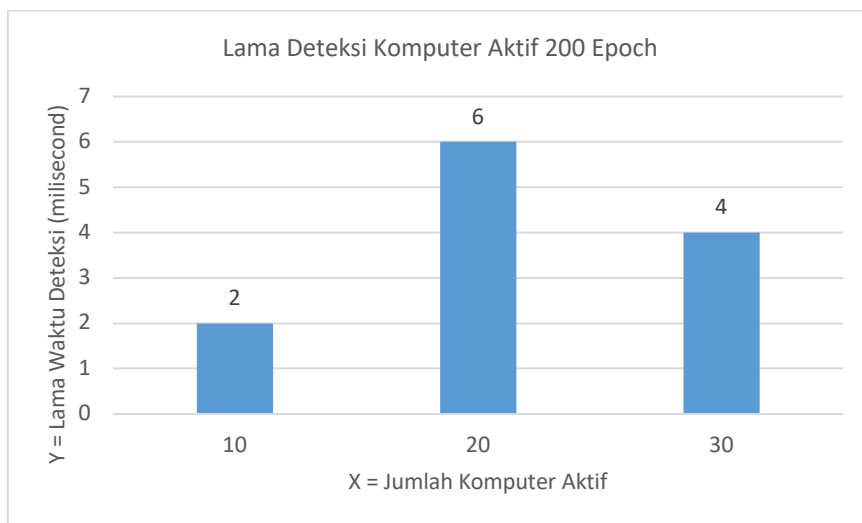
Nilai akurasi pada pengujian dengan 200 epoch juga menunjukkan hasil yang beragam pada setiap kondisi. Pada kondisi pertama, dengan total 10 komputer aktif yang seluruhnya berhasil terdeteksi, diperoleh nilai akurasi sebesar 100%. Pada kondisi kedua, dengan total 20 komputer aktif namun hanya 19 komputer yang berhasil terdeteksi, diperoleh nilai akurasi sebesar 95%. Sementara itu, pada kondisi ketiga dengan total 30 komputer aktif dan hanya 27 komputer berhasil terdeteksi, nilai akurasi yang dicapai adalah 90%. Visualisasi nilai akurasi dari ketiga kondisi tersebut dapat dilihat pada Gambar 16. Selanjutnya, waktu deteksi juga bervariasi pada



setiap kondisi. Pada kondisi pertama, dengan total 10 komputer aktif, waktu deteksi yang dibutuhkan adalah 2 ms. Pada kondisi kedua, dengan 20 komputer aktif, waktu deteksi tercatat 6 ms. Sedangkan pada kondisi ketiga, dengan 30 komputer aktif, waktu deteksi yang dicatat adalah 4 ms. Perbandingan lama deteksi pada ketiga kondisi tersebut dapat dilihat pada Gambar 17.



Gambar 16 Diagram Akurasi Deteksi Komputer Aktif 200 Epoch



Gambar 17 Diagram Lama Deteksi Komputer Aktif 200 Epoch

3.4 Hasil Pengujian

Hasil pengujian model YOLO pada 100, 150, dan 200 epoch menunjukkan bahwa performa model berbeda-beda sesuai jumlah komputer aktif. Untuk 10 komputer aktif, model 200 epoch memberikan akurasi tertinggi 100% dengan waktu deteksi tercepat 2 ms. Pada 20 komputer aktif, model 200 epoch juga unggul dengan akurasi 95%, meskipun waktu deteksi sedikit lebih lama (6 ms). Untuk 30 komputer aktif, model 100 epoch dipilih karena memberikan kombinasi terbaik antara akurasi (96,67%) dan waktu deteksi (3 ms), meskipun model 150 epoch mampu mencapai akurasi 100%, terdapat beberapa komputer yang salah atau tidak terdeteksi. Secara keseluruhan, hasil pengujian menekankan *trade-off* antara jumlah epoch, akurasi, dan kecepatan deteksi. Model dengan epoch lebih tinggi cenderung lebih akurat untuk jumlah komputer yang



lebih sedikit, sedangkan untuk jumlah komputer yang lebih banyak, model dengan epoch lebih rendah menawarkan keseimbangan optimal antara akurasi dan kecepatan.

4. KESIMPULAN

Berdasarkan pengujian di salah satu laboratorium komputer dengan menggunakan beberapa model YOLO, diperoleh kesimpulan sebagai berikut. Untuk pengujian dengan jumlah komputer aktif 10 unit, dipilih model 200 epoch yang memiliki akurasi 100% dalam mengenali komputer aktif dengan waktu 2 ms. Pengujian dengan jumlah komputer aktif 20 unit, dipilih model 200 epoch yang memiliki akurasi 95% dalam mengenali komputer aktif dengan waktu 6 ms. Pengujian dengan jumlah komputer aktif 30 unit, model 150 epoch memiliki akurasi tertinggi 100% dalam mengenali komputer aktif meskipun masih ada beberapa komputer terdeteksi salah. Oleh karena itu, dipilih model 100 epoch dengan akurasi 96,67% dengan waktu 3 ms sebagai model terbaik untuk kondisi ini.

DAFTAR PUSTAKA

- Alberto, J., & Hermanto, D. (2023). Bird Species Classification Using CNN Method and ResNet-50 Architecture. *JATISI (Jurnal Teknik Informatika dan Sistem Informasi)*, 10(3), 34–36. <https://doi.org/10.35957/jatisi.v10i3.4558>
- Alfarizi, D. N., Pangestu, R. A., Aditya, D., Setiawan, M. A., & Rosyani, P. (2023). Penggunaan Metode YOLO pada Deteksi Objek: Sebuah Tinjauan Literatur Sistematis. *Jurnal Artificial Intelligent dan Sistem Penunjang Keputusan*, 1(1), 54–63. <https://jurnalmahasiswa.com/index.php/aidanspk/article/view/144>
- Ali, M. L., & Zhang, Z. (2024). The YOLO Framework: A Comprehensive Review of Evolution, Applications, and Benchmarks in Object Detection. *Computers*, 13(12), Article ID: 336. <https://doi.org/10.3390/computers13120336>
- Amin, R. Al, & Obermaisser, R. (2025). Real-Time Object Detection and Classification Using YOLO for Edge FPGAs. *ArXiv Preprint*. <http://arxiv.org/abs/2507.18174>
- Betti, A., & Tucci, M. (2023). YOLO-S: A Lightweight and Accurate YOLO-like Network for Small Target Detection in Aerial Imagery. *Sensors*, 23(4), Article ID: 1865. <https://doi.org/10.3390/s23041865>
- Erbani, J., Portier, P.-É., Egyed-Zsigmond, E., & Nurbakova, D. (2024). Confusion Matrices: A Unified Theory. *IEEE Access*, 12, 181372–181419. <https://doi.org/10.1109/ACCESS.2024.3507199>
- Gupta, C., Gill, N. S., Gulia, P., Kumar, A., Karamti, H., & Moges, D. M. (2025). An Optimized YOLO NAS Based Framework for Realtime Object Detection. *Scientific Reports*, 15(1), Article ID: 32903. <https://doi.org/10.1038/s41598-025-17919-w>
- Hardi, N., & Sundari, J. (2023). Pengenalan Telapak Tangan Menggunakan Convolutionall Neural Network (CNN). *Reputasi: Jurnal Rekayasa Perangkat Lunak*, 4(1), 10–15. <https://doi.org/10.31294/reputasi.v4i1.1951>
- Hutauruk, J. S. W., Matulatan, T., & Hayaty, N. (2020). Deteksi Kendaraan Secara Real Time Menggunakan Metode YOLO Berbasis Android. *Jurnal Sustainable: Jurnal Hasil Penelitian dan Industri Terapan*, 9(1), 8–14. <https://doi.org/10.31629/sustainable.v9i1.1401>
- Khairunnas, K., Yuniarno, E. M., & Zaini, A. (2021). Pembuatan Modul Deteksi Objek Manusia Menggunakan Metode YOLO untuk Mobile Robot. *Jurnal Teknik ITS*, 10(1), A50–A55. <https://doi.org/10.12962/j23373539.v10i1.61622>
- Laia, F. H., Rosnelly, R., Naswar, A., Buulolo, K., & Lase, M. C. (2023). Deteksi Pengenalan Wajah Orang Berbasis AI Computer Vision. *Jurnal Teknologi Informasi Mura*, 15(1), 62–72. <https://jurnal.univbinainsan.ac.id/index.php/jti/article/view/2024>
- Luthfi, A. (2021). Pendeteksi Senjata Berbahaya pada Percobaan Tindakan Kriminal dengan Menggunakan Metode YOLO (You Only Look Once) [Universitas Islam Negeri Sultan Syarif Kasim Riau]. In *Pekanbaru*. <https://repository.uin-suska.ac.id/41304/>
- Monalia, M., Asfiyanti, N. A., & Putri, S. E. (2022). Computers and Information Technology as a Source of Learning Media for Elementary School Teachers. *International Journal of Natural Science and Engineering*, 5(3), 96–103. <https://doi.org/10.23887/ijnse.v5i3.41862>



- Mulyana, D. I., & Rofik, M. A. (2022). Implementasi Deteksi Real Time Klasifikasi Jenis Kendaraan di Indonesia Menggunakan Metode YOLOV5. *Jurnal Pendidikan Tambusai*, 6(3), 13971–13982. <https://doi.org/10.31004/jptam.v6i3.4825>
- Rahayu, W. I., Prianto, C., & Novia, E. A. (2021). Perbandingan Algoritma K-Means dan Naïve Bayes untuk Memprediksi Prioritas Pembayaran Tagihan Rumah Sakit Berdasarkan Tingkat Kepentingan pada PT. Pertamina (Persero). *Jurnal Teknik Informatika*, 13(2), 1–8. <https://ejurnal.ulbi.ac.id/index.php/informatika/article/view/1383>
- Raup, A., Ridwan, W., Khoeriyah, Y., Supiana, S., & Zaqiah, Q. Y. (2022). Deep Learning dan Penerapannya dalam Pembelajaran. *JlIP - Jurnal Ilmiah Ilmu Pendidikan*, 5(9), 3258–3267. <https://doi.org/10.54371/jiip.v5i9.805>
- Romadloni, P., Adhi Kusuma, B. A., & Baihaqi, W. M. (2022). Komparasi Metode Pembelajaran Mesin untuk Implementasi Pengambilan Keputusan dalam Menentukan Promosi Jabatan Karyawan. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 6(2), 622–628. <https://doi.org/10.36040/jati.v6i2.5238>
- Sadrawi, M., Fugaha, D. R., Heerlie, D. M., Lorell, J., Gautama, N. R. P., & Aminuddin, M. Z. (2023). Artificial Intelligence Based Brain Tumor Localization Using YOLOv5. *Indonesian Journal of Life Sciences*, 5(1), 1–9. <https://doi.org/10.54250/ijls.v5i01.176>
- Saputra, D. H., Imran, B., & Juhartini. (2023). Object Detection untuk Mendeteksi Citra Buah-Buahan Menggunakan Metode YOLO. *Jurnal Kecerdasan Buatan dan Teknologi Informasi*, 2(2), 70–80. <https://doi.org/10.69916/jkbt.v2i2.18>
- Yaseen, M. (2024). What is YOLOv8: An In-Depth Exploration of the Internal Features of the Next-Generation Object Detector. *ArXiv Preprint*. <http://arxiv.org/abs/2408.15857>



Klasifikasi Hewan Anjing, Kucing, dan Harimau Menggunakan Metode Convolutional Neural Network (CNN)

Murdifin Murdifin ^{(1)*}, Shofwatul Uyun ⁽²⁾

Departemen Magister Informatika, UIN Sunan Kalijaga, Yogyakarta, Indonesia
e-mail : 23206052007@student.uin-suka.ac.id, shofwatul.uyun@uin-suka.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 27 Agustus 2024, direvisi 15 Februari 2025, diterima 17 Februari 2025, dan dipublikasikan 30 September 2025.

Abstract

Animal classification is a complex challenge due to variations in shape, color, and patterns across species. Traditional methods, which rely on manual feature extraction, are often ineffective in handling such complexities. Therefore, this study employs Convolutional Neural Networks (CNNs) as a more accurate approach for automatic feature extraction and image classification. This research aims to develop an animal image classification model, specifically for dogs, cats, and tigers, utilizing CNNs. The dataset consists of 4,800 images obtained from Kaggle, which were divided into training, testing, and validation sets. The CNN model was built using TensorFlow/Keras, trained for 50 epochs, and evaluated using accuracy, precision, recall, F1-score, and a confusion matrix. The experimental results show that the model achieved an overall accuracy of 88%, with the highest performance in tiger classification (99% accuracy). However, distinguishing between dogs and cats remains a challenge, with an accuracy of 81% for both classes. The findings indicate that CNNs are effective in automatically classifying animal images, although challenges persist in differentiating visually similar species. This study lays the groundwork for further enhancements, such as refining the model architecture or utilizing data augmentation techniques to boost classification accuracy.

Keywords: *Animal Classification, Convolutional Neural Network (CNN), Deep Learning, Image Processing, Species Identification*

Abstrak

Klasifikasi hewan merupakan tantangan yang kompleks karena variasi bentuk, warna, dan pola antarspesies. Metode tradisional yang mengandalkan fitur manual sering kali kurang efektif dalam menangani kompleksitas ini. Oleh karena itu, penelitian ini menggunakan Convolutional Neural Network (CNN) sebagai pendekatan yang lebih akurat dalam ekstraksi fitur otomatis dan klasifikasi gambar hewan. Penelitian ini bertujuan untuk mengembangkan model klasifikasi gambar hewan, khususnya anjing, kucing, dan harimau, dengan memanfaatkan CNN. Dataset yang digunakan terdiri atas 4.800 gambar yang diperoleh dari Kaggle, kemudian dibagi menjadi data pelatihan, pengujian, dan validasi. Model CNN dibangun menggunakan TensorFlow/Keras, dilatih selama 50 epoch, dan dievaluasi menggunakan metrik akurasi, presisi, *recall*, F1-score, dan confusion matrix. Hasil eksperimen menunjukkan bahwa model mencapai akurasi keseluruhan sebesar 88%, dengan performa terbaik pada klasifikasi harimau (akurasi 99%). Namun, ditemukan tantangan dalam membedakan anjing dan kucing, dengan akurasi masing-masing 81%. Hasil penelitian ini menunjukkan bahwa CNN merupakan metode yang efektif dalam mengklasifikasikan gambar hewan secara otomatis, meskipun masih terdapat kesulitan dalam membedakan spesies yang memiliki karakteristik visual serupa. Studi ini memberikan dasar bagi pengembangan lebih lanjut, seperti peningkatan arsitektur model atau penggunaan teknik augmentasi data untuk meningkatkan akurasi klasifikasi.

Kata Kunci: *Klasifikasi Hewan, Convolutional Neural Network (CNN), Deep Learning, Pengolahan Citra, Identifikasi Spesies*

1. PENDAHULUAN

Klasifikasi adalah kemampuan untuk memberi nama dan mengidentifikasi menurut ukuran, tampilan, atau karakteristik lain (A'yun & Erman, 2019). Kemampuan mengklasifikasi merupakan



Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

salah satu dari kemampuan proses sains yaitu mengelompokkan fakta atau data berdasarkan perbedaan, persamaan, dan karakteristik yang kontras. Tujuan klasifikasi adalah untuk memudahkan pengenalan dan perbandingan (Mariyanti et al., 2022; Mohite et al., 2025). Klasifikasi menggunakan morfologi sebagai dasar utama untuk mengelompokkan, morfologi merujuk pada penampilan luar yang dapat diamati dengan mata telanjang (Adrianto et al., 2022; Ricard et al., 2023).

Hewan telah mengalami banyak perkembangan sepanjang sejarah, hewan juga merupakan makhluk yang sangat dekat pada manusia sesuai dengan jenisnya seperti hewan ternak, hewan peliharaan, dan bahkan hewan liar pun bisa dekat sama manusia. Meskipun memiliki banyak ciri dan karakteristik yang membedakan satu sama lain, beberapa jenis hewan sulit untuk dibedakan, terutama bagi yang baru mengenal dunia hewan. Klasifikasi hewan sebagai bidang ilmu menawarkan berbagai aplikasi penting seperti dalam konservasi alam, penelitian biologi, dan manajemen hewan domestik. Namun, masalah utama dalam klasifikasi hewan adalah kompleksitas dan variasi besar hewan (Norouzzadeh et al., 2018). metode tradisional yang mengandalkan fitur manual seringkali tidak cukup efektif untuk menangani keragaman (Tan & Le, 2019). Untuk mengatasi hal ini, pendekatan otomatis dengan tingkat akurasi yang tinggi menjadi solusi yang diperlukan (Samudra et al., 2023).

Perkembangan teknologi pengenalan citra dan kecerdasan buatan dalam beberapa dekade terakhir telah membuka banyak peluang baru dalam berbagai bidang, seperti klasifikasi dan identifikasi hewan (He et al., 2020). Fukushima (1980), seorang peneliti dari NHK Broadcasting Science Research Laboratories, Kinuta, Setagaya, Tokyo, Jepang, berhasil mengembangkan CNN pertama kali dengan nama NeoCognitron. Penelitian ini akan dilakukan pengelompokkan jenis hewan menggunakan metode Convosional Neural Network (CNN). Metode Convolution Neural Network merupakan metode deep learning yang dapat melakukan proses pembelajaran mandiri untuk pengenalan dan ekstraksi objek (Ahmad et al., 2023). Algoritma CNN diketahui dapat menghasilkan tingkat akurasi yang signifikan dikarenakan metode ini memiliki kedalaman jaringan yang tinggi serta mampu mempelajari sendiri fitur yang terdapat pada citra yang kompleks (Suhendar et al., 2023). Kemampuan yang luar biasa dalam pengenalan pola dan klasifikasi citra secara otomatis telah ditunjukkan oleh Convolutional Neural Network (CNN) (Younis et al., 2022).

Penelitian yang dilakukan oleh Leovincen & Yoannita (2023) dengan judul “Klasifikasi Ras Anjing Berdasarkan Citra Menggunakan Convolutional Neural Network” mengkategorikan 120 ras anjing menggunakan model Convolutional Neural Network (CNN) dengan arsitektur ResNet-50 dan optimizer Adam. Dataset yang digunakan terdiri dari 20.580 gambar yang dibagi menjadi data latih, validasi, dan uji dengan rasio 60:20:20. Gambar-gambar tersebut diubah ukurannya menjadi 224x224 piksel. Penelitian ini mencapai tingkat akurasi sebesar 99,35%.

Penelitian yang dilakukan oleh Micheal & Hartati (2022), berjudul “Klasifikasi Spesies Kupu-Kupu Menggunakan Metode Convolutional Neural Network”, bertujuan untuk mengklasifikasikan spesies kupu-kupu dengan menggunakan metode Convolutional Neural Network (CNN) dengan arsitektur VGG-16 dan LeNet serta optimizer Adam, Adagrad, dan SGD. Dataset penelitian ini berisi 5.455 citra yang dibagi menjadi 4.955 data latih, 250 data uji, dan 250 data validasi. Dataset ini kemudian di-augmentasi sehingga jumlah data latih menjadi 8.000 untuk setiap kelas dan data uji menjadi 800 untuk setiap kelas. Hasil penelitian menunjukkan bahwa tingkat akurasi tertinggi dicapai dengan menggunakan optimizer Adam, dengan VGG-16 mencapai akurasi sebesar 93% dan LeNet mencapai akurasi sebesar 67%.

Penelitian yang dilakukan oleh Laksono et al. (2024) dengan judul “Identifikasi Jenis Ikan Cupang Berdasarkan Gambar Menggunakan Metode Convolutional Neural Network” bertujuan untuk mengidentifikasi jenis ikan cupang. Sistem ini menggunakan metode Convolutional Neural Network (CNN), algoritma deep learning dengan arsitektur keras sequential yang memiliki 1.424.403 parameter, dan biasanya digunakan untuk klasifikasi citra. Dataset yang digunakan terdiri dari 330 data, dengan 300 data latih dan 30 data uji yang terbagi dalam tiga kelas. Sistem



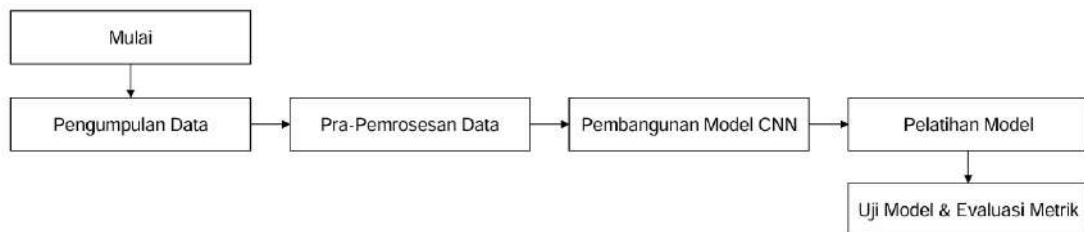
yang dikembangkan berhasil mengidentifikasi tiga jenis ikan cupang dengan akurasi 97% setelah 10 epoch, 93% setelah 15 epoch, dan mencapai akurasi tertinggi 100% setelah 20 epoch.

Tujuan dari penelitian ini adalah untuk mengembangkan model klasifikasi hewan menggunakan CNN, khususnya untuk hewan seperti anjing, kucing, dan harimau. Diharapkan bahwa dengan menggunakan kemampuan CNN untuk mengekstraksi dan mengklasifikasikan fitur otomatis, akan diciptakan model yang sangat akurat dan kuat terhadap berbagai variasi dalam gambar hewan. Perbedaan penelitian ini dengan penelitian sebelumnya yaitu dalam hal pendekatan metodologi, kompleksitas objek, dan tantangan klasifikasi. Leovincen & Yoannita (2023) mengklasifikasikan 120 ras anjing menggunakan ResNet-50, Micheal & Hartati (2022) mengklasifikasikan spesies kupu-kupu dengan VGG-16 dan LeNet, sementara Laksono et al. (2024) mengidentifikasi jenis ikan cupang dengan Keras Sequential. Ketiga penelitian tersebut berfokus pada satu jenis hewan dengan karakteristik serupa dalam satu spesies. Sebaliknya, penelitian ini mengembangkan model CNN untuk klasifikasi lintas spesies dengan tingkat kemiripan visual yang tinggi, yaitu anjing, kucing, dan harimau. Tantangan utama adalah membedakan fitur spesifik antarspesies yang sering tertukar, tidak hanya dalam satu kategori hewan. Selain itu, penelitian ini juga menganalisis metrik evaluasi yang lebih mendalam, termasuk confusion matrix, presisi, recall, dan F1-score, untuk meningkatkan generalisasi model dalam klasifikasi gambar yang lebih kompleks.

2. METODE PENELITIAN

2.1 Tahapan Penelitian

Tahapan Penelitian adalah hubungan antara konsep-konsep penelitian yang akan dilakukan (Hidayatullah et al., 2022). Tahapan pada penelitian ini mencakup pengumpulan dataset gambar hewan, proses pra-pemrosesan data untuk mempersiapkan dataset dan augmentasi, pembangunan arsitektur CNN menggunakan platform TensorFlow/Keras, pelatihan model untuk mengenali pola-pola yang signifikan dalam dataset, evaluasi kinerja model dengan data testing, serta analisis mendalam terhadap hasil yang dihasilkan. Gambar tahapan penelitian ditunjukkan pada Gambar 1.



Gambar 1 Tahapan Penelitian

- 1) **Pengumpulan Data:** Penelitian ini menggunakan dataset 4.800 gambar dari tiga kategori: anjing, kucing, dan harimau. Dataset diperoleh dari Kaggle, dibagi menjadi 3.000 gambar untuk pelatihan, 1.500 untuk pengujian, dan 300 untuk validasi. Augmentasi data seperti rotasi, *flipping*, penyesuaian kecerahan, kontras, *zooming*, dan *cropping* diterapkan untuk meningkatkan generalisasi model.
- 2) **Pra-Pemrosesan Data:** Gambar dikonversi ke resolusi 224×224 piksel dan dinormalisasi ke rentang [0,1] untuk mempercepat konvergensi. Gaussian Blur digunakan untuk mengurangi *noise* dan meningkatkan kualitas data (K et al., 2023; Supiyandi et al., 2024).
- 3) **Pembangunan Model CNN:** Model CNN terdiri dari empat lapisan konvolusi dengan filter 3×3, diikuti pooling 2×2, batch normalization, dan dropout 0,5 untuk mencegah overfitting. Fully connected layer dengan aktivasi ReLU di hidden layer dan softmax pada output digunakan untuk klasifikasi.
- 4) **Pelatihan Model:** Pelatihan dilakukan dengan TensorFlow/Keras menggunakan optimizer Adam (learning rate 0,001), categorical cross-entropy loss, batch size 32, dan 50 epoch. Early



stopping diterapkan untuk menghindari overfitting, dengan akurasi dan loss yang terus meningkat dan menurun.

