



Analisis Pengaruh Recursive Feature Elimination Terhadap Kinerja Model Prediksi Dini *Diabetes Mellitus* di RS PKU Muhammadiyah Bima

Nani Sulistianingsih^{1*}, Siti Agrippina Alodia Yusuf¹, Anggreni¹, Firda Niken Sari¹, M. Pandawan Juniarta¹,
Juryanti Permatasari¹

¹ Program Studi Sistem dan Teknologi Informasi, Universitas Muhammadiyah Mataram, Indonesia

* Korespondensi: nani.sulistianingsih@ummat.ac.id

Sitasi: Sulistianingsih, N.; Siti Agrippina Alodia Yusuf, S. A. A.; Anggreni, A.; Sari, F. N.; Juniarta, M. P.; and Permatasari, J. (2025). Analisis Pengaruh Recursive Feature Elimination Terhadap Kinerja Model Prediksi Dini Diabetes Mellitus di RS PKU Muhammadiyah Bima. JTIM: Jurnal Teknologi Informasi Dan Multimedia, 7(3), 548-560. <https://doi.org/10.35746/jtim.v7i3.774>

Diterima: 15-06-2025

Direvisi: 09-07-2025

Disetujui: 11-07-2025



Copyright: © 2025 oleh para penulis. Karya ini dilisensikan di bawah Creative Commons Attribution-ShareAlike 4.0 International License. (<https://creativecommons.org/licenses/by-sa/4.0/>).

Abstract: Early detection of diabetes mellitus is a crucial step in preventing chronic complications and enabling more effective disease management. This study aims to analyze the impact of the Recursive Feature Elimination (RFE) method on the performance of machine learning-based diabetes prediction models at RS PKU Muhammadiyah Bima. A quantitative approach was employed by implementing the CRISP-DM framework, encompassing data selection, preprocessing, transformation, data mining, and model evaluation. Missing values in height and weight variables were imputed using linear regression based on age and gender features. The transformation process included calculating the Body Mass Index (BMI) as a new feature relevant to diabetes risk. Evaluation was carried out on three classification algorithms—Support Vector Machine (SVM), Logistic Regression (LR), and Random Forest (RF)—both before and after the application of RFE. The results showed that all models experienced significant performance improvements following feature selection with RFE, achieving 100% in all evaluation metrics. Insulin and BMI were consistently selected features, underscoring their contribution to diabetes detection. It can be concluded that RFE effectively reduces model complexity without sacrificing accuracy, thereby supporting the efficient implementation of predictive models in clinical settings.

Keywords: Diabetes Mellitus, RFE, Feature Selection

Abstrak: Deteksi dini diabetes mellitus merupakan langkah krusial dalam upaya pencegahan komplikasi kronis dan pengelolaan penyakit yang lebih efektif. Penelitian ini bertujuan untuk menganalisis pengaruh metode *Recursive Feature Elimination* (RFE) terhadap kinerja model prediksi diabetes berbasis pembelajaran mesin di RS PKU Muhammadiyah Bima. Pendekatan kuantitatif digunakan dengan menerapkan kerangka kerja *CRIPS-DM*, mencakup pemilihan data, pra-pemrosesan, transformasi, penambahan data, evaluasi model. Imputasi nilai hilang pada *variable* tinggi dan berat badan dilakukan menggunakan regresi linier berbasis fitur usia dan jenis kelamin. Proses transformasi mencakup perhitungan *Body Mass Index* (BMI) sebagai fitur baru yang relevan terhadap resiko diabetes. Evaluasi dilakukan terhadap tiga algoritma klasifikasi yaitu, SVM, LR dan *Random Forest* (RF), baik sebelum maupun sesudah penerapan RFE. Hasil menunjukkan bahwa semua model mengalami peningkatan signifikan setelah seleksi fitur dengan RFE mencapai 100%. Fitur yang terpilih secara konsisten adalah insulin dan BMI, menegaskan kontribusi keduanya dalam deteksi diabetes. Sehingga disimpulkan RFE mampu menyederhanakan kompleksitas model tanpa mengorbankan akurasi sehingga mendukung penerapan model prediksi dalam konteks klinis secara efisien.

Kata kunci: Diabetes Mellitus, RFE, Seleksi Fitur

1. Pendahuluan

Diabetes Mellitus merupakan istilah yang sering digunakan untuk menggambarkan kondisi di mana tubuh tidak dapat mengontrol kadar gula (glukosa) dalam darah dengan baik [1]. Kadar gula darah yang terus-menerus tinggi dapat menyebabkan penyakit kronis lainnya seperti *kardiovaskular*, gagal ginjal, mata, saraf, hingga kematian dini [2], [3]. Berdasarkan data International Diabetes Federation (IDF) pada tahun 2021 ada lebih dari 463 juta orang yang terkena penyakit *diabetes mellitus* dengan rentang umur 20-79 tahun [4]. Angka penderita penyakit *diabetes mellitus* kemungkinan akan terus meningkat sekitar 578 juta orang pada tahun 2030 dan 700 juta orang pada tahun 2045 [5]. Termasuk di Indonesia penderita *diabetes mellitus* terus menunjukkan peningkatan. menjadi penyakit dengan penyebab kematian ketiga di Indonesia. Menurut Perkenkes Nomor 30 Tahun 2023, anjuran konsumsi gula harian adalah 4 sendok makan atau setara dengan 50 gram per orang per hari. Konsumsi gula harian yang berlebih dapat menyebabkan resiko kesehatan termasuk *diabetes mellitus* [6]. Faktor-faktor seperti urbanisasi, populasi yang menua, menurunnya aktivitas fisik dan meningkatnya angka kelebihan berat badan atau obesitas memberikan kontribusi signifikan terhadap meningkatnya penderita. Hal ini tentu dapat menurunkan kualitas hidup bagi penderita, tetapi juga menimbulkan beban finansial yang signifikan bagi sistem kesehatan [7].