- 5) Uji Model dan Evaluasi Metrik: Pengujian dilakukan dengan data testing 1.500 gambar. Evaluasi metrik mencakup akurasi, presisi, recall, F1-score, dan categorical cross-entropy loss, dihitung menggunakan rumus pada Pers. (1) sampai (5).

$$Accuracy = \frac{TN + TP}{TN + FN + TP + FP} \quad (1)$$

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (3)$$

$$L = - \sum_i y_i \log \hat{y}_i \quad (3)$$

3. HASIL DAN PEMBAHASAN

3.1 Pengumpulan Data

Dataset yang digunakan dalam penelitian terbagi menjadi tiga jenis data yaitu data *training*, *testing*, dan validasi. Kelas yang digunakan dalam penelitian ini sebanyak tiga kelas jenis hewan seperti ditunjukkan pada Tabel 1 dan contoh dataset pada Gambar 2. Pada Tabel 1, dataset yang digunakan pada data *training* dibagi menjadi 1.000 data/kelas dan total keseluruhan data *training* yang digunakan sebanyak 3.000 data. Data *testing* digunakan untuk memvalidasi model yang sudah dilatih untuk melihat akurasi dalam melakukan klasifikasi. Data testing dibagi menjadi 500 data/kelas dan total keseluruhan data testing sebanyak 1.500 data. Data validasi digunakan untuk mengevaluasi dan menyetel parameter model selama proses pelatihan untuk mencegah overfitting. Data validasi dibagi menjadi 100 data/kelas dan total keseluruhan data validasi sebanyak 300 data.

Tabel 1 Dataset

Jenis Data	Kelas	Jumlah Data	Total Data
Data Training	Anjing	1000	3000
	Harimau	1000	
	Kucing	1000	
Data Testing	Anjing	500	1500
	Harimau	500	
	Kucing	500	
Data Validasi	Anjing	100	300
	Harimau	100	
	Kucing	100	

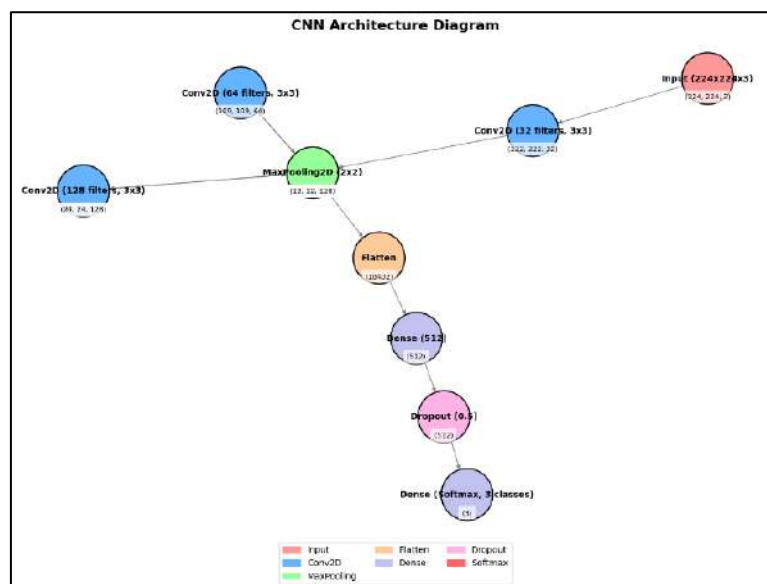




Gambar 2 Contoh Dataset

3.2 Pemodelan CNN

Model CNN ini dirancang untuk klasifikasi citra menjadi tiga kelas (anjing, kucing, dan harimau) dengan input berukuran 224×224 piksel (RGB). Arsitektur terdiri dari empat lapisan konvolusi dengan jumlah filter bertambah secara bertahap (32, 64, 128, 128), menggunakan kernel 3×3 dan aktivasi ReLU untuk mengekstraksi fitur. Setiap lapisan konvolusi diikuti oleh MaxPooling2D (2×2) untuk mengurangi dimensi data dan kompleksitas komputasi. Setelah ekstraksi fitur, lapisan Flatten mengubah data menjadi vektor satu dimensi, yang kemudian diproses oleh Dense (512 neuron, ReLU) dan Dropout (0,5) untuk mengurangi overfitting. Lapisan output menggunakan Dense (3 neuron, Softmax) untuk menghasilkan probabilitas klasifikasi. Model dikompilasi dengan optimizer Adam dan categorical crossentropy, yang optimal untuk klasifikasi multi-kelas. Kombinasi konvolusi, pooling, dan fully connected memastikan model dapat mengenali pola dalam citra secara efektif. Pemodelan ditunjukkan pada Gambar 3.



Gambar 3 Pemodelan CNN

3.3 Pelatihan Model

Model mencapai akurasi pelatihan 86% setelah 50 epoch, dengan penurunan bertahap dalam nilai loss pelatihan. Evaluasi pada data validasi juga menunjukkan peningkatan akurasi hingga sekitar 83% pada akhir pelatihan. Hasil ini menunjukkan bahwa model CNN tidak hanya mampu mempelajari pola-pola yang kompleks dari data pelatihan, tetapi juga efektif dalam menggeneralisasi pada data baru yang tidak digunakan dalam pelatihan. Dengan demikian, model ini menunjukkan potensi yang baik untuk aplikasi klasifikasi gambar dalam konteks studi ini. Hasil Pelatihan Model ditunjukkan pada Gambar 4.



```
Epoch 1/50
94/94 ————— 173s 2s/step - accuracy: 0.4035 - loss: 1.0440 - val_accuracy: 0.5233 - val_loss: 0.9392
Epoch 2/50
94/94 ————— 180s 2s/step - accuracy: 0.5175 - loss: 0.9153 - val_accuracy: 0.4833 - val_loss: 0.9110
Epoch 3/50
94/94 ————— 178s 2s/step - accuracy: 0.5349 - loss: 0.8631 - val_accuracy: 0.5700 - val_loss: 0.8279
Epoch 4/50
94/94 ————— 173s 2s/step - accuracy: 0.6087 - loss: 0.7803 - val_accuracy: 0.6467 - val_loss: 0.6953
Epoch 5/50
94/94 ————— 176s 2s/step - accuracy: 0.6294 - loss: 0.6922 - val_accuracy: 0.6533 - val_loss: 0.6729
Epoch 6/50
94/94 ————— 176s 2s/step - accuracy: 0.6615 - loss: 0.6974 - val_accuracy: 0.7033 - val_loss: 0.5779
Epoch 7/50
94/94 ————— 166s 2s/step - accuracy: 0.6759 - loss: 0.6829 - val_accuracy: 0.6900 - val_loss: 0.5923
Epoch 8/50
94/94 ————— 79s 817ms/step - accuracy: 0.6972 - loss: 0.6267 - val_accuracy: 0.7467 - val_loss: 0.5124
Epoch 9/50
94/94 ————— 69s 714ms/step - accuracy: 0.7295 - loss: 0.5776 - val_accuracy: 0.7567 - val_loss: 0.4863
Epoch 10/50
94/94 ————— 79s 821ms/step - accuracy: 0.7152 - loss: 0.5650 - val_accuracy: 0.7433 - val_loss: 0.5134
Epoch 11/50
94/94 ————— 69s 708ms/step - accuracy: 0.7223 - loss: 0.5510 - val_accuracy: 0.7800 - val_loss: 0.4767
Epoch 12/50
94/94 ————— 79s 821ms/step - accuracy: 0.7312 - loss: 0.5565 - val_accuracy: 0.7567 - val_loss: 0.4922
Epoch 13/50
...
Epoch 49/50
94/94 ————— 77s 797ms/step - accuracy: 0.8666 - loss: 0.3137 - val_accuracy: 0.8267 - val_loss: 0.3393
Epoch 50/50
94/94 ————— 67s 689ms/step - accuracy: 0.8606 - loss: 0.3213 - val_accuracy: 0.8300 - val_loss: 0.3396
```

Gambar 4 Hasil Pelatihan Model

3.4 Evaluasi Model

Akurasi model pada data pelatihan adalah 87.06%, yang menunjukkan bahwa model mampu mengenali dan mengklasifikasikan gambar-gambar dalam data pelatihan dengan tingkat keberhasilan sebesar 87.06%. Nilai loss pada data pelatihan adalah 0,3339, yang menggambarkan besarnya kesalahan prediksi model terhadap data pelatihan. Hasil evaluasi model ditunjukkan pada Gambar 5.

```
47/47 ————— 7s 150ms/step - accuracy: 0.8706 - loss: 0.3339
Test loss: 0.3046361804008484
Test accuracy: 0.875333309173584
```

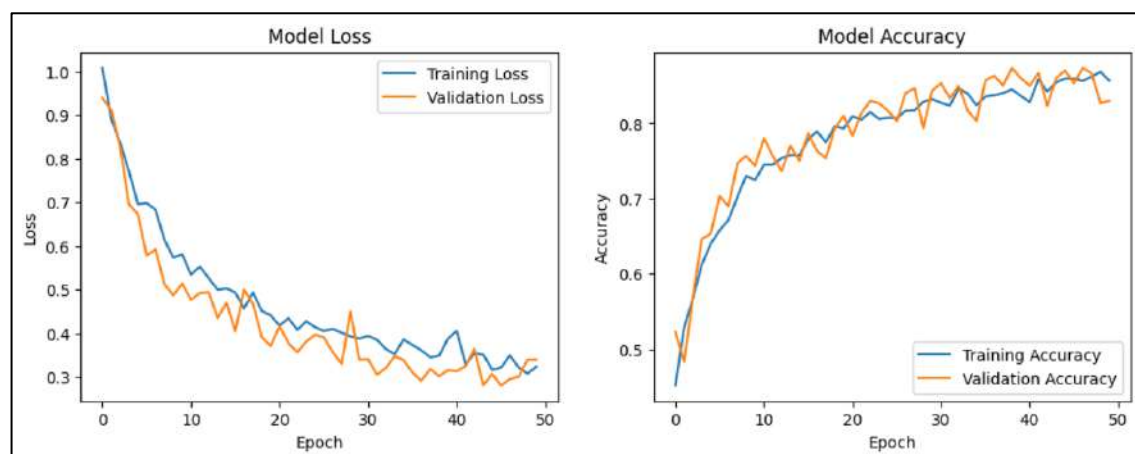
Gambar 5 Hasil Evaluasi Model

3.5 Analisis Hasil

3.5.1 Loss dan Akurasi

Kurva loss menunjukkan penurunan cepat selama 10 epoch pertama untuk data pelatihan dan validasi. Setelah itu, penurunannya lebih bertahap. Pada akhir pelatihan, loss validasi sedikit lebih rendah dibandingkan loss pelatihan, berkisar antara 0,3-0,4. Akurasi model meningkat tajam selama 10 epoch pertama, dari sekitar 45% menjadi 75%. Kemudian, peningkatannya lebih lambat hingga mencapai 90-92% di akhir pelatihan. Kurva akurasi pelatihan dan validasi berdekatan, menunjukkan kemampuan generalisasi model yang baik. Hasil ini menunjukkan bahwa model mampu mempelajari pola-pola dalam data pelatihan secara efektif dan menerapkannya pada data validasi yang belum pernah dilihat sebelumnya. Tidak ada tanda-tanda overfitting yang signifikan, karena performa validasi terus meningkat seiring dengan performa pelatihan. Hasil loss dan akurasi ditunjukkan pada Gambar 6.





Gambar 6 Hasil Evaluasi Model

3.5.2 Metrik Presisi, Recall, F1-Score, dan Akurasi Keseluruhan

Hasil klasifikasi model CNN terhadap data uji menunjukkan metrik evaluasi yang positif untuk ketiga kelas hewan: anjing, harimau, dan kucing. Akurasi keseluruhan model adalah 88%, yang berarti model dapat mengklasifikasikan 88% dari total gambar dengan benar. Untuk anjing, presisi adalah 0,81, recall 0,83, dan F1-score 0,82. Harimau memiliki presisi 0,99, recall 1.00, dan F1-score 0,99, menunjukkan kinerja hampir sempurna. Kucing menunjukkan presisi 0,83, recall 0,80, dan F1-score 0,81. Macro average dari presisi, recall, dan F1-score masing-masing adalah 0,88, menunjukkan kinerja rata-rata yang baik di semua kelas. Selain itu, weighted average dari presisi, recall, dan F1-score juga masing-masing adalah 0,88, mempertimbangkan jumlah sampel di setiap kelas. Hasil ini menunjukkan bahwa model CNN memiliki kemampuan generalisasi yang baik dan dapat mengklasifikasikan gambar-gambar baru dengan akurasi yang tinggi, terutama pada kelas harimau yang menunjukkan kinerja terbaik. Hasil metrik ditunjukkan pada Gambar 6.

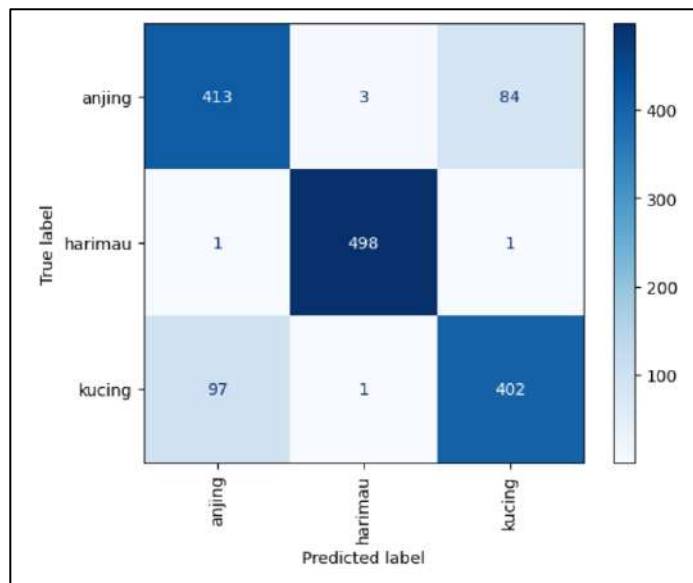
Tabel 2 Hasil Metrik Presisi, Recall, F1-Score, dan Akurasi Keseluruhan

	Presisi	Recall	F1-Score	Support
Anjing	0,81	0,83	0,82	500
Harimau	0,99	1.00	0,99	500
Kucing	0,83	0,80	0,81	500
Akurasi	-	-	0,88	1500
Macro Avg	0,88	0,88	0,88	1500
Weighted Avg	0,88	0,88	0,88	1500

3.5.3 Confusion Matrix

Confusion matrix adalah sebuah metode pengukuran yang digunakan untuk mengevaluasi kinerja atau tingkat keakuratan dari proses klasifikasi (Openg et al., 2022). Matriks ini menunjukkan kinerja model klasifikasi untuk tiga jenis hewan: anjing, harimau, dan kucing. Model memiliki akurasi tertinggi dalam mengidentifikasi harimau, dengan 498 dari 500 kasus dikenali dengan benar. Untuk anjing, model berhasil mengidentifikasi dengan tepat sebanyak 413 kasus, namun terjadi kesalahan klasifikasi di mana 84 anjing salah dikenali sebagai kucing dan 3 sebagai harimau. Kucing memiliki tingkat kesalahan tertinggi, dengan 402 kasus dikenali dengan benar, tetapi 97 salah diklasifikasikan sebagai anjing dan 1 sebagai harimau. Secara keseluruhan, model menunjukkan performa yang baik, terutama dalam mengidentifikasi harimau. Hasil confusion matrix ditunjukkan pada Gambar 7.





Gambar 7 Hasil Confusion Matrix

3.5.4 Akurasi Prediksi Gambar

Model Convolutional Neural Network (CNN) yang telah dilatih digunakan untuk memprediksi kategori dari gambar-gambar individual yang diujikan. Pada Gambar 8, ditampilkan hasil prediksi untuk tiga gambar uji, masing-masing dari kelas anjing, harimau, dan kucing. Untuk gambar anjing, model memberikan probabilitas sebesar 0,6856 (68,56%) untuk kategori “anjing”, dengan nilai probabilitas kecil untuk kategori lainnya. Hal ini menunjukkan bahwa model mampu mengenali ciri khas dari gambar tersebut sebagai gambar anjing, meskipun probabilitas tidak terlalu tinggi dibanding kelas lainnya. Sementara itu, gambar harimau teridentifikasi dengan probabilitas sangat tinggi, yaitu 0,9713 (97,13%), menunjukkan bahwa model sangat yakin terhadap klasifikasi tersebut. Prediksi ini mengindikasikan bahwa ciri visual harimau lebih mudah dibedakan dibandingkan dengan kelas lainnya. Untuk gambar kucing, model memprediksi gambar ini sebagai “kucing” dengan probabilitas 0,8275 (82,75%). Meski demikian, terdapat kemungkinan yang lebih kecil untuk kategori “anjing” dan “harimau”. Prediksi ini menunjukkan kemampuan model dalam mengenali kucing meskipun dengan beberapa tantangan dalam membedakan dari kelas lain.



Gambar 8 Hasil Prediksi Gambar



Prediksi ini mencerminkan bahwa model CNN telah mampu mengklasifikasikan gambar secara akurat pada sebagian besar kasus, meskipun terdapat variasi probabilitas pada beberapa gambar. Hasil ini juga menunjukkan bahwa model memiliki tingkat kepercayaan tinggi pada kategori dengan pola visual yang jelas, seperti harimau. Rata-rata keseluruhan probabilitas untuk ketiga gambar tersebut adalah 0,8281 (82,81%). Nilai ini memberikan gambaran umum tentang tingkat kepercayaan model dalam mengklasifikasikan gambar yang diuji, menunjukkan kinerja yang konsisten pada dataset yang diberikan. Selain itu, rata-rata ini menunjukkan bahwa model berhasil mengidentifikasi pola visual dalam data uji dengan akurasi yang cukup baik, meskipun masih terdapat ruang untuk perbaikan dalam membedakan antara kategori yang memiliki kemiripan visual.

4. KESIMPULAN

Berdasarkan evaluasi yang dilakukan, model Convolutional Neural Network (CNN) berhasil mencapai akurasi keseluruhan sebesar 88% dalam mengklasifikasikan gambar hewan seperti anjing, harimau, dan kucing. Performa model ini terutama mencolok dalam mengenali gambar harimau, dengan tingkat akurasi yang mencapai 99%. Meskipun demikian, model menghadapi kesulitan dalam membedakan antara gambar anjing dan kucing, di mana akurasi untuk kedua jenis hewan tersebut 81%. Selama proses pelatihan, terjadi peningkatan signifikan dalam akurasi dari awal hingga akhir, dengan mencapai 92% pada epoch terakhir dari nilai awal sekitar 45%. Evaluasi menggunakan confusion matrix menunjukkan bahwa sebagian besar kasus berhasil diklasifikasikan dengan tepat, meskipun terdapat kesalahan yang menunjukkan adanya kemiripan fitur antara anjing dan kucing dalam perspektif model. Dalam hal presisi, recall, F1-score, dan akurasi keseluruhan model mampu mencapai nilai yang memuaskan untuk setiap kelas, menunjukkan kemampuan yang baik dalam mengklasifikasikan dengan presisi tinggi dan mampu mengingat kembali data yang relevan. Secara keseluruhan, hasil ini menegaskan bahwa CNN dapat digunakan untuk klasifikasi gambar hewan, dan menunjukkan potensi untuk pengembangan lebih lanjut di masa depan untuk mengatasi tantangan dalam membedakan detail antara jenis-jenis hewan yang serupa.

Sebagai saran bagi pengembang selanjutnya, penelitian ini dapat dikembangkan lebih lanjut dengan memperluas cakupan data, baik dari segi jumlah maupun variasi dataset, untuk meningkatkan generalisasi model. Selain itu, dapat dilakukan eksplorasi terhadap arsitektur jaringan yang lebih kompleks atau metode augmentasi data yang lebih bervariasi guna meningkatkan akurasi. Uji coba pada lingkungan nyata juga disarankan untuk memastikan performa model dalam kondisi sebenarnya serta mengidentifikasi potensi perbaikan di masa depan.

DAFTAR PUSTAKA

- Adrianto, H., Setyawan, Y., Banjarnahor, D. P., Kusumah, I. P., & Messakh, B. D. (2022). Pembekalan Klasifikasi Baru Makhluk Hidup Hewan Kepada Guru-Guru Biologi. *Sebatik*, 26(2), 638–643. <https://doi.org/10.46984/sebatik.v26i2.2152>
- Ahmad, A., Idris, I. S. K., & Bode, A. (2023). Klasifikasi Jenis Buah Tomat Menggunakan Convolutional Neural Network. *Jurnal Ilmiah Ilmu Komputer*, 2(2), 83–89. <https://doi.org/10.37195/balok.v2i2.617>
- A'yun, D. K., & Erman, E. (2019). Kemampuan Siswa Mengklarifikasi Kingdom Animalia Invertebrata: Studi Kasus di SMP Negeri 1 Jabon. *PENSA: E-Jurnal Pendidikan Sains*, 7(3), 361–366. <https://ejournal.unesa.ac.id/index.php/pensa/article/view/32294>
- Fukushima, K. (1980). Neocognitron: A Self-Organizing Neural Network Model for a Mechanism of Pattern Recognition Unaffected by Shift in Position. *Biological Cybernetics*, 36(4), 193–202. <https://doi.org/10.1007/BF00344251>
- He, K., Gkioxari, G., Dollar, P., & Girshick, R. (2020). Mask R-CNN. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 42(2), 386–397. <https://doi.org/10.1109/TPAMI.2018.2844175>



- Hidayatullah, D., Ardiansah, T., & Styawati, S. (2022). Sistem Informasi Reservasi Pelayanan dan Penyewaan Fasilitas Lapangan Futsal Berbasis Web dengan Metode Waterfall. *Jurnal Teknologi dan Sistem Informasi (JTSI)*, 3(3), 64–68. <https://doi.org/10.33365/jtsi.v3i3.1994>
- K, L., Gadde, S., Puttagunta, M. K., Dhanalakshmi, G., & El-Ebiary, Y. A. B. (2023). Efficiency Analysis of Firefly Optimization-Enhanced GAN-Driven Convolutional Model for Cost-Effective Melanoma Classification. *International Journal of Advanced Computer Science and Applications*, 14(11), 742–753. <https://doi.org/10.14569/IJACSA.2023.0141175>
- Laksono, S. A., Rahmat, B., & Nugroho, B. (2024). Identifikasi Jenis Ikan Cupang Berdasarkan Gambar Menggunakan Metode Convolutional Neural Network. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 3331–3338. <https://doi.org/10.36040/jati.v8i3.9676>
- Leovinent, A., & Yoannita, Y. (2023). Klasifikasi Ras Anjing Berdasarkan Citra Menggunakan Convolutional Neural Network. *Jurnal Algoritme*, 3(2), 160–169. <https://doi.org/10.35957/algoritme.v3i2.3389>
- Mariyanti, S., Gayatri, Y., & Wikanta, W. (2022). Pengembangan Atlas Klasifikasi Hewan Vertebrata Berbasis Sumber Daya Hayati Lokal Sebagai Sumber Belajar Biologi di Sekolah Penulis. *Journal of Science, Education, and Studies*, 1(1), 1–9. <https://journal.um-surabaya.ac.id/index.php/J-SES/article/view/14877>
- Micheal, M., & Hartati, E. (2022). Klasifikasi Spesies Kupu Kupu Menggunakan Metode Convolutional Neural Network. *MDP Student Conference (MSC) 2022*, 569–577. <https://jurnal.mdp.ac.id/index.php/msc/article/view/1928>
- Mohite, O. S., Jørgensen, T. S., Booth, T. J., Charusanti, P., Phaneuf, P. V., Weber, T., & Palsson, B. O. (2025). Pangenome Mining of the *Streptomyces* Genus Redefines Species' Biosynthetic Potential. *Genome Biology*, 26(1), Article ID: 9. <https://doi.org/10.1186/s13059-024-03471-9>
- Norouzzadeh, M. S., Nguyen, A., Kosmala, M., Swanson, A., Palmer, M. S., Packer, C., & Clune, J. (2018). Automatically Identifying, Counting, and Describing Wild Animals in Camera-Trap Images with Deep Learning. *Proceedings of the National Academy of Sciences*, 115(25), E5716–E5725. <https://doi.org/10.1073/pnas.1719367115>
- Openg, J. B. J. R., Hiswati, M. E., & Hamzah, H. (2022). Klasifikasi Unggas Ordo Anseriformes Berdasarkan Citra Menggunakan Metode Deep Learning dengan Algoritma Convolutional Neural Network (CNN). *7th Seminar Nasional Teknik Elektro, Informatika dan Sistem Informasi (SINTaKS)*, 1. <https://doi.org/10.35842/sintaks.v1i1.3>
- Ricard, A., Crevier-Denoix, N., Pourcelot, P., Crichan, H., Sabbagh, M., Dumont-Saint-Priest, B., & Danvy, S. (2023). Genetic Analysis of Geometric Morphometric 3D Visuals of French Jumping Horses. *Genetics Selection Evolution*, 55(1), Article ID: 63. <https://doi.org/10.1186/s12711-023-00837-8>
- Samudra, J. T., Rosnelly, R., Situmorang, Z., & Ramadhan, P. S. (2023). Model Klasifikasi Jenis Hewan dengan SVM, KNN, Logistic Regression Menggunakan Pre-Trained VGG 16. *Jurnal SAINTIKOM (Jurnal Sains Manajemen Informatika dan Komputer)*, 22(2), 225–231. <https://doi.org/10.53513/jis.v22i2.8314>
- Suhendar, S., Purnama, A., & Fauzi, E. (2023). Deteksi Penyakit pada Daun Tanaman Ubi Jalar Menggunakan Metode Convolutional Neural Network. *Jurnal Ilmiah Informatika Global*, 14(3), 62–67. <https://doi.org/10.36982/jiig.v14i3.3478>
- Supiyandi, S., Sitorus, A., Fitriah, N., Virul, H., & Rangkuti, S. P. (2024). Pendeteksi Gerakan pada Vidio Menggunakan Python dan OpenCV. *Merkurius: Jurnal Riset Sistem Informasi dan Teknik Informatika*, 2(6), 334–343. <https://doi.org/10.61132/mercurius.v2i6.522>
- Tan, M., & Le, Q. (2019). EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks. In K. Chaudhuri & R. Salakhutdinov (Eds.), *Proceedings of the 36th International Conference on Machine Learning* (Vol. 97, pp. 6105–6114). PMLR. <https://proceedings.mlr.press/v97/tan19a.html>
- Younis, E. M. G., Zaki, S. M., Kanjo, E., & Houssein, E. H. (2022). Evaluating Ensemble Learning Methods for Multi-Modal Emotion Recognition Using Sensor Data Fusion. *Sensors*, 22(15), Article ID: 5611. <https://doi.org/10.3390/s22155611>



Arsitektur Microservice untuk Optimalisasi Aplikasi Eco-Maps dalam Mendukung Kampus Ramah Lingkungan

Alam Rahmatulloh ^{(1)*}, Rohmat Gunawan ⁽²⁾, Randi Rizal ⁽³⁾

Departemen Informatika, Universitas Siliwangi, Tasikmalaya, Indonesia

e-mail : {alam,rohmatgunawan,randirizal}@unsil.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 30 Agustus 2024, direvisi 19 Maret 2025, diterima 22 Maret 2025, dan dipublikasikan 30 September 2025.

Abstract

The implementation of environmentally friendly campus concepts has become increasingly crucial in addressing global environmental challenges. Eco-Maps is an application designed to visualize and manage sustainability efforts on campus, including energy management, waste management, and sustainable transportation initiatives. To enhance efficiency and flexibility, this study discusses the application of a microservice architecture in Eco-Maps. This architecture supports faster and more efficient development, testing, and deployment, while enabling horizontal scalability to manage high complexity and large data volumes. By separating application functions into independent services, microservices facilitate maintenance and updates while minimizing the impact of failures in individual services. This study also reviews the integration of containerization technologies, such as Docker and Kubernetes, to support microservice implementation. Through these technologies, the application can be deployed quickly and consistently across various environments, from development to production. System testing was conducted using load testing and stress testing methods, as shown in Tables 3 and 4. The results demonstrate that the average response time across ten iterations was 745.9 ms, with an average CPU usage of 44.38%. These findings confirm that processing load directly affects CPU efficiency and overall system performance.

Keywords: *Eco-Maps, Microservice Architecture, Sustainable Campus, Docker, Kubernetes, System Performance*

Abstrak

Penerapan konsep kampus ramah lingkungan semakin krusial dalam menghadapi tantangan lingkungan global. Eco-Maps merupakan aplikasi yang dikembangkan untuk memvisualisasikan dan mengelola upaya keberlanjutan di kampus, seperti manajemen energi, pengelolaan limbah, dan transportasi berkelanjutan. Untuk meningkatkan efisiensi dan fleksibilitas, penelitian ini membahas penerapan arsitektur *microservice* pada Eco-Maps. Arsitektur ini mendukung pengembangan, pengujian, dan penyebaran yang lebih cepat serta efisien, sekaligus memungkinkan skalabilitas horizontal dalam mengelola kompleksitas dan volume data yang besar. Dengan memisahkan fungsi aplikasi ke dalam layanan independen, *microservice* memudahkan pemeliharaan dan pembaruan, serta meminimalkan dampak apabila terjadi gangguan pada salah satu layanan. Penelitian ini juga meninjau integrasi teknologi kontainerisasi seperti Docker dan manajemen orkestrasi Kubernetes untuk mendukung implementasi *microservice*. Melalui teknologi tersebut, aplikasi dapat dideploy secara cepat dan konsisten di berbagai lingkungan, mulai dari pengembangan hingga produksi. Pengujian sistem dilakukan menggunakan metode load test dan stress test, sebagaimana ditunjukkan pada Tabel 3 dan Tabel 4. Hasil menunjukkan bahwa rata-rata waktu respons dari sepuluh iterasi adalah 745,9 ms dengan rata-rata penggunaan CPU sebesar 44,38%. Temuan ini menegaskan bahwa beban pemrosesan memiliki pengaruh langsung terhadap efisiensi CPU dan kinerja keseluruhan sistem.

Kata Kunci: *Eco-Maps, Arsitektur Microservice, Kampus Ramah Lingkungan, Docker, Kubernetes, Kinerja Sistem*



1. PENDAHULUAN

Dalam era digital yang terus berkembang pesat, teknologi informasi memainkan peran yang semakin penting dalam mendukung berbagai aspek kehidupan, termasuk dalam upaya mewujudkan kampus ramah lingkungan (Abdullai et al., 2022; Abdulmouti et al., 2022; Faizah & Nugraheni, 2024; Trevisan et al., 2023). Kampus-kampus di seluruh dunia kini berupaya untuk menerapkan konsep keberlanjutan melalui berbagai inisiatif yang bertujuan untuk mengurangi dampak lingkungan, meningkatkan efisiensi energi, dan mengelola sumber daya secara lebih bijak (Khan & Ximei, 2022; Sugiarto et al., 2022). Salah satu alat yang dapat digunakan untuk mendukung inisiatif-inisiatif ini adalah aplikasi berbasis Eco-Maps, sebuah platform digital yang dirancang untuk membantu universitas dalam memantau dan mengelola infrastruktur hijau, penggunaan energi, pengelolaan limbah, serta berbagai kegiatan lain yang berkontribusi terhadap keberlanjutan kampus (Aasa et al., 2020; Amelia et al., 2023; Martins et al., 2021; Ribeiro et al., 2021).

Namun, seiring dengan meningkatnya kompleksitas data dan kebutuhan untuk memberikan layanan yang lebih cepat dan responsif, arsitektur perangkat lunak yang digunakan untuk mengembangkan aplikasi Eco-Maps sering kali menghadapi tantangan besar (Dimov et al., 2022; Palomo et al., 2018; Xiao, 2024). Arsitektur monolitik tradisional, yang menggabungkan semua fungsi aplikasi ke dalam satu kesatuan, sering kali tidak mampu menangani tuntutan ini dengan baik (Iqbal et al., 2023; Sidiq et al., 2024). Sebagai salah satu solusi, banyaknya pengembang perangkat lunak mulai beralih ke arsitektur *microservice*, sebuah pendekatan desain yang membagi aplikasi menjadi layanan-layanan kecil yang dapat dioperasikan secara independen namun tetap dapat saling berintegrasi dengan baik (Santos et al., 2024; Weerasinghe & Perera, 2024).

Arsitektur *microservice* menawarkan sejumlah keunggulan yang sangat relevan dengan kebutuhan pengembangan aplikasi Eco-Maps (Kautsar et al., 2023). Dengan membagi aplikasi menjadi layanan-layanan yang lebih kecil dan lebih terfokus, setiap layanan dapat dikembangkan, diuji, dan dideploy secara terpisah tanpa harus mempengaruhi layanan lain yang sudah berjalan (Telang, 2023). Hal ini tidak hanya memungkinkan pengembangan yang lebih cepat dan efisien, tetapi juga meningkatkan fleksibilitas dan skalabilitas sistem. Dalam konteks kampus ramah lingkungan, arsitektur *microservice* memungkinkan aplikasi Eco-Maps untuk terus berkembang dan beradaptasi dengan kebutuhan yang dinamis, seperti perubahan dalam manajemen energi, pengelolaan transportasi, atau integrasi dengan sistem IoT yang ada.

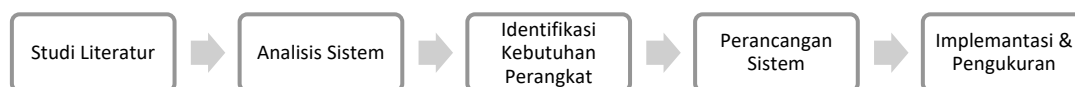
Penerapan arsitektur *microservice* pada aplikasi Eco-Maps juga membuka peluang untuk integrasi yang lebih baik dengan sistem lain yang mungkin sudah ada di kampus. Sebagai contoh, aplikasi ini dapat terhubung dengan sistem manajemen gedung yang cerdas, jaringan sensor IoT, serta platform analitik data yang digunakan untuk memantau dan mengoptimalkan penggunaan sumber daya di kampus. Dengan pendekatan ini, aplikasi Eco-Maps tidak hanya berfungsi sebagai alat yang berdiri sendiri, tetapi juga sebagai bagian dari ekosistem digital yang lebih besar yang mendukung upaya kampus untuk menjadi lebih ramah lingkungan. Oleh karena itu, penelitian ini bertujuan untuk mengeksplorasi penerapan arsitektur *microservice* dalam pengembangan aplikasi Eco-Maps dan bagaimana pendekatan ini dapat mengoptimalkan fungsionalitas serta kinerja aplikasi dalam mendukung visi kampus ramah lingkungan.

2. METODE PENELITIAN

Metode penelitian ini terdiri dari lima tahap utama, yaitu: studi literatur, analisis sistem, identifikasi kebutuhan perangkat, perancangan sistem, serta implementasi dan pengukuran, sebagaimana ditunjukkan pada Gambar 1. Tahap studi literatur dilakukan dengan menelaah berbagai jurnal dan referensi relevan untuk memperoleh dasar teori serta pemahaman mendalam terkait topik penelitian. Selanjutnya, tahap analisis sistem bertujuan untuk mengidentifikasi permasalahan yang ada dan merumuskan solusi yang sesuai. Proses kemudian dilanjutkan dengan identifikasi kebutuhan perangkat, mencakup perangkat keras maupun perangkat lunak yang diperlukan dalam penelitian. Tahap perancangan sistem meliputi penyusunan alur kerja serta struktur



sistem. Terakhir, tahap implementasi dan pengukuran dilakukan untuk menguji sistem yang telah dirancang serta mengevaluasi kinerjanya secara menyeluruh.



Gambar 1 Metodologi Penelitian

2.1 Studi Literatur

Kegiatan pada tahap ini mencakup peninjauan mendalam terhadap penelitian yang ada serta pemahaman menyeluruh mengenai seluruh informasi terkait *microservice*. Selain itu, sumber referensi yang berkaitan dengan teknologi kontainer dan arsitektur berbasis API juga ditelaah secara komprehensif. Penelaahan ini tidak hanya terbatas pada literatur umum tetapi juga mencakup jurnal-jurnal ilmiah yang relevan dan telah diterbitkan, guna memastikan pemahaman yang mendalam dan akurat terhadap konsep-konsep yang mendasari. Dengan pendekatan ini, diharapkan mampu menghasilkan landasan teoritis yang kuat untuk pengembangan lebih lanjut.

2.2 Analisis Sistem

Untuk mengevaluasi kinerja *microservice* dalam penelitian ini sekaligus mendukung optimalisasi aplikasi Eco-Maps dalam mewujudkan kampus ramah lingkungan, dikembangkan sebuah prototipe aplikasi yang mengintegrasikan berbagai layanan pengelolaan lingkungan kampus. Sistem informasi ini, sebagaimana digambarkan pada Gambar 2, terdiri atas empat layanan utama, yaitu Auth Service (layanan 1), Eco-Map Service (layanan 2), Waste Management Service (layanan 3), dan Energy Monitoring Service (layanan 4), serta NATS Streaming Service sebagai event-bus. Seluruh layanan terhubung dengan basis data yang digunakan, yaitu MongoDB.

Auth Service bertugas menangani seluruh aktivitas autentikasi dan manajemen pengguna. Eco-Map Service menyediakan peta interaktif kampus yang menampilkan informasi mengenai area hijau, titik daur ulang, serta lokasi-lokasi penting lainnya. Waste Management Service mengelola proses pengumpulan, pemilahan, dan pembuangan limbah secara efisien. Energy Monitoring Service memantau konsumsi energi di seluruh kampus sekaligus memberikan rekomendasi penghematan energi. Terakhir, NATS Streaming Service berfungsi sebagai penghubung utama untuk pertukaran data antar-layanan dengan mekanisme *publishing* (mengirim data) dan *listening* (menerima data).

2.3 Identifikasi Kebutuhan Perangkat

Proses pengukuran dalam penelitian ini memerlukan identifikasi perangkat keras dan perangkat lunak untuk memastikan sistem dapat berjalan secara optimal. Spesifikasi perangkat keras yang digunakan disajikan secara rinci pada Tabel 1, memuat informasi mengenai komponen utama, seperti prosesor, memori, dan kapasitas penyimpanan. Sedangkan spesifikasi perangkat lunak yang mendukung penelitian ditampilkan pada Tabel 2, mencakup sistem operasi, bahasa pemrograman, serta perangkat pendukung lainnya yang digunakan selama proses implementasi.

Tabel 1 Spesifikasi Perangkat Keras

No.	Nama	Deskripsi
1	CPU	AMD A4-9120 RADEON R3, 4 COMPUTE CORES 2C + 2G @ 2x 2.25GHz
2	GPU	AMD STONEY
3	RAM	8GB
4	Kernel	x86_64 Linux 5.10.70-1-MANJARO
5	Disk	110GB
6	Operating System	Manjaro 21.1.6 Pahvo (Linux)

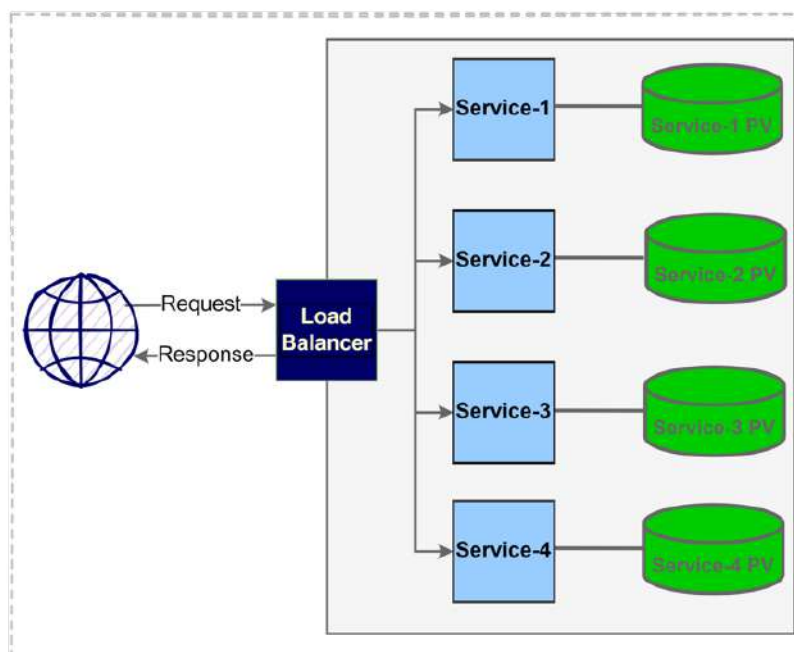


Tabel 2 Spesifikasi Perangkat Lunak

No.	Nama	Deskripsi
1	Docker Latest Version	Application-oriented container
2	Kubernetes v1.22.2	Container orchestrator
3	Scaffold v2 beta20	Kubernetes support
4	NATS Streaming v0.22	Event bus
5	Typescript v4.4.4	Programming language
6	Node.js v17.1.0	Runtime environment
7	Express.js v4.17.1	App framework
8	MongoDB v5.0	Database
9	Jest v27.1.0	Testing framework
10	Stan.js v0.3.2	Client communication tool

2.4 Perancangan Sistem

Pada tahap ini, prototipe aplikasi berbasis layanan mikro dikembangkan dengan menerapkan prinsip komunikasi data dalam arsitektur *microservice*. Prototipe aplikasi terdiri atas empat layanan utama, yaitu Service-1, Service-2, Service-3, dan Service-4, yang masing-masing terhubung ke basis data tersendiri sebagaimana ditunjukkan pada Gambar 2. Gambar 2 memperlihatkan arsitektur *microservice* yang terdiri dari beberapa layanan (Service-1 hingga Service-4) yang dikelola oleh sebuah Load Balancer. Load Balancer berfungsi mendistribusikan permintaan (request) dari pengguna secara merata ke setiap layanan untuk mencegah terjadinya kelebihan beban pada salah satu layanan. Masing-masing layanan memiliki penyimpanan data tersendiri (Persistent Volume) yang terhubung langsung, sehingga memungkinkan penyimpanan data secara persisten. Alur permintaan dan respons ditunjukkan oleh panah: permintaan dari pengguna masuk melalui Load Balancer, kemudian diproses oleh layanan terkait, dan hasilnya dikirim kembali ke pengguna melalui Load Balancer. Dengan rancangan ini, arsitektur *microservice* diharapkan mampu memberikan efisiensi, skalabilitas, serta ketersediaan layanan yang optimal.



Gambar 2 Communication Process in *Microservice*



2.5 Implementasi dan Pengukuran

Pada tahap ini, dilakukan pengujian dengan metode *load* dan *stess test*. Selain itu, waktu respons (ms), tingkat kesalahan (%) (Hong et al., 2018), dan Penggunaan CPU (ms) diukur pada arsitektur layanan mikro berbasis *microservice*. Pengukuran waktu respons dan tingkat kesalahan dilakukan dengan bantuan Apache JMeter dan API node.js untuk mengukur penggunaan CPU. Percobaan dilakukan dengan mengirimkan 100 paket data permintaan dari klien ke aplikasi pada server berbasis layanan mikro. Setiap percobaan diulang sepuluh kali dan hasil percobaan dicatat dalam tabel.

3. HASIL DAN PEMBAHASAN

Pada tahap ini dibahas implementasi sistem Eco-Maps beserta hasil pengukuran waktu respons dan penggunaan CPU dalam arsitektur berbasis *microservice*. Implementasi ini menunjukkan bagaimana arsitektur *microservice* mendukung integrasi serta pengelolaan data lingkungan secara efisien. Pengukuran waktu respons dilakukan untuk mengevaluasi kecepatan sistem dalam menangani permintaan API. Arsitektur *microservice* pada Eco-Maps mencakup empat layanan utama, yaitu:

- 1) Auth Service, dengan endpoint untuk registrasi (POST /auth/register), login (POST /auth/login), dan verifikasi token (GET /auth/verify).
- 2) Eco-Map Service, menyediakan API untuk memperoleh daftar lokasi ramah lingkungan (GET /eco-map/locations) dan rute menuju lokasi tertentu (GET /eco-map/route).
- 3) Waste Management Service, dengan endpoint untuk penjadwalan pengelolaan limbah (POST /waste-management/schedule) dan pengambilan laporan pengelolaan limbah (GET /waste-management/reports).
- 4) Energy Monitoring Service, menyediakan API untuk memantau konsumsi energi (GET /energy-monitoring/usage) dan melaporkan anomali energi (POST /energy-monitoring/report).

Komunikasi antar-layanan dilakukan melalui NATS Streaming menggunakan event `auth.user.created`, `waste.management.schedule.created`, dan `energy.monitoring.report.submit` untuk mendukung sinkronisasi data secara *event-driven*. Selain itu, pengukuran penggunaan CPU dilakukan untuk menilai efisiensi pemrosesan dalam konteks arsitektur *microservice*. Hasil pengujian menunjukkan bahwa sistem Eco-Maps berbasis *microservice* mampu memberikan kinerja optimal, baik dari segi kecepatan respons maupun efisiensi penggunaan sumber daya CPU, sehingga efektif dalam mendukung operasional sistem secara keseluruhan.

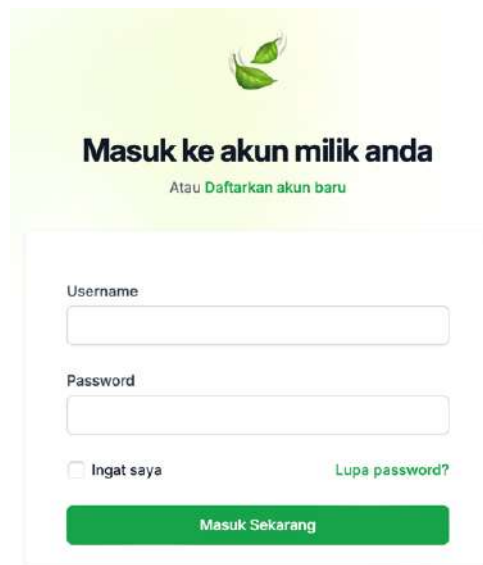
3.1 Implementasi Sistem Eco-Maps

Pada tahap ini, implementasi sistem Eco-Maps dimulai dari proses login sebagaimana ditunjukkan pada Gambar 3, kemudian dilanjutkan dengan tampilan Menu Utama yang terdiri atas beberapa komponen penting. Setelah berhasil masuk, pengguna diarahkan ke menu utama yang mencakup fitur-fitur berikut (Gambar 4):

- 1) Dashboard, menampilkan informasi sistem secara keseluruhan.
- 2) User Management, untuk pengelolaan data pengguna.
- 3) Master Category, untuk pengaturan kategori dalam sistem.
- 4) Master Soal, untuk manajemen pertanyaan yang digunakan dalam proses evaluasi.
- 5) Master Instrument, untuk pengelolaan instrumen atau alat ukur yang digunakan.
- 6) Hasil Eco-Maps, untuk menampilkan hasil analisis berdasarkan data yang diproses oleh sistem.

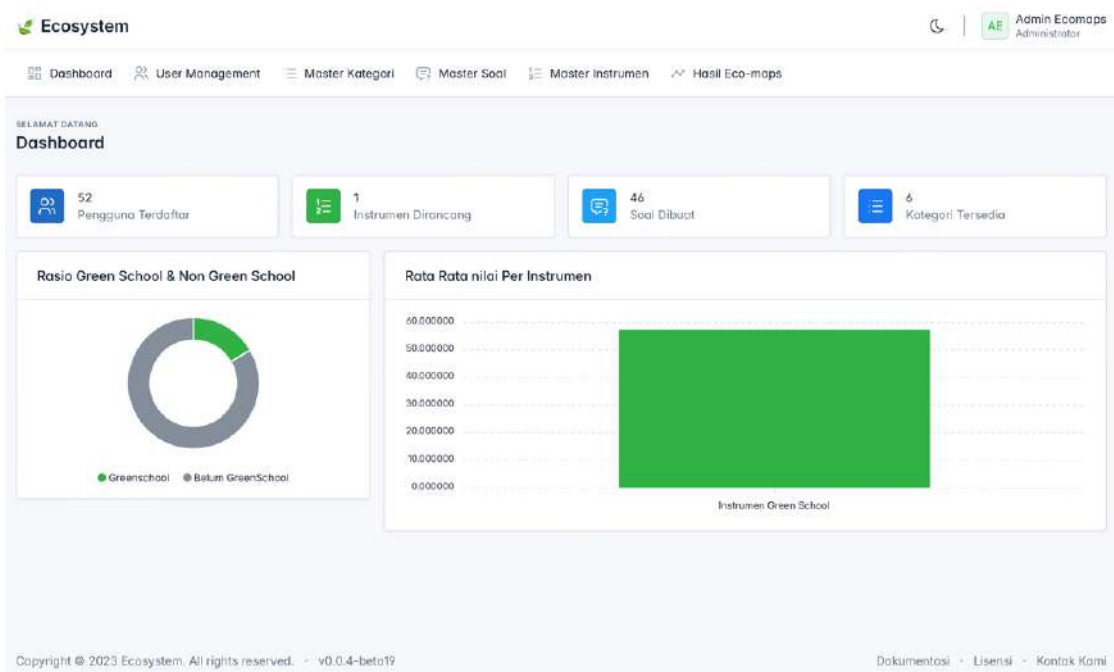
Struktur ini dirancang dengan mempertimbangkan aspek kemudahan navigasi dan efisiensi penggunaan. Hasil pengujian menunjukkan bahwa sistem Eco-Maps dengan rancangan tersebut mampu memberikan performa yang responsif serta mendukung kebutuhan pengguna dalam manajemen data dan analisis secara efektif.





The login form is titled "Masuk ke akun milik anda" (Log in to your account) with a sub-link "Atau Daftarkan akun baru" (Or register a new account). It contains two input fields for "Username" and "Password". Below the password field is a checkbox labeled "Ingat saya" (Remember me) and a link "Lupa password?" (Forgot password?). A green button labeled "Masuk Sekarang" (Log in now) is at the bottom.

Gambar 3 Login Eco-Maps



Gambar 4 Menu Utama Eco-Maps

3.2 Pengujian menggunakan Stress and Load Test

Tabel 3 dan Tabel 4 menyajikan hasil pengujian Load Test dan Stress Test pada implementasi arsitektur *microservice* aplikasi Eco-Maps, yang terdiri dari empat layanan utama, yaitu Auth Service (Layanan 1), Eco-Map Service (Layanan 2), Waste Management Service (Layanan 3), dan Energy Monitoring Service (Layanan 4), serta layanan NATS Streaming sebagai event-bus. Hasil pengujian menunjukkan bahwa setiap layanan memiliki performa yang bervariasi sesuai dengan fungsinya. Auth Service dan NATS Streaming menunjukkan performa terbaik, dengan waktu respons rata-rata masing-masing 120 ms dan 110 ms, serta tingkat kesalahan minimal (*error rate* 10% dan 5%) bahkan pada pengujian beban tinggi.



Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

Sementara itu, Eco-Map Service dan Waste Management Service mengalami *bottleneck* signifikan, ditandai dengan waktu respons tinggi (150 ms dan 180 ms) serta peningkatan *error rate* hingga 15% dan 20% saat stress test. Hal ini disebabkan oleh tingginya konsumsi memori dan latensi database yang mencapai 120 ms. Untuk penelitian selanjutnya, disarankan penerapan caching pada Eco-Map dan Waste Management, serta auto-scaling untuk menangani lonjakan beban. Selain itu, Energy Monitoring Service membutuhkan peningkatan kapasitas memori untuk menjaga performa saat beban tinggi. Secara keseluruhan, meskipun arsitektur *microservice* memberikan fleksibilitas dan skalabilitas, optimasi lebih lanjut pada database, alokasi sumber daya, dan sistem caching diperlukan untuk meningkatkan performa di seluruh layanan.

Tabel 3 Hasil Pengujian dengan Load Test

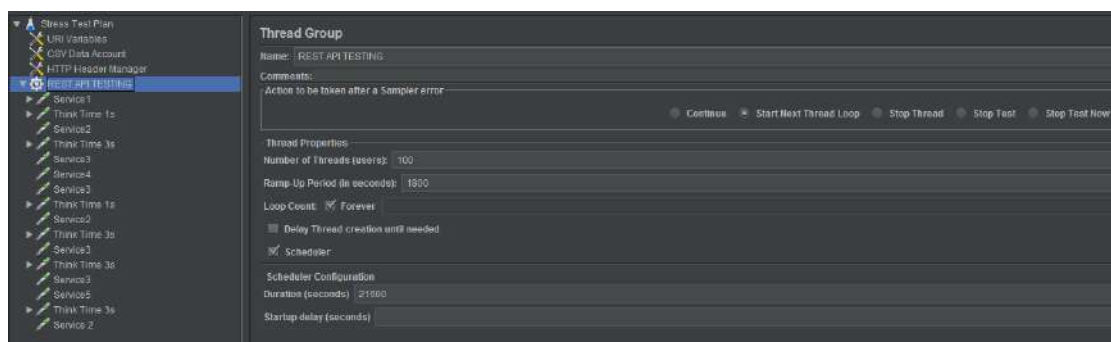
Parameter Uji	Layanan 1	Layanan 2	Layanan 3	Layanan 4	Event Bus
Jumlah Request	300/sec	700/sec	600/sec	500/sec	1000/sec
Response Time	100 ms	130 ms	160 ms	150 ms	90 ms
Error rate	0.2%	0.5%	0.8%	0.3%	0.1%
Throughput (Reg/Sec)	300	690	580	490	1,000
Memory Usage	200 MB	300 MB	350 MB	250 MB	150 MB
Database Latency	50 ms	80 ms	120 ms	70 ms	-
Service Latency	20 ms	40 ms	60 ms	30 ms	10 ms
Network Bandwidth	100 Mbps	150 Mbps	120 Mbps	110 Mbps	200 Mbps

Tabel 4 Hasil pengujian dengan Stress Test

Parameter Uji	Layanan 1	Layanan 2	Layanan 3	Layanan 4	Event Bus
Jumlah Request	500/sec	1000/sec	800/sec	700/sec	1200/sec
Response Time	120 ms	150 ms	180 ms	170 ms	110 ms
Error rate	10%	15%	20%	12%	5%
Throughput (Reg/Sec)	490	970	760	680	1190
Memory Usage	300 MB	450 MB	500 MB	400 MB	200 MB
Database Latency	50 ms	80 ms	120 ms	70 ms	-
Service Latency	20 ms	40 ms	60 ms	30 ms	10 ms
Network Bandwidth	100 Mbps	150 Mbps	120 Mbps	110 Mbps	200 Mbps

3.3 Hasil Pengukuran Waktu Respon dan Penggunaan CPU

Percobaan dilakukan dengan mengirimkan 100 paket data permintaan dari klien ke server, dengan menerapkan *microservice*. Percobaan diulang sebanyak sepuluh kali, dan dihitung nilai rata-ratanya. Waktu respons diukur menggunakan alat Apache JMeter, seperti yang ditunjukkan pada Gambar 5. Setelah sepuluh kali percobaan, data yang diperoleh dicatat, kemudian dihitung persentase perubahan data antara arsitektur komunikasi data berbasis *microservice*, seperti yang ditunjukkan pada Tabel 3.



Gambar 5 Konfigurasi jumlah permintaan pada JMeter



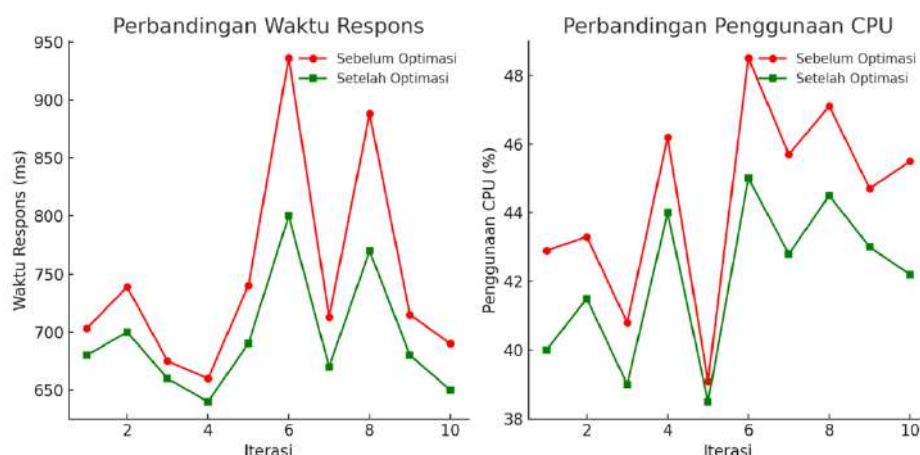
Tabel 5 Hasil Pengukuran Waktu Respon dan Penggunaan CPU

Percobaan	Waktu (ms)	CPU Usage (%)
1	703	42,9
2	739	43,3
3	675	40,8
4	660	46,2
5	740	39,1
6	936	48,5
7	713	45,7
8	888	47,1
9	715	44,7
10	690	45,5

Dari data pada Tabel 3, yang menunjukkan hasil pengukuran waktu respons dan penggunaan CPU dengan penerapan *microservice*, terlihat bahwa waktu respons bervariasi antara 660 ms hingga 936 ms, dengan iterasi keempat mencatat waktu tercepat (660 ms) dan iterasi keenam mencatat waktu terlama (936 ms). Penggunaan CPU juga menunjukkan variasi, mulai dari 39,1% pada iterasi kelima hingga 48,5% pada iterasi keenam. Secara umum, ada kecenderungan bahwa waktu respons yang lebih tinggi terkait dengan penggunaan CPU yang lebih tinggi, mengindikasikan bahwa beban pemrosesan API dapat meningkatkan kebutuhan CPU dan mempengaruhi kinerja.

Perbandingan performa dalam penelitian ini sudah mempertimbangkan kesamaan spesifikasi perangkat keras dan lunak, namun masih ada potensi ketimpangan dalam distribusi sumber daya, terutama pada layanan yang mengalami *bottleneck* seperti Eco-Map dan Waste Management. Variabilitas waktu respons dan penggunaan CPU menunjukkan adanya faktor eksternal yang dapat mempengaruhi hasil, seperti manajemen sumber daya Kubernetes dan latensi database. Untuk meningkatkan *fairness*, perlu diterapkan optimasi seperti caching, auto-scaling, dan *load balancing*. Secara keilmuan, penelitian ini berdampak bagi akademisi dalam pengembangan arsitektur *microservices*, pengembang sistem dalam optimasi performa layanan, serta institusi pendidikan dan industri dalam implementasi smart campus.

Pada Gambar 6 menunjukkan perbandingan kinerja sistem sebelum dan setelah optimasi. Grafik kiri menampilkan waktu respons, di mana setelah optimasi terlihat penurunan waktu respons yang lebih stabil dan lebih rendah. Grafik kanan menunjukkan penggunaan CPU, yang mengalami efisiensi lebih baik setelah optimasi, mengindikasikan distribusi beban yang lebih optimal. Perbaikan ini dapat dicapai melalui penerapan caching, auto-scaling, dan *load balancing* untuk mengatasi *bottleneck* dalam layanan.



Gambar 6 Perbandingan Waktu Respon dan Penggunaan CPU



4. KESIMPULAN

Berdasarkan hasil percobaan, pengujian Load Test menunjukkan bahwa sistem memiliki performa yang stabil, dengan waktu respons rata-rata antara 100–160 ms dan *error rate* rendah (0,1%–0,8%). Pada Stress Test, layanan Auth Service dan NATS Streaming tetap menunjukkan performa optimal dengan *error rate* masing-masing 10% dan 5%, sedangkan Eco-Map Service dan Waste Management Service mengalami *bottleneck*, dengan waktu respons 150–180 ms dan *error rate* hingga 15%–20%, yang disebabkan oleh tingginya konsumsi memori dan latensi database hingga 120 ms.

Hasil pengukuran juga mencakup waktu respons dan penggunaan CPU. Waktu respons rata-rata dari sepuluh percobaan dengan penerapan API adalah 745,9 ms, dengan kisaran antara 660 ms hingga 936 ms. Iterasi keenam mencatat waktu respons tertinggi (936 ms), sedangkan iterasi keempat mencatat waktu respons terendah (660 ms). Penggunaan CPU bervariasi, dengan nilai terendah sebesar 39,1% pada iterasi kelima dan tertinggi 48,5% pada iterasi keenam. Pola ini menunjukkan bahwa waktu respons yang lebih lama cenderung terjadi saat penggunaan CPU tinggi, menandakan bahwa beban pemrosesan berdampak langsung pada efisiensi CPU dan kinerja keseluruhan sistem.

Meskipun arsitektur *microservice* memberikan fleksibilitas dan skalabilitas, optimasi lebih lanjut pada database, alokasi sumber daya, dan implementasi sistem caching diperlukan untuk meningkatkan performa seluruh layanan. Selain itu, penerapan auto-scaling juga diperlukan untuk menangani beban tinggi. Untuk penelitian selanjutnya, disarankan fokus pada optimasi arsitektur API-driven guna mengurangi fluktuasi waktu respons dan penggunaan CPU. Penelitian lanjutan dapat mencakup pengelolaan beban kerja yang lebih baik, eksplorasi teknologi optimasi seperti caching dan *load balancing*, serta perbandingan performa dengan arsitektur lain, misalnya Event-Driven Architecture (EDA). Pendekatan ini diharapkan dapat menemukan metode yang lebih efektif untuk meningkatkan efisiensi dan konsistensi kinerja sistem berbasis API-driven.

UCAPAN TERIMA KASIH

Ucapan terima kasih kepada Lembaga Penelitian dan Pengabdian kepada Masyarakat (LPPM) Universitas Siliwangi atas dukungan dana dan berbagai bantuan yang diberikan dalam skema penelitian ini. Kami juga berterima kasih kepada seluruh tim penelitian atas kerja keras, kolaborasi, dan dedikasi tinggi dalam setiap tahap penelitian ini. Kami berharap hasilnya dapat memberikan kontribusi berarti bagi perkembangan ilmu pengetahuan dan teknologi serta bermanfaat bagi masyarakat.

DAFTAR PUSTAKA

- Aasa, O. P., Jesuleye, O. A., & Ajayi, M. O. (2020). Towards Greening Decisions on the University Campus: Initiatives, Importance and Barriers. *International Journal of Engineering and Management Research*, 10(3), 82–88. <https://doi.org/10.31033/ijemr.10.3.13>
- Abdullai, L., Porras, J., & Haque, S. (2022). Building a National Smart Campus to Support Sustainable Business Development: An Ecosystem Approach. *ArXiv Preprint*. <http://arxiv.org/abs/2209.13613>
- Abdilmouti, H., Skaf, Z., & Alblooshi, S. (2022). Smart Green Campus: The Campus of Tomorrow. *2022 Advances in Science and Engineering Technology International Conferences (ASET)*, 1–8. <https://doi.org/10.1109/ASET53988.2022.9735087>
- Amelia, A., Sundawa, B. V., Yusoff, M., Mohamad Zain, J. B., Azis, A., Suprianto, Roslina, R., & Pribadi, B. A. (2023). Performance Analysis of Environmental Monitoring System (EMS) towards POLMEDs Green Campus. *International Journal on Advanced Science, Engineering and Information Technology*, 13(5), 1727–1732. <https://doi.org/10.18517/ijaseit.13.5.19481>
- Dimov, A., Emanuilov, S., Bontchev, B., Dankov, Y., & Papapostolu, T. (2022). Architectural Approaches to Overcome Challenges in the Development of Data-Intensive Systems. In T.



- Ahram (Ed.), *Human Factors in Software and Systems Engineering* (Vol. 61, pp. 38–43). AHFE International. <https://doi.org/10.54941/ahfe1002521>
- Faizah, A. N., & Nugraheni, N. (2024). Pendidikan Berkelanjutan Berbasis Konservasi dan Teknologi Sebagai Aksi Nyata dalam Mewujudkan SDGs. *Jurnal Penelitian Ilmu-Ilmu Sosial*, 1(10), 73–80. <https://doi.org/10.5281/zenodo.11152410>
- Hong, X. J., Sik Yang, H., & Kim, Y. H. (2018). Performance Analysis of RESTful API and RabbitMQ for Microservice Web Application. *2018 International Conference on Information and Communication Technology Convergence (ICTC)*, 257–259. <https://doi.org/10.1109/ICTC.2018.8539409>
- Iqbal, M., Syahputra, A. K., & Handoko, W. (2023). Penerapan Monolithic Architecture pada Aplikasi Ujian Online BERBASIS Berbasis Web. *Jurnal Pemberdayaan Sosial dan Teknologi Masyarakat*, 2(2), 213–218. <https://doi.org/10.54314/jpstm.v2i2.1107>
- Kautsar, I. A., Maika, M. R., Budiman, A. N., Setyawan, A. B., & Awali, J. Y. (2023). Microservice Based Architecture: The Development of Rapid Prototyping Supportive Tools for Project Based Learning. *2023 IEEE World Engineering Education Conference (EDUNINE)*, 1–6. <https://doi.org/10.1109/EDUNINE57531.2023.10102884>
- Khan, A., & Ximei, W. (2022). Digital Economy and Environmental Sustainability: Do Information Communication and Technology (ICT) and Economic Complexity Matter? *International Journal of Environmental Research and Public Health*, 19(19), Article ID: 12301. <https://doi.org/10.3390/ijerph191912301>
- Martins, P., Lopes, S. I., Rosado da Cruz, A. M., & Curado, A. (2021). Towards a Smart & Sustainable Campus: An Application-Oriented Architecture to Streamline Digitization and Strengthen Sustainability in Academia. *Sustainability*, 13(6), Article ID: 3189. <https://doi.org/10.3390/su13063189>
- Palomo, I., Willemen, L., Drakou, E., Burkhard, B., Crossman, N., Bellamy, C., Burkhard, K., Campagne, C. S., Dangol, A., Franke, J., Kulczyk, S., Le Clec'h, S., Abdul Malak, D., Muñoz, L., Narusevicius, V., Ottoy, S., Roelens, J., Sing, L., Thomas, A., ... Verweij, P. (2018). Practical Solutions for Bottlenecks in Ecosystem Services Mapping. *One Ecosystem*, 3, Article ID: e20713. <https://doi.org/10.3897/oneeco.3.e20713>
- Ribeiro, J. M. P., Hoeckesfeld, L., Dal Magro, C. B., Favretto, J., Barichello, R., Lenzi, F. C., Secchi, L., de Lima, C. R. M., & de Andrade Guerra, J. B. S. O. (2021). Green Campus Initiatives as Sustainable Development Dissemination at Higher Education Institutions: Students' Perceptions. *Journal of Cleaner Production*, 312, Article ID: 127671. <https://doi.org/10.1016/j.jclepro.2021.127671>
- Santos, L. C. dos, da Silva, M. L. P., & dos Santos Filho, S. G. (2024). Sustainability in Industry 4.0: Edge Computing Microservices as a New Approach. *Sustainability*, 16(24), Article ID: 11052. <https://doi.org/10.3390/su162411052>
- Sidiq, M. A. Z., Anshori, M. I., & Yaqin, R. A. (2024). Penerapan Arsitektur Monolitik pada Aplikasi Jasa Service Online Tekku Berbasis Web. *JUKI: Jurnal Komputer dan Informatika*, 6(1), 27–36. <https://doi.org/10.53842/juki.v6i1.418>
- Sugiarto, A., Lee, C.-W., & Huruta, A. D. (2022). A Systematic Review of the Sustainable Campus Concept. *Behavioral Sciences*, 12(5), Article ID: 130. <https://doi.org/10.3390/bs12050130>
- Telang, T. (2023). Microservices Architecture. In *Beginning Cloud Native Development with MicroProfile, Jakarta EE, and Kubernetes* (pp. 111–142). Apress. https://doi.org/10.1007/978-1-4842-8832-0_5
- Trevisan, L. V., Eustachio, J. H. P. P., Dias, B. G., Filho, W. L., & Pedrozo, E. Á. (2023). Digital Transformation Towards Sustainability in Higher Education: State-of-the-Art and Future Research Insights. *Environment, Development and Sustainability*, 26(2), 2789–2810. <https://doi.org/10.1007/s10668-022-02874-7>
- Weerasinghe, L. D. S. B., & Perera, I. (2024). Reference Architecture for Microservices with an Optimized Inter-Service Communication Strategy. *2024 International Research Conference on Smart Computing and Systems Engineering (SCSE)*, 1–6. <https://doi.org/10.1109/SCSE61872.2024.10550466>
- Xiao, X. (2024). Architectural Tactics to Improve the Environmental Sustainability of Microservices: A Rapid Review. *ArXiv Preprint*. <http://arxiv.org/abs/2407.16706>



Analisis Efektivitas Metode Filtering dan Intersection dalam Analisis Data Permukaan Bangunan dengan QGIS

Prana Wijaya Pratama Nandana ^{(1)*}, Muhammad Faisal ⁽²⁾

Departemen Teknik Informatika, UIN Maulana Malik Ibrahim, Malang, Indonesia

e-mail : 210605110120@student.uin-malang.ac.id, mfaisal@ti.uin-malang.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 18 September 2024, direvisi 20 Februari 2025, diterima 21 Februari 2025, dan dipublikasikan 30 September 2025.

Abstract

This study evaluates the efficiency of two methods for processing geospatial building surface data, namely Filtering and Intersection, using a case study in Blitar Regency. The data for this research was obtained by comparing two sources: OpenStreetMap (OSM), which has a data completeness rate of 60%, and Google Open Building, with a data completeness rate of 90%. From these two sources, the data with the highest completeness, which is from Google Open Building, was selected for further analysis. The data processing was carried out using QGIS software, chosen for its capability to support various geospatial analysis methods. The comparison of the two methods was based on three main criteria: processing time, resource efficiency, and scalability. The results showed that the Filtering method outperforms in all these aspects. Filtering can complete processing in an average of 1.6 seconds, significantly faster than the Intersection method, which requires an average of 7 minutes and 50 seconds. In terms of resource efficiency, Filtering is also more economical, with an average CPU usage of 18.85% and memory usage of 121.4 MB, compared to the Intersection method's 34.05% CPU usage and 236.4 MB of memory. Additionally, the Filtering method demonstrated better scalability, capable of handling larger datasets with fewer resources and less time. Therefore, the Filtering method is recommended for geospatial data processing that prioritizes speed, efficiency, and the ability to handle large and complex datasets.

Keywords: Filtering, Intersection, Kabupaten Blitar, Comparing, QGIS

Abstrak

Efisiensi dua metode pengolahan data geospasial permukaan bangunan, yaitu Filtering dan Intersection, dievaluasi dalam penelitian ini dengan menggunakan studi kasus di Kabupaten Blitar. Data untuk penelitian ini diperoleh dengan melakukan perbandingan antara dua sumber data, yakni OpenStreetMap (OSM), yang memiliki tingkat kelengkapan 60%, dan Google Open Building, yang tingkat kelengkapannya mencapai 90%. Dari kedua sumber ini, data dengan tingkat kelengkapan tertinggi, yaitu dari Google Open Building, dipilih untuk analisis lebih lanjut. Pengolahan data dilakukan menggunakan perangkat lunak QGIS, yang dipilih karena kemampuannya dalam mendukung berbagai metode analisis geospasial. Perbandingan kedua metode dilakukan berdasarkan tiga kriteria utama: waktu pemrosesan, efisiensi penggunaan sumber daya, dan skalabilitas. Hasil analisis menunjukkan bahwa metode Filtering memiliki keunggulan dalam semua aspek ini. Metode Filtering mampu menyelesaikan pemrosesan rata-rata dalam waktu 1,6 detik, jauh lebih cepat dibandingkan dengan metode Intersection yang membutuhkan rata-rata 7 menit 50 detik. Dari segi efisiensi sumber daya, Filtering juga lebih hemat dengan penggunaan rata-rata CPU sebesar 18,85% dan memori sebesar 121,4 MB, dibandingkan dengan penggunaan CPU sebesar 34,05% dan memori sebesar 236,4 MB oleh metode Intersection. Selain itu, metode Filtering juga menunjukkan skalabilitas yang lebih baik, mampu menangani *dataset* yang lebih besar dengan waktu dan sumber daya yang lebih sedikit. Oleh karena itu, metode Filtering lebih direkomendasikan untuk pengolahan data geospasial yang memprioritaskan kecepatan, efisiensi, serta kemampuan dalam menangani *dataset* yang besar dan kompleks.

Kata Kunci: Filtering, Intersection, Kabupaten Blitar, Perbandingan, QGIS



1. PENDAHULUAN

Data geospasial adalah data yang berkaitan dengan lokasi atau tempat di permukaan bumi (Lochana M et al., 2024). Data ini mencakup informasi tentang objek, fenomena, atau kondisi yang memiliki aspek geografis, seperti koordinat, bentuk, dan ukuran (Ivanov, 2020). Data geospasial digunakan untuk membuat peta, menganalisis, dan memahami berbagai aspek lingkungan dan aktivitas manusia (F & Saleh, 2023). Secara umum, data geospasial terbagi menjadi dua jenis utama: data vektor, yang menggambarkan fitur geografis seperti titik (misalnya, lokasi bangunan), garis (seperti jalan), dan area (misalnya, wilayah atau lahan), serta data raster, yang berupa grid atau piksel dengan nilai tertentu, seperti ketinggian tanah atau gambar satelit (Singla et al., 2021). Data ini sering digunakan dalam Sistem Informasi Geografis (SIG) untuk berbagai keperluan, seperti analisis spasial, perencanaan tata ruang, pengelolaan sumber daya alam, pemantauan lingkungan, dan berbagai aplikasi lainnya yang membutuhkan pemahaman tentang lokasi dan distribusi geografis (Wahab & Kurniawan, 2023).

Agar data geospasial dapat dimanfaatkan, diperlukan perangkat lunak atau alat untuk pengolahan data geospasial salah satunya adalah QGIS. QGIS (Quantum GIS) adalah perangkat lunak gratis yang digunakan untuk membuat, mengedit, dan menganalisis data peta atau data geospasial (Flenniken et al., 2020). Dengan QGIS, pengguna dapat membuat peta interaktif, mengolah data lokasi, dan melakukan analisis geografis, seperti menghitung jarak atau melihat tumpang tindih antara area (Duarte et al., 2021). QGIS juga menawarkan berbagai metode untuk menganalisis dan mengolah data spasial, termasuk metode Filtering dan Intersection, yang sering digunakan dalam pengolahan data permukaan bangunan (Wegmann et al., 2020).

Teknik Filtering dalam pengolahan data geospasial telah terbukti efektif dalam berbagai penelitian, terutama dalam analisis fenomena Urban Heat Island (UHI). Seperti penelitian yang dilakukan oleh Darmawan et al. (2024) menggunakan platform Google Earth Engine untuk memonitor UHI di Indonesia. Dalam penelitian tersebut, teknik Filtering diterapkan pada pengolahan data citra satelit MODIS untuk memastikan proses yang lebih cepat dan efisien, dengan hanya memproses data yang relevan untuk estimasi suhu permukaan dan intensitas UHI. Hasil dari penelitian ini menunjukkan bahwa teknik Filtering dapat mengurangi noise dalam data geospasial dan meningkatkan akurasi visualisasi serta analisis UHI. Oleh karena itu, teknik Filtering dianggap sebagai pilihan yang tepat untuk pengolahan data permukaan bangunan dalam penelitian ini.

Kemudian untuk metode Intersection diterapkan pada penelitian oleh Wismarini & Khristianto (2016) membahas implementasi metode superimpose, khususnya teknik Intersection, dalam pemodelan spasial tingkat kerentanan banjir di Kota Semarang. Dengan menggunakan perangkat lunak Sistem Informasi Geografis (SIG) seperti ArcView 3.3, penelitian ini memanfaatkan analisis *overlay* untuk menggabungkan berbagai data indikator banjir, termasuk kemiringan lereng, struktur tanah, tata guna lahan, dan curah hujan. Metode Intersection digunakan untuk menghasilkan peta tematik yang menunjukkan zonasi tingkat rentan banjir, dengan mempertimbangkan area yang benar-benar tergenang banjir. Hasil penelitian ini menunjukkan efektivitas metode Intersection dalam menghasilkan peta yang tidak hanya menggambarkan potensi banjir, tetapi juga memberikan identifikasi lebih rinci mengenai tingkat kerentanannya. Peta tematik yang dihasilkan memuat informasi yang sangat berguna untuk identifikasi zona dengan tingkat kerentanan banjir yang berbeda, memberikan data yang sangat informatif untuk perencanaan mitigasi bencana. Hasil akhir penelitian ini berupa peta digital zonasi tingkat rentan banjir yang dapat menginformasikan tingkat kerentanan di setiap zona, dengan data kualitatif yang didasarkan pada indikator geomorfologi dan data curah hujan.

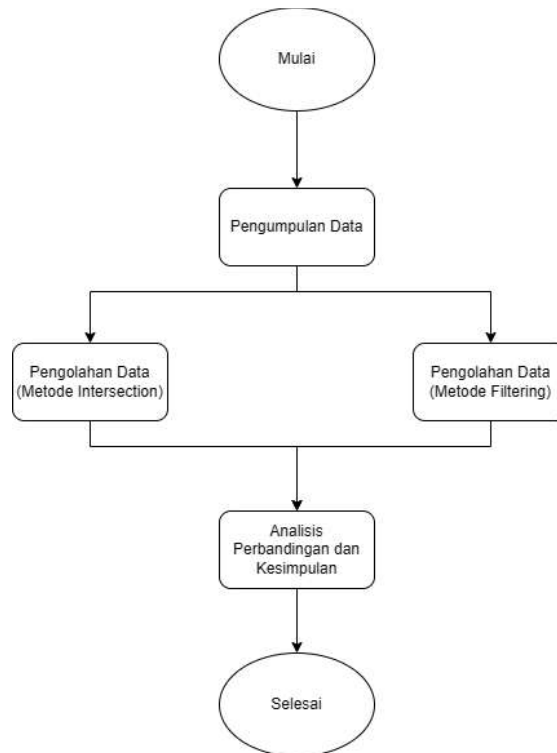
Karena terdapat berbagai metode pada software QGIS, maka penting untuk memilih metode yang sesuai dengan pengolahan data geospasial khususnya data permukaan bangunan (Kukulska et al., 2018). Oleh karena itu, berdasarkan penelitian sebelumnya pada penelitian ini akan dilakukan analisis pada metode Intersection dan metode Filtering untuk mendapatkan metode mana yang memiliki efektivitas paling tinggi dalam pengolahan data geospasial



khususnya data permukaan bangunan. Hasil dari penelitian ini diharapkan dapat memberikan kontribusi praktis dan teoretis. Secara praktis, penelitian ini akan membantu pengguna software GIS dalam memilih metode pengolahan data yang paling efisien untuk proyek mereka, yang dapat menghemat waktu dan sumber daya. Secara teoretis, penelitian ini akan memperkaya literatur mengenai analisis metode pengolahan data geospasial, khususnya dalam konteks pengolahan data permukaan bangunan.

2. METODE PENELITIAN

2.1 Alur Penelitian



Gambar 1 Alur Penelitian

Gambar 1 menunjukkan alur pada penelitian ini. Penelitian dimulai dengan mengidentifikasi tujuan dan mengumpulkan data geospasial dari OpenStreetMap (OSM) dan Google Open Building, diikuti dengan validasi dan perbandingan kelengkapan data untuk menentukan sumber yang terbaik (Sutanto & Aditya, 2021). Studi kasus yang digunakan dalam penelitian ini adalah Kabupaten Blitar. Data yang terpilih kemudian diolah menggunakan QGIS dengan menerapkan metode Filtering dan Intersection (Khajedehi et al., 2024). Hasil dari kedua metode dianalisis dan dibandingkan untuk menilai efisiensi dari masing-masing metode. Selanjutnya, kesimpulan akan ditarik untuk menentukan metode yang lebih efisien untuk digunakan di antara kedua metode tersebut.

2.2 Metode Pengumpulan Data

Dalam penelitian ini, metode pengumpulan data dilakukan dengan mengidentifikasi dan membandingkan dua sumber data geospasial yang utama, yaitu OpenStreetMap (OSM) dan Google Open Building (Fram et al., 2015). Proses ini diawali dengan mengunduh data bangunan dari kedua sumber tersebut melalui perangkat lunak QGIS dan layanan berbasis web terkait, seperti website *Hot Export Tool* untuk OSM dan website Open Buildings untuk Google Open Building (Brovelli et al., 2018). Data yang diperoleh kemudian dianalisis untuk menentukan kelengkapan dan keakuratannya (Zhou, 2018).



Untuk memastikan bahwa data yang digunakan dalam penelitian ini valid dan akurat, dilakukan beberapa tahapan proses validasi sebelum analisis lebih lanjut. Langkah pertama adalah memeriksa kelengkapan dan konsistensi data dari dua sumber, yaitu OpenStreetMap (OSM) dan Google Open Building. Data dari kedua sumber dibandingkan berdasarkan jumlah objek bangunan yang tersedia dalam cakupan wilayah studi di Kabupaten Blitar. Hasil perbandingan menunjukkan bahwa Google Open Building memiliki kelengkapan data sebesar 90%, sedangkan OSM hanya 60% dari jumlah keseluruhan jumlah bangunan, sehingga Google Open Building dipilih sebagai sumber utama.

Setelah sumber data ditentukan, dilakukan pembersihan dan penyaringan data. Langkah ini mencakup penghapusan duplikasi bangunan, pengecekan kesalahan geometri, serta validasi atribut data seperti luas permukaan dan koordinat lokasi. Selanjutnya, diterapkan *filter* spasial untuk memastikan bahwa hanya bangunan yang berada dalam batas administratif Kabupaten Blitar yang digunakan dalam analisis. Hal ini dilakukan dengan menggunakan fungsi *clip* pada QGIS, yang memungkinkan pemotongan data berdasarkan batas wilayah yang telah ditentukan.

Selain itu, dilakukan validasi visual dengan membandingkan hasil *overlay* data bangunan terhadap citra satelit dari Google Earth. Langkah ini bertujuan untuk memastikan bahwa bangunan yang teridentifikasi dalam *dataset* sesuai dengan kondisi nyata di lapangan. Jika ditemukan ketidaksesuaian atau bangunan yang tidak relevan, data tersebut dihapus dari analisis. Dengan pendekatan ini, penelitian memastikan bahwa hanya data bangunan yang valid, akurat, dan sesuai dengan wilayah studi yang digunakan dalam perbandingan metode Filtering dan Intersection.

2.3 Metode Pengolahan Data

Dalam penelitian ini, metode Filtering dan Intersection dipilih karena keduanya merupakan pendekatan yang umum digunakan dalam analisis data geospasial pada Sistem Informasi Geografis (SIG), terutama dalam pemrosesan data bangunan. Kedua metode ini memungkinkan pemrosesan dan analisis data spasial dengan pendekatan yang berbeda, sehingga diperlukan evaluasi lebih lanjut terkait efisiensi dan skalabilitasnya. Perbandingan dilakukan untuk melihat bagaimana efektivitas masing-masing metode dalam menangani data dalam konteks pemrosesan geospasial.

Metode Intersection adalah teknik dalam GIS yang digunakan untuk menemukan area yang tumpang tindih antara dua lapisan data spasial, seperti lapisan bangunan dan lapisan batas administrasi (Memduhoglu & Basaraner, 2022). Dalam konteks penelitian ini, metode Intersection diterapkan untuk memproses data permukaan bangunan di Kabupaten Blitar. Rumus konseptual yang digunakan dalam metode ini ditunjukkan pada Pers. (1). Rumus tersebut menjelaskan bahwa himpunan $A \cap B$ terdiri atas elemen x yang berada sekaligus pada himpunan A dan B. Dalam konteks penelitian ini, himpunan A merepresentasikan area bangunan, sedangkan himpunan B menunjukkan batas wilayah administratif. Dengan demikian, $A \cap B$ adalah area bangunan yang terletak di dalam batas wilayah yang dianalisis. Hasil dari metode Intersection ini akan memberikan informasi mengenai bangunan yang berada di wilayah yang dibutuhkan (El-Ashmawy, 2019). Efisiensi dari metode ini kemudian akan dibandingkan dengan metode Filtering dalam penelitian ini (El-Kareem & El-Emary, 2019).

$$A \cap B = \{ x \mid x \in A \text{ dan } x \in B \} \quad (1)$$

Metode Filtering diterapkan untuk menyaring data berdasarkan kriteria tertentu sehingga hanya data yang memenuhi syarat yang akan diproses lebih lanjut (Raju, 2004). Dalam penelitian ini, metode Filtering digunakan untuk memilih id bangunan yang berada dalam batas administrasi wilayah yang diinginkan serta bangunan yang memenuhi kriteria tertentu, seperti luas permukaan minimum atau jenis bangunan (Afnarius et al., 2020). Filtering dilakukan dengan menggunakan alat *filter* yang ada di QGIS. Rumus konseptual yang digunakan dalam metode ini ditunjukkan pada Pers. (2). Rumus tersebut menjelaskan bahwa $F(x)$ merupakan himpunan data x yang



memenuhi kriteria K . Dalam konteks penelitian ini, x merepresentasikan objek data seperti bangunan, sedangkan K adalah kriteria tertentu yang digunakan sebagai dasar penyaringan, seperti batas administrasi. Tujuan dari metode Filtering ini adalah untuk mempercepat proses pengolahan data dengan hanya memfokuskan pada data yang relevan (Rajab & Wang, 2020). Efisiensi metode Filtering akan dievaluasi dan dibandingkan dengan metode Intersection, baik dari segi kecepatan proses maupun hasil akhir yang diperoleh.

$$F(x) = \{ x \mid x \text{ memenuhi kriteria } K \} \quad (2)$$

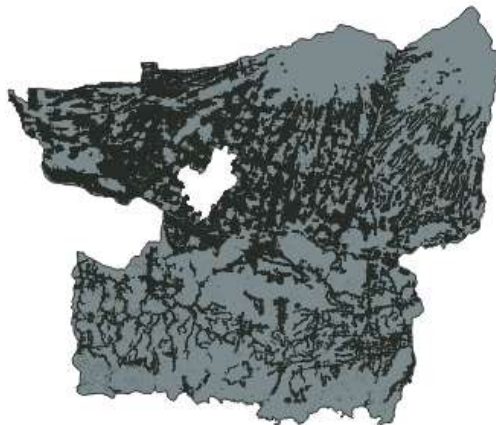
2.4 Perbandingan Metode Intersection dan Filtering

Pada penelitian ini untuk membandingkan hasil dari metode Intersection dan Filtering dengan menggunakan tiga kriteria utama, yaitu waktu pemrosesan, efisiensi sumber daya, dan skalabilitas. Waktu pemrosesan diukur untuk mengetahui seberapa cepat masing-masing metode dapat menyelesaikan tugas pengolahan data geospasial (Goodman et al., 2019). Efisiensi sumber daya dievaluasi dengan menilai bagaimana setiap metode mengoptimalkan penggunaan sumber daya sistem dalam perangkat lunak QGIS, termasuk penggunaan CPU dan memori, serta dampak keseluruhan terhadap kinerja sistem (Auradkar et al., 2022). Skalabilitas dinilai dengan menguji kemampuan kedua metode dalam menangani data dalam skala besar, serta bagaimana kinerja masing-masing metode ketika volume data yang diproses meningkat (Yogi et al., 2024).

Untuk memastikan hasil perbandingan yang adil dan bebas dari gangguan eksternal, seluruh eksperimen dilakukan pada perangkat dengan spesifikasi tetap, yaitu prosesor Intel Core i5-7200U (2.50 GHz), RAM 12 GB DDR4, penyimpanan 1 TB HDD, serta sistem operasi Windows 10 Pro 64-bit. Selain itu, proses latar belakang yang tidak diperlukan, seperti pembaruan otomatis dan aplikasi yang berjalan di *background*, dinonaktifkan sebelum pengujian dimulai. Setiap eksperimen dijalankan sebanyak tiga kali, dan hasil rata-rata digunakan untuk memastikan konsistensi pengukuran. Dengan pendekatan ini, penelitian dapat memberikan evaluasi yang objektif terhadap keunggulan dan keterbatasan masing-masing metode dalam pengolahan data geospasial.

3. HASIL DAN PEMBAHASAN

Pada penelitian ini, peneliti menggunakan metode Filtering dan Intersection untuk mengolah data bangunan di wilayah Kabupaten Blitar. Peneliti membandingkan hasil dari kedua metode ini berdasarkan waktu pemrosesan, efisiensi sumber daya, dan skalabilitas. Untuk memudahkan pemahaman, Gambar 2 menunjukkan seperti apa data permukaan bangunan di Kabupaten Blitar, sementara Tabel 3 di Lampiran A menjelaskan atribut-atribut dari data geospasial yang digunakan.



Gambar 2 Data Geospasial Permukaan Bangunan Kabupaten Blitar



Gambar 2 menampilkan representasi geospasial permukaan bangunan di Kabupaten Blitar. Area berwarna hitam menunjukkan distribusi dan kepadatan bangunan, sedangkan area berwarna abu-abu merepresentasikan batas wilayah administratif Kabupaten Blitar. Sementara itu, Tabel 3 di Lampiran A menyajikan atribut-atribut penting dari data geospasial terkait permukaan bangunan di wilayah tersebut. Tabel ini memuat informasi terstruktur dan mendetail mengenai aspek geografis serta karakteristik bangunan, dengan setiap kolom memiliki peran khusus dalam memberikan informasi yang lebih komprehensif. Penjelasan untuk masing-masing kolom adalah sebagai berikut:

- 1) **Latitude**: Menunjukkan koordinat garis lintang dari lokasi tertentu.
- 2) **Longitude**: Menunjukkan koordinat garis bujur dari lokasi tertentu.
- 3) **Area_in_me**: Menggambarkan luas area dalam satuan meter persegi (m^2) untuk setiap titik bangunan.
- 4) **Confidence**: Berisi tingkat kepercayaan atau akurasi dari data yang tercatat, dinyatakan dalam bentuk desimal.
- 5) **Full_plus**: Menampilkan kode alfanumerik yang berfungsi sebagai identifikasi unik atau kode alamat khusus untuk setiap lokasi.
- 6) **ID_KAB**: Merupakan kode identifikasi kabupaten yang sesuai dengan data pada tabel.

Secara keseluruhan, Tabel 3 memberikan informasi geospasial yang rinci, mulai dari koordinat geografis, luas area, hingga tingkat kepercayaan data, sehingga mendukung analisis spasial terkait distribusi bangunan di Kabupaten Blitar.

3.1 Waktu Pemrosesan

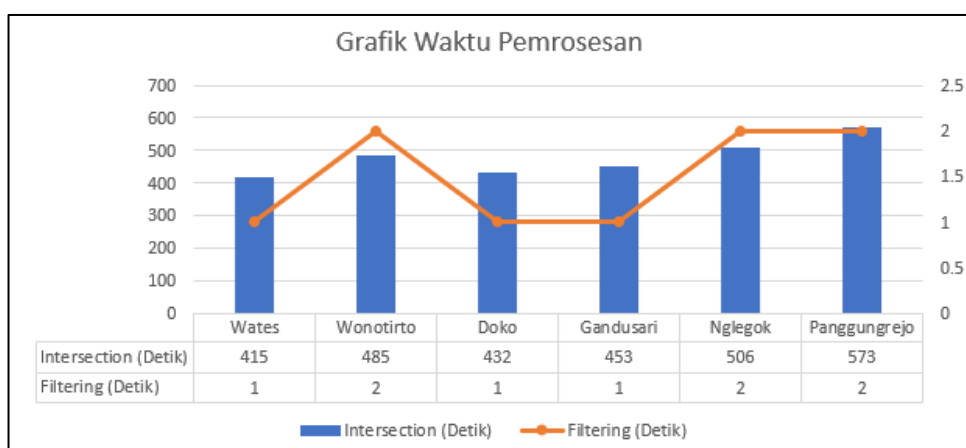
Tabel 1 menyajikan analisis komparatif mengenai durasi pemrosesan data menggunakan dua metode yang berbeda, yaitu Intersection dan Filtering, sejumlah kecamatan yang berada di wilayah Kabupaten Blitar dijadikan studi kasus. Kecamatan yang disertakan dalam studi ini meliputi Wates, Wonotirto, Doko, Gandusari, Nglegok, dan Panggungrejo. Dalam tabel ini, metode Intersection menunjukkan variasi waktu pemrosesan yang cukup signifikan, dengan durasi yang berkisar antara 6 hingga 9 menit. Sebagai contoh, Kecamatan Wates memerlukan waktu 6 menit 55 detik, sementara Kecamatan Panggungrejo membutuhkan waktu paling lama, yakni 9 menit 33 detik. Variasi ini mencerminkan kompleksitas dan volume data yang berbeda di setiap kecamatan, yang mempengaruhi waktu pemrosesan ketika menggunakan metode Intersection.

Tabel 1 Hasil Perbandingan Waktu Pemrosesan Data

Kecamatan	Durasi Pemrosesan	
	Intersection	Filtering
Wates	06 Menit 55 Detik (415 Detik)	1 Detik
Wonotirto	08 Menit 05 Detik (485 Detik)	2 Detik
Doko	07 Menit 12 Detik (432 Detik)	1 Detik
Gandusari	07 Menit 33 Detik (453 Detik)	1 Detik
Nglegok	08 Menit 26 Detik (506 Detik)	2 Detik
Panggungrejo	09 Menit 33 Detik (573 Detik)	2 Detik

Sebaliknya, metode Filtering menunjukkan keunggulan yang nyata dalam hal efisiensi waktu pemrosesan. Di semua kecamatan yang diteliti, metode ini mampu menyelesaikan proses hanya dalam 1 hingga 2 detik, sebuah pencapaian yang sangat efisien dibandingkan dengan metode Intersection. Perbedaan waktu yang mencolok ini menegaskan bahwa metode Filtering adalah pilihan yang lebih efektif ketika kecepatan pemrosesan menjadi prioritas utama. Dengan hasil yang menunjukkan efisiensi waktu yang luar biasa dari metode Filtering, dapat disimpulkan bahwa metode ini lebih cocok digunakan dalam situasi di mana pemrosesan cepat sangat dibutuhkan, mengurangi waktu dan sumber daya yang diperlukan untuk analisis data geospasial di setiap kecamatan yang diteliti. Untuk memberikan gambaran yang lebih jelas mengenai perbedaan hasil antara kedua metode, visualisasi hasil dalam bentuk grafik dapat dilihat pada Gambar 3, yang menunjukkan perbandingan waktu pemrosesan antara Filtering dan Intersection pada berbagai kecamatan.





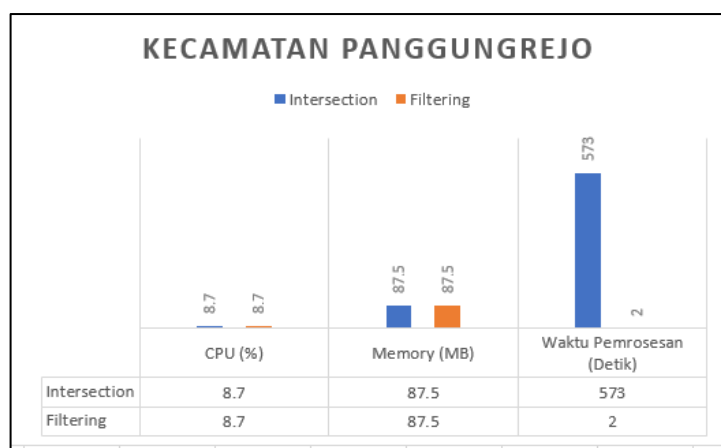
Gambar 3 Grafik Waktu Pemrosesan

3.2 Efisiensi Sumber Energi

Dalam hal efisiensi sumber daya yang ditunjukkan pada Tabel 2, metode Filtering secara keseluruhan lebih unggul dibandingkan metode Intersection. Pengujian menunjukkan bahwa metode Filtering lebih hemat dalam penggunaan CPU, memori, dan waktu pemrosesan, serta memiliki konsumsi daya yang lebih optimal dibandingkan dengan Intersection, yang cenderung memerlukan sumber daya lebih besar. Hasil ini mengindikasikan bahwa metode Filtering lebih efisien untuk pemrosesan data geospasial yang membutuhkan optimalisasi sumber daya. Untuk melihat perbandingan lebih jelas antara kedua metode, maka akan ditunjukkan menggunakan grafik untuk setiap poin pengukuran sumber daya.

Tabel 2 Hasil Pengukuran Sumber Daya Menggunakan *Task Manager*

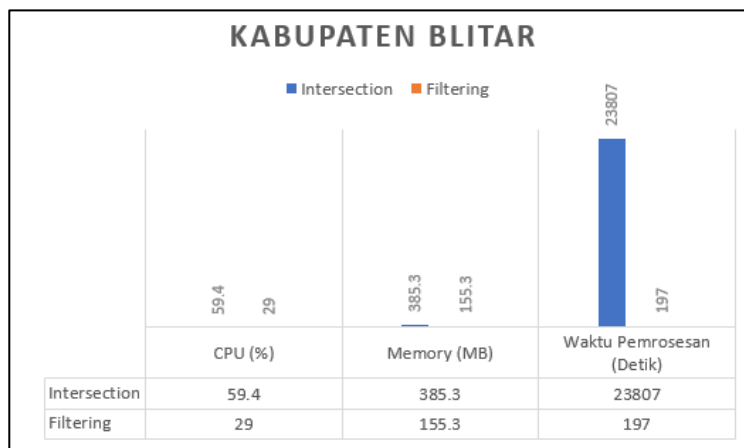
Wilayah	Metode	CPU	Memory	Waktu Pemrosesan	Power Usage
Kecamatan Panggungrejo	Intersection	8,7%	87,5 MB	09 Menit 33 Detik (573 Detik)	Very High
	Filtering	8,7%	87,5 MB	2 Detik	Moderate
Kabupaten Blitar	Intersection	59,4%	385,3 MB	06 Jam 38 Menit 27 Detik (23.907 Detik)	Very High
	Filtering	29,0%	155.3 MB	03 Menit 17 Detik (197 Detik)	Very High



Gambar 4 Grafik Penggunaan Sumberdaya pada Wilayah Kecamatan Panggungrejo



Pada Gambar 4 menampilkan grafik terkait penggunaan sumber daya sistem (CPU, Memori, Waktu Pemrosesan, dan Penggunaan Daya) untuk dua metode yang diterapkan pada wilayah “Kecamatan Panggungrejo.” Pada metode pertama, yaitu Intersection, penggunaan CPU tercatat sebesar 8,7%, dengan memori sebesar 87,5 MB, dan waktu pemrosesan mencapai 09 menit 33 detik. Metode ini juga menunjukkan penggunaan daya yang sangat tinggi. Sedangkan pada metode kedua, yaitu Filtering, penggunaan CPU tetap pada angka yang sama, yaitu 8,7%, dengan memori yang identik, yakni 87,5 MB. Namun, waktu pemrosesan untuk metode Filtering jauh lebih cepat, yaitu hanya 2 detik, dengan penggunaan daya yang lebih optimal dibandingkan metode Intersection. Perbedaan waktu pemrosesan dan penggunaan daya ini memberikan gambaran tentang efisiensi masing-masing metode dalam hal kecepatan dan konsumsi energi.



Gambar 5 Grafik Penggunaan Sumberdaya pada Wilayah Kabupaten Blitar

Pada Gambar 5 menampilkan grafik mengenai penggunaan sumber daya sistem (CPU, Memori, Waktu Pemrosesan, dan Penggunaan Daya) untuk dua metode yang diterapkan pada wilayah “Kabupaten Blitar.” Pada metode pertama, yaitu Intersection, penggunaan CPU tercatat sebesar 59,4%, dengan memori sebesar 385,3 MB, dan waktu pemrosesan yang sangat lama, yaitu 06 jam 38 menit 27 detik. Metode ini juga menunjukkan penggunaan daya yang sangat tinggi. Sementara itu, pada metode kedua, yaitu Filtering, penggunaan CPU tercatat lebih rendah, yaitu 29,0%, dengan memori yang juga lebih kecil, yaitu 155,3 MB. Waktu pemrosesan untuk metode Filtering jauh lebih cepat, yaitu 03 menit 17 detik, meskipun penggunaan daya untuk kedua metode tersebut masih sangat tinggi. Hal ini menunjukkan bahwa meskipun metode Filtering lebih efisien dalam hal waktu, metode Intersection memerlukan lebih banyak sumber daya sistem, baik dari segi CPU maupun memori.

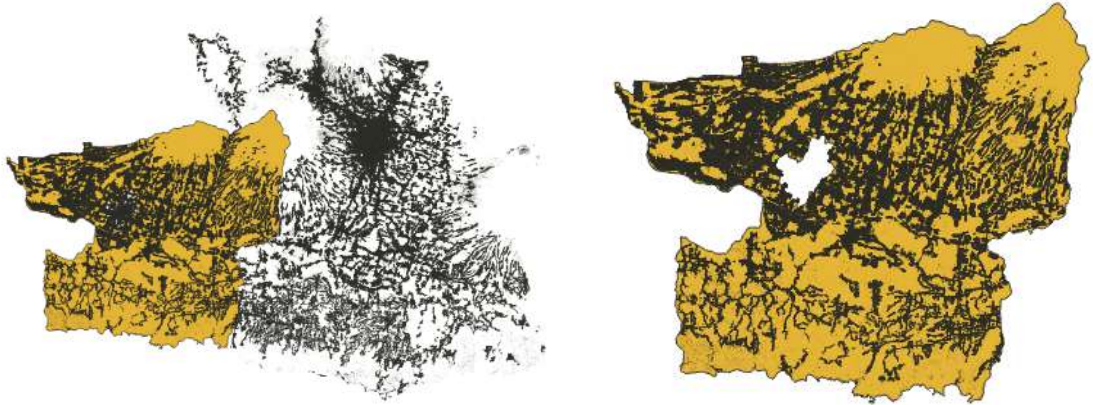
3.3 Skalabilitas

Dalam hal skalabilitas, metode Intersection dan Filtering menunjukkan perbedaan yang signifikan ketika diterapkan pada *dataset* yang lebih besar maupun lebih kecil. Pengukuran *resource* sistem menunjukkan bahwa metode Intersection cenderung memerlukan lebih banyak waktu pemrosesan dan resource dibandingkan dengan metode Filtering. Sebagai contoh, pada wilayah Kecamatan Panggungrejo, saya mengambil data bangunan dari *dataset* bangunan se-Kabupaten Blitar. Metode Intersection menggunakan 8,7% CPU dan 87,5 MB Memory dengan waktu pemrosesan 9 menit 33 detik. Di sisi lain, metode Filtering pada wilayah yang sama menggunakan *resource* yang sama, yaitu 8,7% CPU dan 87,5 MB Memory, namun waktu pemrosesannya jauh lebih cepat, hanya 2 detik.

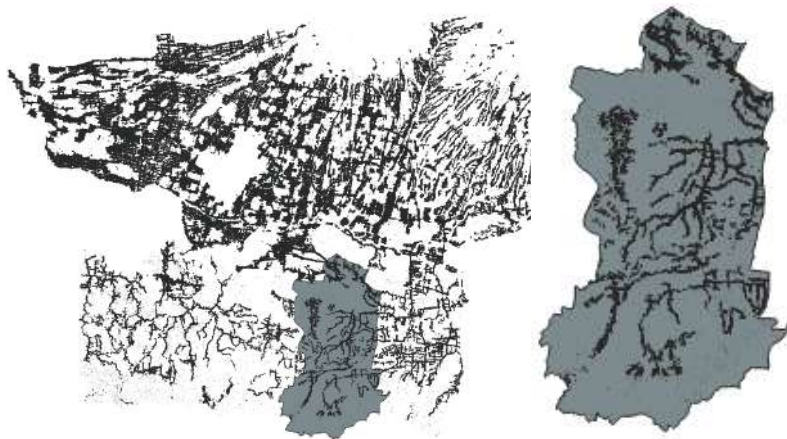
Ketika diterapkan pada *dataset* yang lebih besar, seperti data bangunan se-Kabupaten Blitar yang diproses dari *dataset* bangunan di wilayah Malang Raya, Blitar Raya, dan Kota Batu, perbedaan ini semakin mencolok. Metode Intersection memerlukan waktu 6 jam 38 menit 27 detik dengan penggunaan CPU sebesar 59,4% dan Memory 385,3 MB. Sementara itu, metode



Filtering hanya memerlukan waktu 3 menit 17 detik dengan penggunaan CPU sebesar 29% dan Memory 155,3 MB. Seperti yang dicantumkan pada Tabel 2. Dengan demikian, metode Filtering lebih unggul dalam hal skalabilitas, baik pada *dataset* kecil maupun besar, karena membutuhkan lebih sedikit waktu dan *resource* sistem untuk menyelesaikan pemrosesan. Selanjutnya untuk memperjelas bagaimana gambaran data bangunan yang telah melalui pemrosesan data dapat dilihat di Gambar 6 dan 7.



Gambar 6 Sebelum dan Sesudah Pemrosesan Data pada Kabupaten Blitar



Gambar 7 Sebelum dan Sesudah Pemrosesan Data pada Kecamatan Panggungrejo

3.4 Pembahasan

Hasil penelitian menunjukkan adanya perbedaan signifikan dalam durasi pemrosesan antara metode Filtering dan Intersection dalam analisis data geospasial permukaan bangunan di Kabupaten Blitar. Dari segi waktu pemrosesan, metode Filtering secara konsisten menunjukkan hasil yang jauh lebih cepat dibandingkan Intersection di enam kecamatan yang diuji. Berdasarkan tabel perbandingan, Filtering mampu menyelesaikan proses rata-rata hanya dalam 1–2 detik per kecamatan, sementara Intersection memerlukan waktu antara 6 hingga 9 menit. Sebagai contoh, di Kecamatan Wates, Filtering hanya membutuhkan 1 detik dibandingkan 6 menit 55 detik (415 detik) pada Intersection. Hasil serupa juga terlihat di Kecamatan Panggungrejo, di mana Filtering menyelesaikan proses dalam 2 detik, sementara Intersection memakan waktu 9 menit 33 detik (573 detik). Variasi waktu pemrosesan pada Intersection dipengaruhi oleh kompleksitas data di masing-masing kecamatan, sedangkan Filtering tetap konsisten terlepas dari jumlah dan kerumitan data. Efisiensi ini juga tercermin di kecamatan lain seperti Wonotiro, Doko, Gandusari,



dan Nglegok, di mana Filtering rata-rata hanya memerlukan 1–2 detik, jauh lebih singkat dibandingkan durasi pemrosesan Intersection yang berkisar antara 7 hingga 9 menit.

Selain keunggulan dalam kecepatan pemrosesan, metode Filtering juga lebih efisien dalam penggunaan sumber daya. Rata-rata penggunaan CPU untuk Filtering sebesar 18,85% dengan konsumsi memori 121,4 MB, sedangkan Intersection memerlukan CPU sebesar 34,05% dan memori 236,4 MB. Dalam skala yang lebih besar, Filtering menunjukkan skalabilitas yang lebih tinggi, mampu menangani pemrosesan data se-Kabupaten Blitar dalam waktu hanya 3 menit 17 detik, sementara Intersection memakan waktu 6 jam 38 menit 27 detik. Implikasi dari hasil ini sangat penting bagi pengolahan data geospasial dalam skala besar, terutama dalam aplikasi yang membutuhkan pemrosesan cepat dan efisien, seperti pemetaan wilayah bencana atau pengelolaan tata ruang kota. Pengurangan waktu dan sumber daya memungkinkan pengguna untuk menghemat biaya operasional serta meningkatkan responsivitas sistem dalam situasi yang memerlukan analisis real-time.

Penelitian oleh Bonny & Soudan (2015) memperkuat efektivitas teknik pemfilteran dalam pengolahan data berskala besar. Mereka membahas penggunaan teknik pemfilteran untuk mempercepat perbandingan urutan pada basis data besar dengan mengeliminasi data yang tidak relevan sebelum proses pencocokan dilakukan. Hasilnya, mereka berhasil mempercepat proses perbandingan urutan hingga 50% dibandingkan metode tradisional yang mengharuskan pencocokan dilakukan terhadap seluruh basis data. Pendekatan ini tidak hanya mengurangi waktu komputasi secara signifikan, tetapi juga mempertahankan akurasi hasil. Temuan tersebut sejalan dengan hasil penelitian ini yang menunjukkan bahwa metode Filtering tidak hanya unggul dalam pengolahan data geospasial tetapi juga konsisten dengan pendekatan efisien yang diadopsi dalam penelitian terdahulu. Hal ini menegaskan bahwa Filtering merupakan solusi yang lebih efektif untuk pengolahan data kompleks dalam berbagai konteks aplikasi, baik geospasial maupun basis data lainnya.

4. KESIMPULAN

Penelitian ini menunjukkan bahwa metode Filtering memiliki keunggulan signifikan dibandingkan metode Intersection dalam pemrosesan data geospasial, terutama dari segi kecepatan, efisiensi sumber daya, dan skalabilitas. Dari hasil pengujian, metode Filtering mampu menyelesaikan proses dalam rata-rata 1,6 detik, jauh lebih cepat dibandingkan metode Intersection yang membutuhkan 7 menit 50 detik. Selain itu, metode Filtering lebih hemat dalam penggunaan sumber daya dengan rata-rata 18,85% CPU dan 121,4 MB memori, sedangkan metode Intersection memerlukan 34,05% CPU dan 236,4 MB memori, serta menunjukkan konsumsi daya yang tinggi secara konsisten. Pada *dataset* yang lebih besar, Filtering tetap unggul, seperti dalam pemrosesan data bangunan se-Kabupaten Blitar, di mana metode ini hanya memerlukan 3 menit 17 detik, dibandingkan dengan 6 jam 38 menit 27 detik yang dibutuhkan oleh Intersection. Berdasarkan hasil tersebut, dapat disimpulkan bahwa metode Filtering lebih unggul untuk pemrosesan data geospasial yang membutuhkan efisiensi waktu, optimalisasi sumber daya, dan skalabilitas yang lebih baik. Sebaliknya, metode Intersection kurang ideal dalam situasi di mana kecepatan dan efisiensi sistem menjadi faktor utama. Kemudian untuk penelitian selanjutnya, disarankan untuk mengeksplorasi penerapan metode Filtering dalam analisis geospasial berbasis machine learning guna meningkatkan efisiensi klasifikasi dan pemetaan data GIS.

DAFTAR PUSTAKA

- Afnarius, S., Syukur, M., Ekaputra, E. G., Parawita, Y., & Darman, R. (2020). Development of GIS for Buildings in the Customary Village of Minangkabau Koto Gadang, West Sumatra, Indonesia. *ISPRS International Journal of Geo-Information*, 9(6), Article ID: 365. <https://doi.org/10.3390/ijgi9060365>
- Auradkar, P. K., Raykar, A., Agarwal, I., Sitaram, D., & R., M. (2022). Accuracy Assessment and Performance Analysis of Raster to Vector Conversions on LULC Data – India. *Journal of Engineering, Design and Technology*, 20(6), 1787–1809. <https://doi.org/10.1108/JEDT-04-2021-0224>



- Bonny, T., & Soudan, B. (2015). Filtering Technique for High Speed Database Sequence Comparison. *Proceedings of the 2015 IEEE 9th International Conference on Semantic Computing (IEEE ICSC 2015)*, 73–76. <https://doi.org/10.1109/ICOSC.2015.7050781>
- Brovelli, M. A., Wu, H., Minghini, M., Molinari, M. E., Kilsedar, C. E., Zheng, X., Shu, P., & Chen, J. (2018). Open Source Software and Open Educational Material on Land Cover Maps Intercomparison and Validation. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XLII-4, 61–68. <https://doi.org/10.5194/isprs-archives-XLII-4-61-2018>
- Darmawan, S., Nurulhakim, N. N., & Hernawati, R. (2024). Kecerdasan Buatan Berbasis Geospasial (GeoAI) Menggunakan Google Earth Engine untuk Monitoring Fenomena Urban Heat Island di Indonesia. *ELKOMIKA: Jurnal Teknik Energi Elektrik, Teknik Telekomunikasi, & Teknik Elektronika*, 12(2), 303–320. <https://doi.org/10.26760/elkomika.v12i2.303>
- Duarte, L., Queirós, C., & Teodoro, A. C. (2021). Comparative Analysis of QGIS Plugins for Web Maps Creation. *La Granja*, 34(2), 8–26. <https://doi.org/10.17163/lgr.n34.2021.01>
- El-Ashmawy, Dr. K. L. A. (2019). Three-Dimensional Intersection Method for Monitoring and Analysis of Horizontal and Vertical Movements of Buildings. *International Journal of Advances in Scientific Research and Engineering*, 05(10), 162–170. <https://doi.org/10.31695/IJASRE.2019.33553>
- El-Kareem, M. M. A., & El-Emary, I. M. M. (2019). Comparative Study Between the Algorithm Using Frequency Domain and the Algorithm Using Linear Programming. *Indian Journal of Science and Technology*, 12(18), 1–5. <https://doi.org/10.17485/ijst/2019/v12i18/144589>
- F, S. N., & Saleh, S. R. (2023). Penggunaan Model SIG dalam Analisis Fisik Lingkungan di Kota Metro. *Jurnal Perencanaan Wilayah dan Kota*, 17(2), 42–54. <https://doi.org/10.29313/jpwk.v17i2.346>
- Flenniken, J. M., Stuglik, S., & Iannone, B. V. (2020). Quantum GIS (QGIS): An Introduction to a Free Alternative to More Costly GIS Platforms. *EDIS*, 2020(2), Article ID: 7. <https://doi.org/10.32473/edis-fr428-2020>
- Fram, C., Chistopoulou, K., & Ellul, C. (2015). Assessing the Quality of OpenStreetMap Building Data and Searching for a Proxy Variable to Estimate OSM Building Data Completeness. In N. Malleon, N. Addis, H. Durham, A. Heppenstall, R. Lovelace, P. Norman, & R. Oldroyd (Eds.), *GIS Research UK 2015 (GISRUK2015)* (pp. 195–205). GISRUK 2015 Proceedings. <https://api.semanticscholar.org/CorpusID:196089344>
- Goodman, S., BenYishay, A., Lv, Z., & Runfola, D. (2019). GeoQuery: Integrating HPC Systems and Public Web-Based Geospatial Data Tools. *Computers and Geosciences*, 122, 103–112. <https://doi.org/10.1016/j.cageo.2018.10.009>
- Ivanov, S. (2020). Spatial Data Models. *Journal Scientific and Applied Research*, 20(1), 40–46. <https://doi.org/10.46687/jsar.v20i1.303>
- Khajedehi, M. H., Prataviera, E., Bordignon, S., & De Carli, M. (2024). Exploiting Geographic Open Data to Improve Urban Building Energy Simulations: The Padova City Center Case Study. *E3S Web of Conferences*, 523, Article ID: 05007. <https://doi.org/10.1051/e3sconf/202452305007>
- Kukulska, A., Salata, T., Cegielska, K., & Szylar, M. (2018). Methodology of Evaluation and Correction of Geometric Data Topology in QGIS Software. *Acta Scientiarum Polonorum. Formatio Circumiectus*, 17(1), 125–138. <https://doi.org/10.15576/ASP.FC/2018.17.1.125>
- Lochana M, Manvith P, & Anusha U A. (2024). A Literature Review of Exploratory Analysis of Geolocal Data. *International Journal of Advanced Research in Science, Communication and Technology*, 593–597. <https://doi.org/10.48175/IJARSCT-15379>
- Memduhoglu, A., & Basaraner, M. (2022). An Approach for Multi-Scale Urban Building Data Integration and Enrichment Through Geometric Matching and Semantic Web. *Cartography and Geographic Information Science*, 49(1), 1–17. <https://doi.org/10.1080/15230406.2021.1952108>
- Rajab, M., & Wang, D. (2020). Practical Challenges and Recommendations of Filter Methods for Feature Selection. *Journal of Information & Knowledge Management*, 19(01), Article ID: 2040019. <https://doi.org/10.1142/S0219649220400195>



- Raju, P. L. N. (2004). Spatial Data Analysis. In M. V. K. Sivakumar, P. S. Roy, K. Harmsen, & S. K. Saha (Eds.), *Proceedings of the Training Workshop* (pp. 151–174). World Meteorological Organisation. <https://api.semanticscholar.org/CorpusID:14763553>
- Singla, S., Eldawy, A., Diao, T., Mukhopadhyay, A., & Scudiero, E. (2021). Experimental Study of Big Raster and Vector Database Systems. *2021 IEEE 37th International Conference on Data Engineering (ICDE)*, 2243–2248. <https://doi.org/10.1109/ICDE51399.2021.00231>
- Sutanto, A. I., & Aditya, T. P. (2021). Pendekatan Otomatisasi Evaluasi Kualitas Kelengkapan pada Informasi Geospasial. *Jurnal Geosaintek*, 7, 27–36. <https://journal.its.ac.id/index.php/geosaintek/article/view/7387>
- Wahab, L., & Kurniawan, A. (2023). Pemanfaatan Sistem Informasi Geografis untuk Pemetaan Lahan Pertanian di Kecamatan Kembaran, Banyumas, Jawa Tengah. *Jurnal Agroindustri Terapan Indonesia*, 1(1), 1–10. <https://doi.org/10.31962/jati.v1i1.119>
- Wegmann, M., Schwalb-Willmann, J., & Dech, S. (2020). *An Introduction to Spatial Data Analysis*. Pelagic Publishing. <https://doi.org/10.53061/HCED6492>
- Wismarini, T. D., & Khristianto, T. (2016). Implementasi Superimpose dalam Pemodelan Spasial Tingkat Rentan Banjir di Semarang. *Jurnal Teknologi Informasi DINAMIK*, 21(2), 124–138. <https://doi.org/10.35315/dinamik.v21i2.6092>
- Yogi, K. S., V, D. G., K M, M., Sujithra, L. R., Prasad, K., & Midhun, P. (2024). Scalability and Performance Evaluation of Machine Learning Techniques in High-Volume Social Media Data Analysis. *2024 11th International Conference on Reliability, Infocom Technologies and Optimization (Trends and Future Directions) (ICRITO)*, 1–6. <https://doi.org/10.1109/ICRITO61523.2024.10522361>
- Zhou, Q. (2018). Exploring the Relationship Between Density and Completeness of Urban Building Data in OpenStreetMap for Quality Estimation. *International Journal of Geographical Information Science*, 32(2), 257–281. <https://doi.org/10.1080/13658816.2017.1395883>



LAMPIRAN A

Tabel 3 Atribut Data

No.	Latitude	Longitude	Area in me	Confidence	Full plus	ID KAB
1	-8.003240350000000	112.015856420000006	29.569900000000001	0.7849	6P3JX2W8+P857	3505
2	-8.087264050000000	112.066084869999997	316.466099999999983	0.8518	6P3JW378+3CWG	3505
3	-8.070966739999999	112.311355509999999	142.856699999999989	0.7595	6P3JW8H6+JG97	3505
...	...	...	...	...	...	...
1032658	-8.058731470000000	112.311500159999994	138.684500000000014	0.8187	6P3JW8R6+GJ2J	3505



Analisis Sistem Deteksi Citra untuk Optimalisasi Pengawasan Lalu Lintas Udara Menggunakan Algoritma YOLOv5

Astika Ayuningtyas ^{(1)*}, Imam Riadi ⁽²⁾, Anton Yudhana ⁽³⁾

¹ Departemen Informatika, Institut Teknologi Dirgantara Adisutjipto, Yogyakarta, Indonesia

² Departemen Sistem Informasi, Universitas Ahmad Dahlan, Yogyakarta, Indonesia

³ Departemen Teknik Elektro, Universitas Ahmad Dahlan, Yogyakarta, Indonesia

e-mail : astika@stta.ac.id, imam.riadi@is.uad.ac.id, eyudhana@ee.uad.ac.id.

* Penulis korespondensi.

Artikel ini diajukan 31 Oktober 2024, direvisi 21 Januari 2025, diterima 23 Januari 2025, dan dipublikasikan 30 September 2025.

Abstract

This study aims to develop an image detection system capable of identifying manned and unmanned aircraft objects to support air traffic surveillance. The increasing flight activity, both from commercial aircraft and drones, requires a more optimal surveillance system to connect the airspace efficiently. In this study, a Convolutional Neural Network (CNN) model utilizing the You Only Look Once version 5 (YOLOv5) method is employed to detect and classify objects in real-time from aircraft images. The methodology employed includes collecting aerial image data, labeling the data, and training object detection models using YOLOv5. The dataset used consists of 2,520 images of manned aircraft (warplanes) and 5,422 images of unmanned aircraft (drones). The experimental results demonstrate that the YOLOv5 model achieves high detection accuracy for both manned and unmanned aircraft, with a relatively fast inference time, thereby supporting the development of an effective air traffic surveillance system. This system is expected to be an integral part of a more sophisticated and responsive air traffic surveillance solution.

Keywords: Convolutional Neural Network, YOLOv5, Deep Learning, Aircraft Detection, Air Traffic Surveillance

Abstrak

Penelitian ini bertujuan untuk mengembangkan sistem deteksi citra yang mampu mengidentifikasi objek pesawat berawak dan tidak berawak guna mendukung pengawasan lalu lintas udara. Meningkatnya aktivitas penerbangan, baik dari pesawat komersial maupun drone, diperlukan sistem pengawasan yang lebih optimal untuk menghubungkan ruang udara secara efisien. Dalam penelitian ini, menggunakan model Convolutional Neural Network (CNN) dengan metode You Only Look Once versi 5 (YOLOv5) diterapkan untuk mendeteksi dan mengklasifikasikan objek dari citra pesawat secara *real-time*. Metodologi yang digunakan meliputi pengumpulan data citra udara, pelabelan data, dan pelatihan model deteksi objek dengan YOLOv5. Dataset yang digunakan sebesar 2.520 gambar dengan kelas pesawat berawak (*warplane*) dan 5.422 gambar dengan kelas pesawat tidak berawak (*drone*). Hasil eksperimen menunjukkan bahwa model YOLOv5 memiliki akurasi deteksi yang tinggi terhadap pesawat berawak dan tidak berawak, dengan waktu inferensi yang relatif cepat, sehingga dapat mendukung pengembangan sistem pengawasan lalu lintas udara secara efektif. Sistem ini diharapkan dapat menjadi bagian integral dari solusi pengawasan lalu lintas udara yang lebih canggih dan responsif.

Kata Kunci: Convolutional Neural Network, YOLOv5, Deep Learning, Deteksi Pesawat, Pengawasan Lalu Lintas Udara

1. PENDAHULUAN

Industri penerbangan mengalami perkembangan pesat dalam beberapa dekade terakhir. Menurut data Internasional Air transport Association (IATA), jumlah penumpang udara diperkirakan akan terus tumbuh sebesar 3,7 % per tahun hingga tahun 2037, mencapai 8,2 miliar penumpang (Wittmer & Müller, 2021). Pengawasan lalu lintas udara adalah salah satu aspek yang sangat penting dalam industri penerbangan untuk menjamin keselamatan dan efisiensi operasional. Dalam beberapa dekade terakhir, peningkatan jumlah penerbangan komersial dan



Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

penggunaan drone secara signifikan telah meningkatkan kompleksitas pengelolaan ruang udara. Otoritas penerbangan global menghadapi tantangan besar dalam memastikan bahwa seluruh objek udara, termasuk pesawat terbang dan drone, dapat dideteksi dan dipantau secara efektif untuk menghindari tabrakan, memastikan kelancaran alur lalu lintas udara, dan merespons ancaman keamanan potensial. Teknologi pengawasan yang ada saat ini, seperti radar, memiliki keterbatasan, baik dalam hal jangkauan, akurasi, maupun efisiensi dalam mendeteksi pesawat, terutama pada lingkungan yang kompleks dan penuh tantangan, seperti kondisi cuaca buruk atau wilayah udara yang padat.

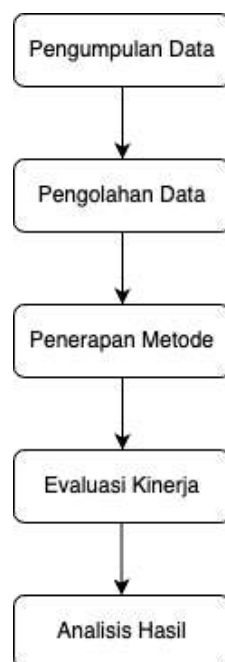
Seiring dengan kemajuan pesat dalam bidang kecerdasan buatan (Artificial Intelligence atau AI) dan pembelajaran mendalam (*deep learning*), teknologi ini memberikan peluang baru untuk meningkatkan kemampuan deteksi objek (Akbar et al., 2022; Anne & Gueye, 2024; Argho et al., 2024; Arnia et al., 2024; Arooj et al., 2024; Asadi et al., 2024; Chowdhury et al., 2023; D L et al., 2024; Das & Baruah, 2023; Ellouze et al., 2024; Fahmi et al., 2023; Freitas et al., 2024; Giménez-Gallego et al., 2024; Inamullah et al., 2024; H. Liu et al., 2024; Oyedele et al., 2024; Peryanto et al., 2020a, 2020b; Ramakrishnan et al., 2024; Reddy et al., 2024; S et al., 2024; D. K. Saha et al., 2024; Saifullah et al., 2023; Sharma et al., 2024; Sileo et al., 2024; Sotoude et al., 2023; Tan & Liu, 2023; Ye et al., 2023; Yi et al., 2024; Zhang et al., 2024), termasuk pesawat. Salah satu algoritma deteksi objek yang sangat efisien dan telah banyak digunakan adalah You Only Look Once version 5 (YOLOv5). YOLOv5 merupakan salah satu model deteksi objek berbasis *deep learning* yang terkenal karena kemampuannya untuk melakukan deteksi objek secara *real-time* dengan kecepatan dan akurasi yang sangat tinggi (Aydin & Singha, 2023; Goud et al., 2023; Huang et al., 2023; Jung & Choi, 2022; Lian-suo et al., 2024; Murthy et al., 2022; Sapakova et al., 2024; Wang et al., 2023). Algoritma ini bekerja dengan cara membagi gambar menjadi grid dan memprediksi objek dalam setiap grid secara simultan. YOLOv5 sangat ideal untuk aplikasi pengawasan lalu lintas udara karena kemampuan mendeteksi objek dengan cepat (Ren et al., 2023; Xu et al., 2024), bahkan dalam kondisi lingkungan yang kompleks (Fransisca & Santoso, 2023; Karimah et al., 2024; Rishi et al., 2024). Dalam konteks deteksi pesawat, algoritma ini memungkinkan sistem untuk mendeteksi pesawat dalam berbagai kondisi, seperti jarak, sudut pandang, serta perubahan ukuran dan posisi. Hal ini menawarkan keunggulan signifikan dalam hal; 1) Kecepatan Inferensi: YOLOv5 dirancang untuk melakukan deteksi dalam waktu yang sangat singkat, menjadikannya pilihan yang sangat tepat untuk aplikasi yang membutuhkan respons cepat, seperti pengawasan *real-time*; 2) Akurasi Deteksi: Dengan kemampuan deteksi multi-objek (Ahmed et al., 2024; Qiu et al., 2024); 3) Efisiensi Komputasi: Algoritma ini dioptimalkan untuk dapat berjalan pada perangkat keras dengan spesifikasi yang bervariasi, mulai dari sistem berbasis GPU hingga perangkat kecil seperti *edge device* (Feng et al., 2024; G. Liu et al., 2023), yang menjadikannya fleksibel untuk diterapkan dalam sistem pengawasan lalu lintas udara di berbagai lokasi.

Penerapan YOLOv5 dalam sistem deteksi pesawat menawarkan solusi modern untuk mengatasi tantangan pengawasan lalu lintas udara, seperti deteksi objek kecil (drone), peningkatan jumlah pesawat, serta kebutuhan untuk merespons situasi kritis secara cepat dan tepat. YOLOv5 memungkinkan sistem untuk memantau lalu lintas udara secara lebih efektif dibandingkan dengan teknologi konvensional. Hal ini diharapkan dapat mengurangi risiko kecelakaan udara, meningkatkan efisiensi manajemen ruang udara, dan memperkuat sistem keamanan udara secara keseluruhan. Dalam penelitian ini, menganalisis bagaimana implementasi algoritma YOLOv5 dapat mengoptimalkan sistem deteksi pesawat. Penelitian ini diharapkan dapat memberikan kontribusi signifikan dalam meningkatkan efisiensi dan keamanan pengawasan lalu lintas udara global. Selain itu, ke depannya dapat mendukung inisiatif modernisasi pengelolaan lalu lintas udara global, yang diharapkan mampu meningkatkan efisiensi, keamanan, dan kelancaran penerbangan di masa depan. Penelitian dan pengembangan sistem deteksi pesawat berbasis YOLOv5 ini diharapkan dapat memberikan kontribusi signifikan dalam menciptakan solusi pengawasan lalu lintas udara yang lebih cerdas dan efektif, serta mengantisipasi tantangan di masa depan dalam industri penerbangan.



2. METODE PENELITIAN

Penelitian ini bertujuan untuk menganalisis kinerja algoritma YOLOv5 dalam sistem deteksi pesawat. Metodologi penelitian ini mencakup beberapa tahap utama, yaitu pengumpulan data, pengolahan data, penerapan metode, evaluasi kinerja, dan analisis hasil. Gambar 1 menampilkan tahapan-tahapan metodologi yang dilakukan dalam penelitian. Pada tahap pertama dilakukan pengumpulan data citra berupa pesawat berawak dan tidak berawak yang diambil dari sumber data primer yaitu hasil eksperimen peneliti dengan menggunakan bantuan *drone* untuk menangkap objek berupa model pesawat berawak dan tidak berawak (*drone*). Data primer diambil menggunakan Drone DJI mini 4 Pro dengan *Effective Pixels*: 48 MP, pengambilan dilakukan pada Bulan Oktober 2024. Selain data primer, jenis data lainnya yang digunakan dalam penelitian ini adalah jenis sekunder yang didapatkan dari dataset *public* yaitu Kaggle. Data sekunder dikumpulkan dari Bulan September sampai dengan Oktober 2024. Data sekunder digunakan untuk mempersiapkan dataset untuk proses pelatihan.



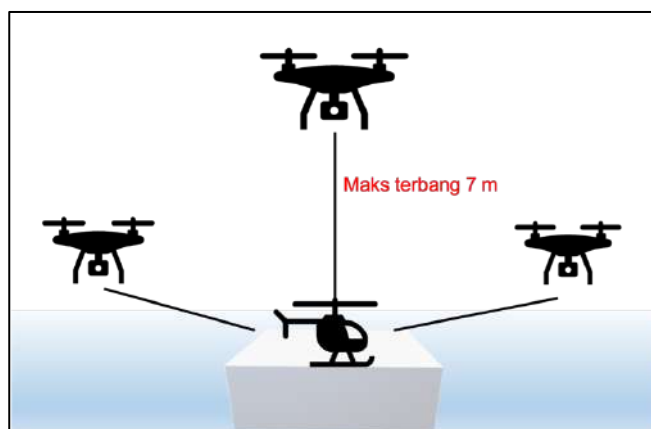
Gambar 1 Tahapan Metodologi Penelitian

Penggunaan jenis data yang berbeda ini memperkaya data pelatihan dan pengujian dalam proses deteksi, sehingga diharapkan mendapat akurasi yang semakin tinggi. Proses pengumpulan data yang dilakukan pertama adalah menyiapkan *drone* dan model pesawat yang akan digunakan untuk simulasi objek yang akan dideteksi. Pengambilan data citra ini dilakukan di Lapangan Utama Institut Teknologi Dirgantara Adisutjipto (ITDA) dengan batasan ketinggian pengambilan obyek adalah 7 (tujuh) meter. Jarak ketinggian ini masih dalam kategori yang diijinkan untuk terbang di kawasan Lanud Adisutjipto. Kondisi pengambilan objek dilakukan pada pagi, siang dan malam hari. Hal tersebut dimaksudkan untuk mendapatkan dataset yang lebih bervariasi dengan kondisi pencahayaan yang beragam dalam jumlah yang besar. Salah satu contoh detail informasi pengambilan citra ditampilkan pada Gambar 2.



Gambar 2 Menunjukkan Ilustrasi Pengambilan Citra Pesawat Terbang





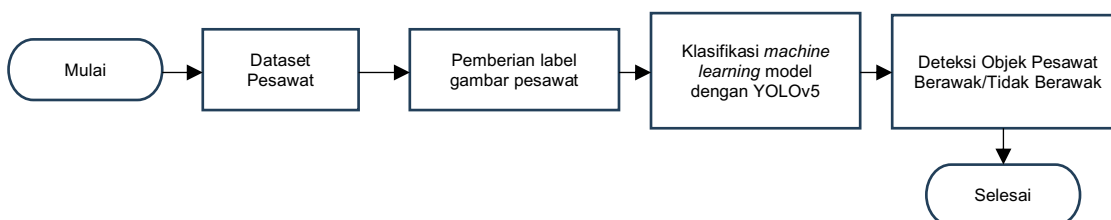
Gambar 3 Ilustrasi Pengambilan Citra Pesawat

Gambar 3 menunjukkan ilustrasi pengambilan citra pesawat terbang. Data yang dikumpulkan berupa gambar dan video objek pesawat. Penelitian ini berfokus pada keberhasilan metode YOLOv5 untuk deteksi objek pesawat. Oleh karena itu, penelitian ini mempertimbangkan kualitas gambar yang dihasilkan oleh kedua jenis dataset yang digunakan. Dataset yang dikumpulkan dari sumber data sekunder untuk kelas pesawat berawak sebanyak 2.520 gambar yang terdiri 2205 *train*, 210 *valid*, dan 105 *test*. Dataset untuk kelas pesawat tidak berawak untuk sumber data primer dan sekunder jumlah keseluruhannya adalah 5.422 gambar dengan 4.717 *train*, 470 *valid*, dan 235 *test*; secara lengkapnya tersaji pada Tabel 1.

Tabel 1 Kelas Dataset

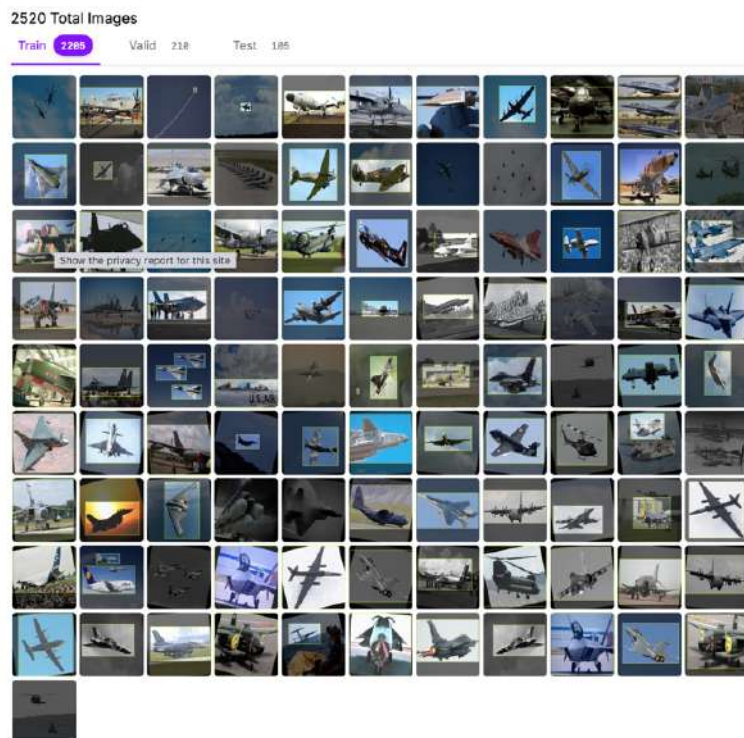
No.	Sumber Data	Kelas	Train	Valid	Test	Definisi Kelas
1	Primer	Pesawat Tidak Berawak (drone)	2593	267	134	Citra Drone
2	Sekunder	Pesawat Berawak	2205	210	105	Citra Pesawat Tempur/Militer
3	Sekunder	Pesawat Tidak Berawak (drone)	2124	203	101	Citra Drone

Selanjutnya pada pengumpulan citra berupa gambar yang didapatkan dari sumber pribadi dan *public* dilakukan pengolahan yaitu diberikan pelabelan untuk memuat kelas objeknya. Proses pelabelan dilakukan dengan cara membuat nama kelas dan *bounding box* pada setiap objek gambar. Tahapan selanjutnya, dilakukan proses *pre-processing* dan *augmentations* untuk meningkatkan akurasi metode YOLOv5 dalam pendeteksian. Secara lengkap ditampilkan pada Gambar 4 yang dijelaskan melalui blok diagram dengan poin berupa *input*, proses dan *output* (klasifikasi dengan YOLOv5). Arsitektur YOLOv5 akan diproses berdasarkan dataset yang sebelumnya telah dimasukkan ke arsitektur tersebut dengan sebelumnya melewati proses *training*. Hasil deteksi berupa pesawat berawak (misalnya: *warplane*) dan pesawat tidak berawak (*drone*). Dataset yang siap dilatih merupakan dataset yang sudah melalui tahap sebelumnya dan sudah diberikan label pada setiap citra.

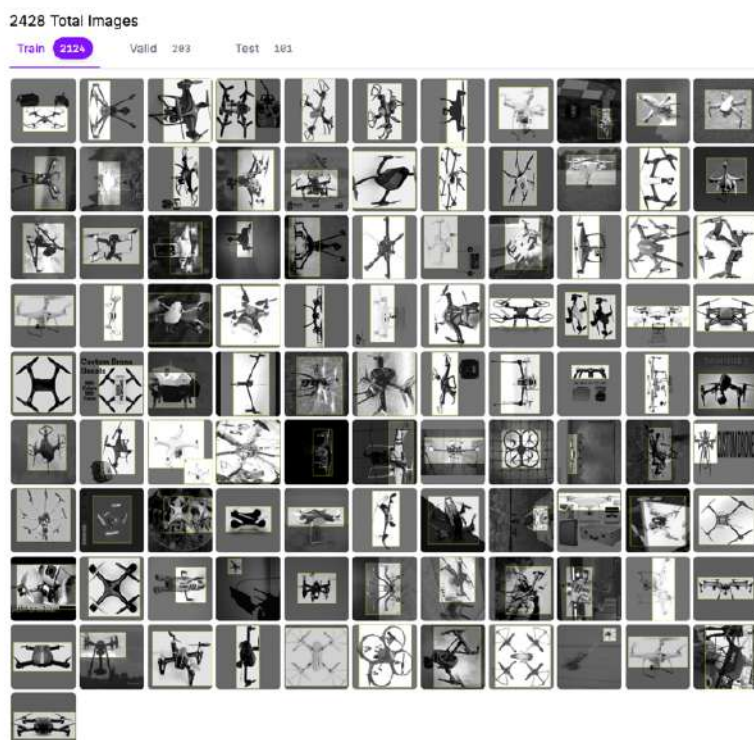


Gambar 4 Blok Diagram Proses Deteksi





Gambar 5 Proses Pembuatan Dataset Pesawat Berawak



Gambar 6 Proses Pembuatan Dataset Pesawat Tidak Berawak

Gambar 5 adalah proses pembuatan dataset dari jenis objek yang akan dideteksi yaitu pesawat berawak yang didapatkan dari dataset *public* yaitu Kaggle (Adam, 2022; Bilel, 2022), cuplikan yang ditampilkan contoh untuk proses anotasi pada pesawat berawak jenis *warplane*. Dataset



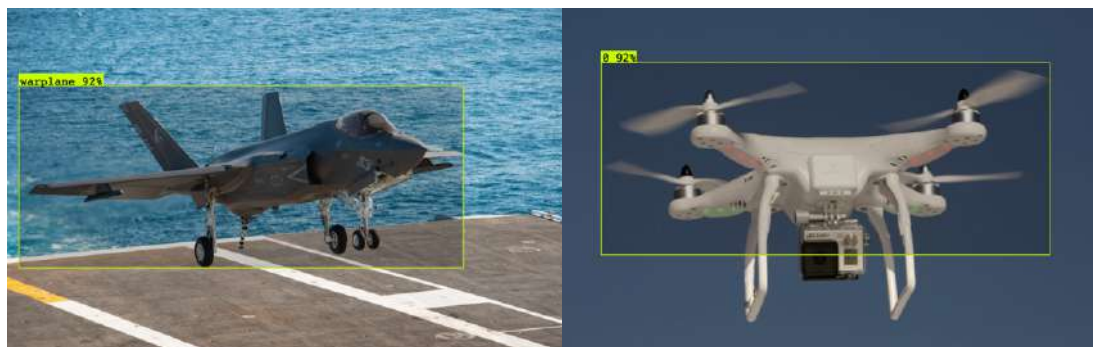
Artikel ini didistribusikan mengikuti lisensi Atribusi-NonKomersial CC BY-NC sebagaimana tercantum pada <https://creativecommons.org/licenses/by-nc/4.0/>.

berupa citra *drone* yang digunakan berjumlah 2520 (dua ribu lima ratus dua puluh). Pada Gambar 6 merupakan dataset dari jenis objek yang akan dideteksi yaitu pesawat tidak berawak yang didapatkan dari dataset *public* yaitu Kaggle (Aksoy et al., 2019; B. Saha, 2022). Dataset berupa citra *drone* yang digunakan berjumlah 2428 (dua ribu empat ratus dua puluh delapan), dimana akan dibagi 3 (tiga) kelompok yaitu data *train*, *valid* dan *test*. Saat proses anotasi menggunakan *platform* Roboflow sehingga pengumpulan data menjadi lebih mudah saat *pre-processing*. Proses anotasi adalah tahapan untuk pemberian label pada setiap objek gambar (Saifullah et al., 2023; Sunardi et al., 2017; Surya et al., 2019; Yudhana et al., 2017). Selanjutnya setiap gambar yang telah dianotasi dapat dikonversi ukurannya atau menambahkan variasi fitur pada gambar untuk meningkatkan akurasi YOLOv5 saat proses deteksi.

Pada tahap ketiga adalah penerapan Metode YOLOv5 dengan melakukan proses pelatihan model menggunakan data yang telah dipersiapkan sebelumnya. Setelah proses pelatihan selesai, langkah selanjutnya adalah melakukan pengujian untuk menilai kemampuan model dalam mengenali objek sesuai dengan tujuan penelitian. Berdasarkan hasil pengujian tersebut, dilakukan evaluasi kinerja untuk mengukur performa model melalui berbagai metrik yang relevan, sehingga dapat diperoleh gambaran mengenai efektivitas metode yang digunakan. Tahapan akhir dari proses ini adalah menarik kesimpulan berdasarkan hasil evaluasi, yang kemudian menjadi dasar untuk memberikan rekomendasi atau pengembangan lebih lanjut.

3. HASIL DAN PEMBAHASAN

Penelitian menggunakan *platform* Roboflow dan menerapkan metode YOLOv5. Tujuan menggunakan metode ini untuk menentukan tempat sebuah citra pada objek yang hadir dan mengelompokkan berdasarkan jenis objeknya yaitu pesawat berawak (misalnya: *warplane*) atau pesawat tidak berawak (misalnya: *drone*). Gambar 7 menunjukkan dua contoh hasil deteksi pesawat berawak jenis *warplane* dan *drone* dengan akurasi 92%.



Gambar 7 Contoh Cuplikan Hasil Deteksi Objek

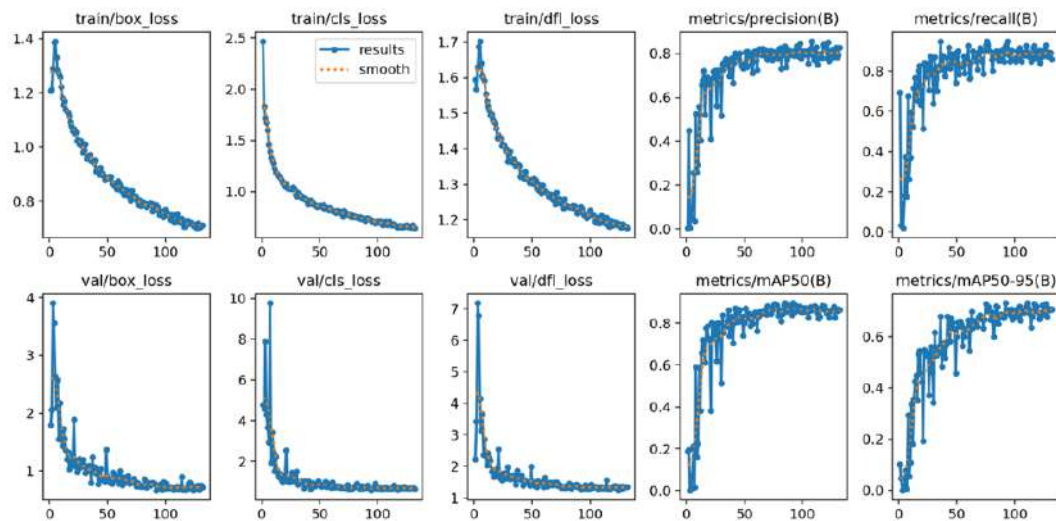
Pada Tabel 2 menunjukkan contoh pengaplikasian Metode YOLOv5 yang telah diterapkan dalam penelitian ini. Diambil sampel cuplikan 2% dari total 4% data *testing* (yang berjumlah 105 data). Hasil secara keseluruhan untuk deteksi objek pesawat tidak berawak pada set validasi memiliki tingkat presisi untuk model pesawat berawak (jenis *warplane*) yaitu 89,6% dengan nilai *precision* 80,4%, ini berarti hanya 19,6% prediksi yang salah (*false positive*), sedangkan untuk model pesawat tidak berawak (*drone*) yaitu 93,5% dengan nilai *precision* 93,2%, artinya hampir semua prediksi tersebut benar, dengan 6,8% kesalahan. Dalam penelitian, akurasi merupakan faktor penting dalam mengevaluasi kinerja suatu sistem deteksi (Akbar et al., 2022; Fahmi et al., 2023; Muis et al., 2023; Sulistyio et al., 2018; Sunardi et al., 2022; Tjahyanto & Atletiko, 2022; Umar et al., 2018). Pada penelitian ini diperoleh hasil bahwa model *warplane* dengan akurasi lebih rendah mampu memberikan nilai kebenaran deteksi lebih tinggi dibandingkan model *drone* yang memiliki akurasi 93,5%. Hal ini dapat disebabkan oleh satu faktor yaitu ketidakseimbangan data (Molnar, 2020; Pramanik et al., 2021; Schenk & Kern, 2024). Faktor lain yang mempengaruhi akurasi pada deteksi objek adalah kompleksitas latar belakang dan variasi pencahayaan seperti pada penelitian (Khairunisa et al., 2024; Pratama et al., 2024; Wahyudi et al., 2024).



Tabel 2 Cuplikan Hasil Deteksi Objek Pesawat

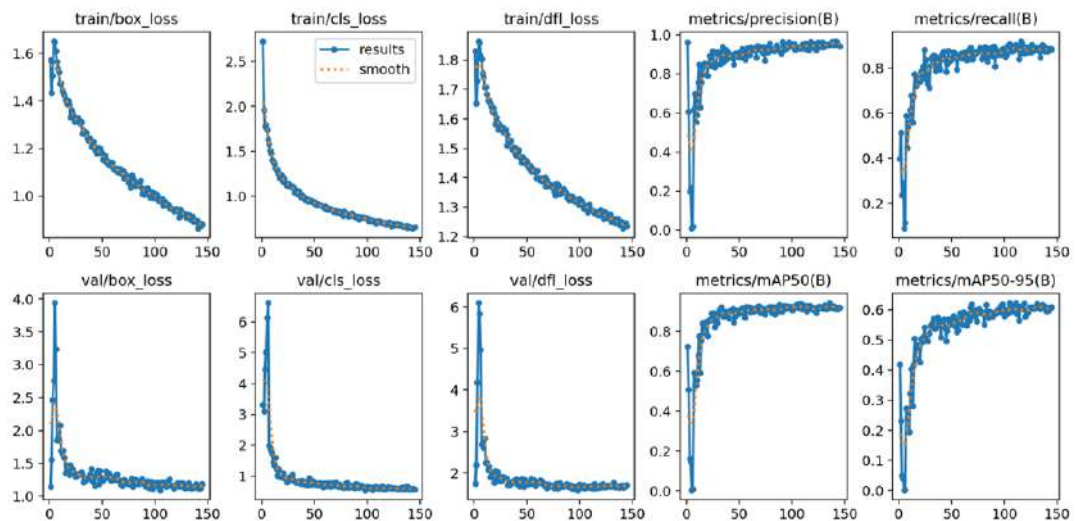
Kode Citra	Kelas	Definisi Kelas	Hasil Deteksi	Akurasi
10007	Pesawat Berawak	Warplane	Terdeteksi	91%
10017	Pesawat Berawak	Warplane	Terdeteksi	92%
10033	Pesawat Berawak	Warplane	Terdeteksi	88%
10043	Pesawat Berawak	Warplane	Terdeteksi	53%
10046	Pesawat Berawak	Warplane	Terdeteksi	94%
10055	Pesawat Berawak	Warplane	Terdeteksi	70%
10063	Pesawat Berawak	Warplane	Tidak Terdeteksi	0%
10076	Pesawat Berawak	Warplane	Terdeteksi	82%
10087	Pesawat Berawak	Warplane	Terdeteksi	88%
10097	Pesawat Berawak	Warplane	Terdeteksi	92%
10104	Pesawat Berawak	Warplane	Terdeteksi	85%
10114	Pesawat Berawak	Warplane	Terdeteksi	93%
10127	Pesawat Berawak	Warplane	Terdeteksi	85%
10132	Pesawat Berawak	Warplane	Terdeteksi	66%
10142	Pesawat Berawak	Warplane	Terdeteksi	90%
10146	Pesawat Berawak	Warplane	Terdeteksi	93%
10161	Pesawat Berawak	Warplane	Terdeteksi	87%
10254	Pesawat Berawak	Warplane	Terdeteksi	76%
10166	Pesawat Berawak	Warplane	Terdeteksi	74%
10353	Pesawat Berawak	Warplane	Terdeteksi	75%
10472	Pesawat Berawak	Warplane	Terdeteksi	83%
10001	Pesawat Tidak Berawak	Drone	Terdeteksi	88%
10002	Pesawat Tidak Berawak	Drone	Terdeteksi	68%
10003	Pesawat Tidak Berawak	Drone	Terdeteksi	73%
10004	Pesawat Tidak Berawak	Drone	Terdeteksi	92%
10005	Pesawat Tidak Berawak	Drone	Terdeteksi	82%
10006	Pesawat Tidak Berawak	Drone	Terdeteksi	83%
10007	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10008	Pesawat Tidak Berawak	Drone	Terdeteksi	72%
10009	Pesawat Tidak Berawak	Drone	Terdeteksi	75%
10010	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10011	Pesawat Tidak Berawak	Drone	Terdeteksi	74%
10012	Pesawat Tidak Berawak	Drone	Terdeteksi	82%
10013	Pesawat Tidak Berawak	Drone	Terdeteksi	75%
10014	Pesawat Tidak Berawak	Drone	Terdeteksi	86%
10015	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10016	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10017	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10018	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10019	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10020	Pesawat Tidak Berawak	Drone	Terdeteksi	69%
10021	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10022	Pesawat Tidak Berawak	Drone	Terdeteksi	87%
10023	Pesawat Tidak Berawak	Drone	Terdeteksi	83%
10024	Pesawat Tidak Berawak	Drone	Terdeteksi	69%
10025	Pesawat Tidak Berawak	Drone	Terdeteksi	90%
10026	Pesawat Tidak Berawak	Drone	Terdeteksi	86%
10254	Pesawat Berawak	Warplane	Terdeteksi	76%
10166	Pesawat Berawak	Warplane	Terdeteksi	74%
10353	Pesawat Berawak	Warplane	Terdeteksi	75%
10472	Pesawat Berawak	Warplane	Terdeteksi	83%
10027	Pesawat Tidak Berawak	Drone	Terdeteksi	75%
10028	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%
10029	Pesawat Tidak Berawak	Drone	Terdeteksi	72%
10030	Pesawat Tidak Berawak	Drone	Tidak Terdeteksi	0%





Gambar 8 Hasil Evaluasi Data Training Objek Warplane

Pada Gambar 8 menunjukkan nilai box loss selama proses pelatihan. Box loss mengukur seberapa baik model menggambarkan kotak pembatas (*bouding box*) yang mengelilingi objek dalam gambar. Nilai box loss menurun seiring dengan peningkatan *epoch*, menunjukkan bahwa model semakin baik dalam memprediksi lokasi pembatas kotak. Penurunan yang konsisten pada semua jenis loss (baik pada pelatihan maupun validasi) terlihat pada Gambar 9 yang mengindikasikan bahwa model berhasil mengurangi kesalahan deteksi secara keseluruhan. Ini berarti bahwa model tidak hanya belajar dengan baik dari pelatihan data, tetapi juga dapat melakukan generalisasi yang baik terhadap validasi data, yang merupakan indikator penting untuk memastikan bahwa model dapat digunakan dalam situasi nyata.



Gambar 9 Hasil Evaluasi Data Training Objek Drone

4. KESIMPULAN

Hasil penelitian dapat disimpulkan bahwa nilai akurasi pengujian sangat baik: 89,6% untuk citra *warplane* dan 93,5% untuk citra drone. Penelitian ini juga menunjukkan bahwa akurasi yang tinggi tidak menjamin tingginya tingkat pengenalan yang benar. Beberapa variabel independen yang mempengaruhi nilai akurasi antara lain kualitas gambar, kualitas dataset, dan lokasi pengambilan



gambar dari berbagai sudut. Semakin beragamnya model kumpulan data, nilai akurasi yang dihasilkan pun bisa semakin meningkat. Banyaknya daerah pada gambar juga mempengaruhi nilai akurasi. Ini karena ketika objek ditumpuk satu sama lain, objek di belakangnya tidak akan dikenali. Penelitian ke depannya dapat menambah dataset dengan berbagai variasi seperti melakukan pengolahan citra pada dataset yang dikumpulkan untuk meningkatkan kualitas gambar yang lebih jelas dan mempunyai banyak sudut pandang pengambilan. Penelitian ini hanya menggunakan metode YOLOv5 dalam mendeteksi objek, perlu diperinci kembali untuk penelitian ke depannya mengembangkan metode tersebut untuk mendeteksi objek (pesawat) dengan menggunakan kumpulan dataset yang lebih banyak dan beragam kondisi. Penelitian lebih lanjut dapat dilakukan dengan mempertimbangkan faktor eksternal mengenai kualitas gambar yang dihasilkan untuk dideteksi.

DAFTAR PUSTAKA

- Adam, G. (2022). *Aerospace Images*. Kaggle. <https://www.kaggle.com/datasets/gatewayadam/aerospace-images>
- Ahmed, M. W., Almujally, N. A., Alazeb, A., Algarni, A., & Park, J. (2024). Enhanced Object Detection and Classification via Multi-Method Fusion. *Computers, Materials & Continua*, 79(2), 3315–3331. <https://doi.org/10.32604/cmc.2024.046501>
- Akbar, S. A., Ghazali, K. H., Hasan, H., Mohamed, Z., Aji, W. S., & Yudhana, A. (2022). Rapid Bacterial Colony Classification Using Deep Learning. *Indonesian Journal of Electrical Engineering and Computer Science*, 26(1), 352–361. <https://doi.org/10.11591/ijeecs.v26.i1.pp352-361>
- Aksoy, M. Ç., Orak, A. S., Özkan, H. M., & Selimoğlu, B. (2019). *Amateur Unmanned Air Vehicle Detection*. Kaggle. <https://doi.org/10.34740/kaggle/dsv/1019970>
- Anne, S., & Gueye, A. D. (2024). CNN and XGBoost for Automatic Segmentation of Stroke Lesions Using CT Data. *Procedia Computer Science*, 237, 72–79. <https://doi.org/10.1016/j.procs.2024.05.081>
- Argho, A. G., Maswood, M. M. S., Mahmood, Md. I., & Mondol, N. (2024). EfficientCovNet: A CNN-Based Approach to Detect Various Pulmonary Diseases Including COVID-19 Using Modified EfficientNet. *Intelligent Systems with Applications*, 21, Article ID: 200315. <https://doi.org/10.1016/j.iswa.2023.200315>
- Arnia, F., Saddami, K., Roslidar, R., Muharar, R., & Munadi, K. (2024). Towards Accurate Diabetic Foot Ulcer Image Classification: Leveraging CNN Pre-Trained Features and Extreme Learning Machine. *Smart Health*, 33, Article ID: 100502. <https://doi.org/10.1016/j.smhl.2024.100502>
- Arooj, S., Altaf, S., Ahmad, S., Mahmood, H., & Mohamed, A. S. N. (2024). Enhancing Sign Language Recognition Using CNN and SIFT: A Case Study on Pakistan Sign Language. *Journal of King Saud University - Computer and Information Sciences*, 36(2), Article ID: 101934. <https://doi.org/10.1016/j.jksuci.2024.101934>
- Asadi, R., Queguineur, A., Wiikinkoski, O., Mokhtarian, H., Aihkisalo, T., Revuelta, A., & Ituarte, I. F. (2024). Process Monitoring by Deep Neural Networks in Directed Energy Deposition: CNN-Based Detection, Segmentation, and Statistical Analysis of Melt Pools. *Robotics and Computer-Integrated Manufacturing*, 87, Article ID: 102710. <https://doi.org/10.1016/j.rcim.2023.102710>
- Aydin, B., & Singha, S. (2023). Drone Detection Using YOLOv5. *Eng*, 4(1), 416–433. <https://doi.org/10.3390/eng4010025>
- Bilel, K. (2022). *Military Aircraft Recognition Dataset*. Kaggle. <https://www.kaggle.com/datasets/khlaifiabilel/military-aircraft-recognition-dataset>
- Chowdhury, A. A., Hasan Mahmud, S. M., Shahjalal Hoque, K. K., Ahmed, K., Bui, F. M., Lio, P., Moni, M. A., & Al-Zahrani, F. A. (2023). StackFBAs: Detection of Fetal Brain Abnormalities Using CNN with Stacking Strategy from MRI Images. *Journal of King Saud University - Computer and Information Sciences*, 35(8), Article ID: 101647. <https://doi.org/10.1016/j.jksuci.2023.101647>
- D L, S., K, V., N, A., & Vashistha, S. (2024). Tomato Leaf Disease Detection Using CNN. *Procedia Computer Science*, 235, 2975–2984. <https://doi.org/10.1016/j.procs.2024.04.281>



- Das, K., & Baruah, A. K. (2023). Object Detection on Scene Images: A Novel Approach. *Procedia Computer Science*, 218, 153–163. <https://doi.org/10.1016/j.procs.2022.12.411>
- Ellouze, A., Kadri, N., Alaerjan, A., & Ksantini, M. (2024). Combined CNN-LSTM Deep Learning Algorithms for Recognizing Human Physical Activities in Large and Distributed Manners: A Recommendation System. *Computers, Materials & Continua*, 79(1), 351–372. <https://doi.org/10.32604/cmc.2024.048061>
- Fahmi, M., Yudhana, A., & Sunardi, S. (2023). Pemilahan Sampah Menggunakan Model Klasifikasi Support Vector Machine Gabungan dengan Convolutional Neural Network. *JURIKOM (Jurnal Riset Komputer)*, 10(1), 76–81. <https://mti.uad.ac.id/download/pemilahan-sampah-menggunakan-model-klasifikasi-support-vector-machine-gabungan-dengan-convolutional-neural-network/>
- Feng, C., Wang, C., Zhang, D., Kou, R., & Fu, Q. (2024). Enhancing Dense Small Object Detection in UAV Images Based on Hybrid Transformer. *Computers, Materials & Continua*, 78(3), 3993–4013. <https://doi.org/10.32604/cmc.2024.048351>
- Fransisca, V., & Santoso, H. (2023). Penerapan Gamma Correction dalam Peningkatan Pendeteksian Objek Malam pada Algoritma YOLOv5. *Building of Informatics, Technology and Science (BITS)*, 5(1), 59–69. <https://doi.org/10.47065/bits.v5i1.3553>
- Freitas, N. R., Vieira, P. M., Tinoco, C., Anacleto, S., Oliveira, J. F., Vaz, A. I. F., Laguna, M. P., Lima, E., & Lima, C. S. (2024). Multiple Mask and Boundary Scoring R-CNN with cGAN Data Augmentation for Bladder Tumor Segmentation in WLC Videos. *Artificial Intelligence in Medicine*, 147, Article ID: 102723. <https://doi.org/10.1016/j.artmed.2023.102723>
- Giménez-Gallego, J., Martínez-del-Rincon, J., González-Teruel, J. D., Navarro-Hellín, H., Navarro, P. J., & Torres-Sánchez, R. (2024). On-Tree Fruit Image Segmentation Comparing Mask R-CNN and Vision Transformer Models. Application in a Novel Algorithm for Pixel-Based Fruit Size Estimation. *Computers and Electronics in Agriculture*, 222, Article ID: 109077. <https://doi.org/10.1016/j.compag.2024.109077>
- Goud, P. A. K., Raj, G. M., Rahul, K., & Lakshmi, A. V. (2023). Military Aircraft Detection Using YOLOv5. In *Intelligent Communication Technologies and Virtual Mobile Networks* (pp. 865–878). Springer. https://doi.org/10.1007/978-981-99-1767-9_63
- Huang, M., Yan, W., Dai, W., & Wang, J. (2023). EST-YOLOv5s: SAR Image Aircraft Target Detection Model Based on Improved YOLOv5s. *IEEE Access*, 11, 113027–113041. <https://doi.org/10.1109/ACCESS.2023.3323575>
- Inamullah, I., Hassan, S., Belhaouari, S. B., & Amin, I. (2024). Deciphering the Impact of Diversity in CNN-Based Ensembles on Overcoming Data Imbalance and Scarcity in Medical Datasets: A Case Study on Diabetic Retinopathy. *Informatics in Medicine Unlocked*, 49, Article ID: 101557. <https://doi.org/10.1016/j.imu.2024.101557>
- Jung, H.-K., & Choi, G.-S. (2022). Improved YOLOv5: Efficient Object Detection Using Drone Images Under Various Conditions. *Applied Sciences*, 12(14), Article ID: 7255. <https://doi.org/10.3390/app12147255>
- Karimah, S., Pangestu, R. T., Febriansyah, A., & Irwan, I. (2024). Implementasi Metode YOLOv5 pada Sistem Pendeteksi Jentik Nyamuk Berbasis IoT. *Jurnal Inovasi Teknologi Terapan*, 2(2), 417–425. <https://doi.org/10.33504/jitt.v2i2.184>
- Khairunisa, N., Carudin, C., & Jamaludin, A. (2024). Analisis Perbandingan Algoritma CNN dan YOLO dalam Mengidentifikasi Kerusakan Jalan. *Jurnal Informatika dan Teknik Elektro Terapan*, 12(3), 1756–1768. <https://doi.org/10.23960/jitet.v12i3.4434>
- Lian-suo, W. E. I., Shen-hao, H., & Long-yu, M. (2024). MTD-YOLOv5: Enhancing Marine Target Detection with Multi-Scale Feature Fusion in YOLOv5 Model. *Heliyon*, 10(4), Article ID: e26145. <https://doi.org/10.1016/j.heliyon.2024.e26145>
- Liu, G., Hu, Y., Chen, Z., Guo, J., & Ni, P. (2023). Lightweight Object Detection Algorithm for Robots with Improved YOLOv5. *Engineering Applications of Artificial Intelligence*, 123, Article ID: 106217. <https://doi.org/10.1016/j.engappai.2023.106217>
- Liu, H., Wang, W., Liu, H., Yi, S., Yu, Y., & Yao, X. (2024). A Degradation Type Adaptive and Deep CNN-Based Image Classification Model for Degraded Images. *Computer Modeling in Engineering & Sciences*, 138(1), 459–472. <https://doi.org/10.32604/cmes.2023.029084>
- Molnar, C. (2020). *Interpretable Machine Learning* (3rd ed.). Lean Publishing.



- Muis, A., Sunardi, S., & Yudhana, A. (2023). Comparison Analysis of Brain Image Classification Based on Thresholding Segmentation with Convolutional Neural Network. *Journal of Applied Engineering and Technological Science (JAETS)*, 4(2), 664–673. <https://doi.org/10.37385/jaets.v4i2.1583>
- Murthy, J. S., Siddesh, G. M., Lai, W.-C., Parameshachari, B. D., Patil, S. N., & Hemalatha, K. L. (2022). ObjectDetect: A Real-Time Object Detection Framework for Advanced Driver Assistant Systems Using YOLOv5. *Wireless Communications and Mobile Computing*, 2022(1), Article ID: 9444360. <https://doi.org/10.1155/2022/9444360>
- Oyedeji, O. A., Khan, S., & Erkoyuncu, J. A. (2024). Application of CNN for Multiple Phase Corrosion Identification and Region Detection. *Applied Soft Computing*, 164, Article ID: 112008. <https://doi.org/10.1016/j.asoc.2024.112008>
- Peryanto, A., Yudhana, A., & Umar, R. (2020a). Klasifikasi Citra Menggunakan Convolutional Neural Network dan K Fold Cross Validation. *Journal of Applied Informatics and Computing*, 4(1), 45–51. <https://doi.org/10.30871/jaic.v4i1.2017>
- Peryanto, A., Yudhana, A., & Umar, R. (2020b). Rancang Bangun Klasifikasi Citra dengan Teknologi Deep Learning Berbasis Metode Convolutional Neural Network. *Format: Jurnal Ilmiah Teknik Informatika*, 8(2), 138–147. <https://doi.org/10.22441/format.2019.v8.i2.007>
- Pramanik, P. K. D., Pal, S., Mukhopadhyay, M., & Singh, S. P. (2021). Big Data Classification: Techniques and Tools. In *Applications of Big Data in Healthcare* (pp. 1–43). Elsevier. <https://doi.org/10.1016/B978-0-12-820203-6.00002-3>
- Pratama, B. K., Sri Lestanti, & Yusniarsi Primasari. (2024). Implementasi Algoritma You Only Look Once (YOLO) untuk Mendeteksi Bahasa Isyarat SIBI. *ProTekInfo (Pengembangan Riset dan Observasi Teknik Informatika)*, 11(2), 7–14. <https://doi.org/10.30656/protekinfo.v11i2.9105>
- Qiu, K., Poon, P.-L., Zhao, S., Towey, D., & Yu, L. (2024). Improving the Validation of Multiple-Object Detection Using a Complex-Network-Community-Based Relevance Metric. *Knowledge-Based Systems*, 299, Article ID: 112027. <https://doi.org/10.1016/j.knosys.2024.112027>
- Ramakrishnan, A. B., Sridevi, M., Vasudevan, S. K., Manikandan, R., & Gandomi, A. H. (2024). Optimizing Brain Tumor Classification with Hybrid CNN Architecture: Balancing Accuracy and Efficiency Through oneAPI Optimization. *Informatics in Medicine Unlocked*, 44, Article ID: 101436. <https://doi.org/10.1016/j.imu.2023.101436>
- Reddy, S. P. K., Harikiran, J., & Chandana, B. S. (2024). Deep CNN Based Multi Object Detection and Tracking in Video Frames with Mean Distributed Feature Set. *Procedia Computer Science*, 235, 723–734. <https://doi.org/10.1016/j.procs.2024.04.069>
- Ren, Z., Zhang, H., & Li, Z. (2023). Improved YOLOv5 Network for Real-Time Object Detection in Vehicle-Mounted Camera Capture Scenarios. *Sensors*, 23(10), Article ID: 4589. <https://doi.org/10.3390/s23104589>
- Rishi, G. H., Kumar, P. R., Varshit, N. C., & Sailaja, K. (2024). Real-Time Military Aircraft Detection Using YOLOv5. *2024 Second International Conference on Data Science and Information System (ICDSIS)*, 1–7. <https://doi.org/10.1109/ICDSIS61070.2024.10594590>
- S, V., M, B., & V, K. (2024). Satellite Image Classification Using CNN with Particle Swarm Optimization Classifier. *Procedia Computer Science*, 233, 979–987. <https://doi.org/10.1016/j.procs.2024.03.287>
- Saha, B. (2022). *Drone-Bird Classification*. Kaggle. <https://www.kaggle.com/datasets/imbikramsaha/drone-bird-classification>
- Saha, D. K., Joy, A. M., & Majumder, A. (2024). YoTransViT: A Transformer and CNN Method for Predicting and Classifying Skin Diseases Using Segmentation Techniques. *Informatics in Medicine Unlocked*, 47, Article ID: 101495. <https://doi.org/10.1016/j.imu.2024.101495>
- Saifullah, S., Drezewski, R., Yudhana, A., Pranolo, A., Kaswijanti, W., Suryotomo, A. P., Putra, S. A., Khaliduzzaman, A., Prabuwo, A. S., & Japkowicz, N. (2023). Nondestructive Chicken Egg Fertility Detection Using CNN-Transfer Learning Algorithms. *Jurnal Ilmiah Teknik Elektro Komputer dan Informatika*, 9(3), 854–871. <https://doi.org/10.26555/jiteki.v9i3.26722>



- Sapakova, S., Sapakov, A., & Yilibule, Y. (2024). A YOLOv5-Based Model for Real-Time Mask Detection in Challenging Environments. *Procedia Computer Science*, 231, 267–274. <https://doi.org/10.1016/j.procs.2023.12.202>
- Schenk, P. O., & Kern, C. (2024). Connecting Algorithmic Fairness to Quality Dimensions in Machine Learning in Official Statistics and Survey Production. *AStA Wirtschafts- und Sozialstatistisches Archiv*, 18(2), 131–184. <https://doi.org/10.1007/s11943-024-00344-2>
- Sharma, V., Kannan, S., Tanya, S., & Panda, N. (2024). Detecting Plant Diseases at Scale: A Distributed CNN Approach with PySpark and Hadoop. *Procedia Computer Science*, 235, 1044–1057. <https://doi.org/10.1016/j.procs.2024.04.099>
- Sileo, M., Capece, N., Gruosso, M., Nigro, M., Bloisi, D. D., Pierri, F., & Erra, U. (2024). Vision-Enhanced Peg-in-Hole for Automotive Body Parts Using Semantic Image Segmentation and Object Detection. *Engineering Applications of Artificial Intelligence*, 128, Article ID: 107486. <https://doi.org/10.1016/j.engappai.2023.107486>
- Sotoude, D., Hoseinkhani, M., & Amiri Tehranizadeh, A. (2023). Context-Aware Fusion of Transformers and CNNs for Medical Image Segmentation. *Informatics in Medicine Unlocked*, 43, Article ID: 101396. <https://doi.org/10.1016/j.imu.2023.101396>
- Sulistyo, W. Y., Riadi, I., & Yudhana, A. (2018). Analisis Deteksi Keaslian Citra Menggunakan Teknik Error Level Analysis dengan Forensicallybeta. *Seminar Nasional Informatika 2018 (SEMNASIF 2018)*, 1(1), 154–159. <https://jurnal.upnyk.ac.id/index.php/semnasif/article/view/2632>
- Sunardi, S., Yudhana, A., & Saifullah, S. (2017). Identity Analysis of Egg Based on Digital and Thermal Imaging: Image Processing and Counting Object Concept. *International Journal of Electrical and Computer Engineering (IJECE)*, 7(1), 200–208. <https://doi.org/10.11591/ijece.v7i1.pp200-208>
- Sunardi, S., Yudhana, A., & WindraPutri, A. R. (2022). Mass Classification of Breast Cancer Using CNN and Faster R-CNN Model Comparison. *Kinetik: Game Technology, Information System, Computer Network, Computing, Electronics, and Control*, 243–250. <https://doi.org/10.22219/kinetik.v7i3.1462>
- Surya, R. A., Fadlil, A., & Yudhana, A. (2019). Identification of Pekalongan Batik Images Using Backpropagation Method. *Journal of Physics: Conference Series*, 1373(1), Article ID: 012049. <https://doi.org/10.1088/1742-6596/1373/1/012049>
- Tan, G., & Liu, Y. (2023). Picture Processing Optimization Technology Based on Mask R-CNN Algorithm. *Procedia Computer Science*, 228, 647–654. <https://doi.org/10.1016/j.procs.2023.11.075>
- Tjahyanto, A., & Atletiko, F. J. (2022). Peningkatan Kinerja Pengklasifikasi Objek Bawah Laut dengan Deep Learning. *MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, 21(3), 753–760. <https://doi.org/10.30812/matrik.v21i3.1466>
- Umar, R., Riadi, I., & Miladiah, M. (2018). Sistem Identifikasi Keaslian Uang Kertas Rupiah Menggunakan Metode K-Means Clustering. *Techno.Com*, 17(2), 179–185. <https://doi.org/10.33633/tc.v17i2.1681>
- Wahyudi, A. A., Khumaidi, A., Rahmat, M. B., Riananda, D. P., Syai'in, M., & Endrasmono, J. (2024). Implementasi Robot Operating System (ROS) untuk Meningkatkan Akurasi Deteksi Bola Menggunakan YOLO V5 pada KRSBI-Beroda. *Jurnal Elektronika dan Otomasi Industri*, 11(2), 648–661. <https://doi.org/10.33795/elkolind.v11i2.5234>
- Wang, M., Yang, W., Wang, L., Chen, D., Wei, F., KeZiErBieKe, H., & Liao, Y. (2023). FE-YOLOv5: Feature Enhancement Network Based on YOLOv5 for Small Object Detection. *Journal of Visual Communication and Image Representation*, 90, Article ID: 103752. <https://doi.org/10.1016/j.jvcir.2023.103752>
- Wittmer, A., & Müller, A. (2021). The Environment of Aviation. In *Aviation Systems* (pp. 79–117). Management of the Integrated Aviation Value Chain, 79–117. https://doi.org/10.1007/978-3-030-79549-8_3
- Xu, S., Zhang, M., Chen, J., & Zhong, Y. (2024). YOLO-HyperVision: A Vision Transformer Backbone-Based Enhancement of YOLOv5 for Detection of Dynamic Traffic Information. *Egyptian Informatics Journal*, 27, Article ID: 100523. <https://doi.org/10.1016/j.eij.2024.100523>



- Ye, Z., Yang, K., Lin, Y., Guo, S., Sun, Y., Chen, X., Lai, R., & Zhang, H. (2023). A Comparison Between Pixel-Based Deep Learning and Object-Based Image Analysis (OBIA) for Individual Detection of Cabbage Plants Based on UAV Visible-Light Images. *Computers and Electronics in Agriculture*, 209, Article ID: 107822. <https://doi.org/10.1016/j.compag.2023.107822>
- Yi, D., Ahmedov, H. B., Jiang, S., Li, Y., Flinn, S. J., & Fernandes, P. G. (2024). Coordinate-Aware Mask R-CNN with Group Normalization: A Underwater Marine Animal Instance Segmentation Framework. *Neurocomputing*, 583, Article ID: 127488. <https://doi.org/10.1016/j.neucom.2024.127488>
- Yudhana, A., Sunardi, S., & Saifullah, S. (2017). Segmentation Comparing Eggs Watermarking Image and Original Image. *Bulletin of Electrical Engineering and Informatics*, 6(1), 47–53. <https://doi.org/10.11591/eei.v6i1.595>
- Zhang, J., Hou, C., Yang, X., Yang, X., Yang, W., & Cui, H. (2024). Advancing Face Detection Efficiency: Utilizing Classification Networks for Lowering False Positive Incidences. *Array*, 22, Article ID: 100347. <https://doi.org/10.1016/j.array.2024.100347>





9 772527 583007



LABORATORIUM AGAMA
MASJID SUNAN KALIJAGA