RS PKU Muhammadiyah Bima merupakan salah satu fasilitas kesehatan di Kota Bima yang menangani pasien penderita. Namun, sama seperti rumah sakit lainnya, tantangan yang dihadapi adalah diagnosis dini yang akurat dan prediksi resiko diabetes yang efisien. Berdasarkan data Dinas Kesehatan Provinsi NTB jumlah penderita yang mendapatkan layanan kesehatan pada 2023 sebanyak 71.031 orang yang tersebar di 10 kabupaten dan kota. Kota dan Kabupaten Bima menjadi penyumbang penderita tertinggi ke 5 di NTB dengan total jumlah 8.927 orang [8]. Hal ini tentu menjadi tantangan bagi RS PKU Muhammadiyah Kota Bima dalam memberikan layanan yang cepat dan akurat terkait dengan penanganan penyakit *diabetes mellitus*.

Penerapan Recursive Feature Elimination (RFE) dalam memprediksi *Diabetes Mellitus* sangat penting untuk meningkatkan kinerja model dengan mengurangi dimensi dan meningkatkan relevansi fitur. RFE bekerja dengan melatih model awal dan mengurutkan fitur berdasarkan tingkat kepentingannya terhadap prediksi. Fitur yang paling tidak relevan dieliminasi secara rekursif, kemudian model dilatih ulang hingga subset fitur optimal diperoleh. Pendekatan ini terbukti efektif dalam meningkatkan akurasi model, sebagaimana ditunjukkan oleh penelitian yang dilakukan [9] yang menggunakan kombinasi RFE dan algoritma *Random Forest* pada dataset PIM India dan mencapai akurasi 81,25%. Selain meningkatkan performa, RFE juga memperkuat interpretabilitas model dan mengurangi risiko *overfitting*, terutama pada data berdimensi tinggi [10]. Penelitian lain [11] mengusulkan kerangka kerja dinamis berbasis RFE untuk menyempurnakan pemilihan fitur, dengan hasil akurasi yang lebih tinggi. Namun, perspektif alternatif juga perlu dipertimbangkan, karena beberapa studi seperti [12] menemukan bahwa pendekatan lain seperti *Principal Component Loading Selection* dapat menghasilkan kinerja sebanding, tergantung pada karakteristik data dan algoritma yang digunakan.

Pemilihan fitur yang tepat melalui RFE terbukti secara signifikan meningkatkan kinerja model prediksi diabetes. Pada berbagai penelitian, algoritma pembelajaran mesin seperti *Support Vector Machine* (SVM), *Logistic Regression*, dan *Random Forest* menunjukkan peningkatan akurasi setelah seleksi fitur dilakukan. Misalnya [13] melaporkan bahwa akurasi *Random Forest* meningkat hingga 99,75% setelah penerapan RFE. Lebih lanjut, [14] menunjukkan peningkatan presisi sebesar 96,7%, recall 84,7%, dan F1-score 90,2% untuk model *Random Forest* pasca seleksi fitur. Analisis *confusion matrix* juga memberikan wawasan penting tentang kinerja klasifikasi, memungkinkan identifikasi kesalahan

prediksi dan ketidakseimbangan kelas [15]. Model lain seperti SVM dan *Logistic Regression* juga mengalami peningkatan performa setelah fitur tidak relevan dieliminasi [16]. Namun, perlu diperhatikan bahwa peningkatan performa tidak selalu bersifat universal, karena efektivitas seleksi fitur dapat bergantung pada hubungan antar fitur dan sifat data yang digunakan.

Prediksi dini *Diabetes Mellitus* tidak sebatas dari sisi teknik, tetapi juga memiliki nilai penting dalam praktik klinis, terutama untuk pencegahan dan intervensi dini. Melalui metode RFE, sejumlah fitur penting telah diidentifikasi yang berkaitan erat dengan risiko diabetes, termasuk kreatinin serum, urea darah, *trigliserida*, *hemoglobin*, dan indeks massa tubuh (BMI) [17], [18]. Kreatinin dan urea mencerminkan fungsi ginjal dan metabolisme tubuh yang terganggu pada pasien diabetes, sedangkan kadar *trigliserida* tinggi mengindikasikan resisten insulin. Hemoglobin dan BMI juga berperan penting dalam pemantauan jangka panjang glukosa dan risiko obesitas sebagai faktor komorbid [19]. Pada konteks klinis, fitur-fitur ini dapat digunakan untuk menyaring pasien berisiko tinggi secara lebih akurat. Model prediktif yang dihasilkan dapat diintegrasikan ke dalam sistem rekam medis elektronik, memungkinkan penilaian risiko secara real-time dan mendukung pengambilan keputusan berbasis data. Namun demikian, penting untuk menyeimbangkan antara kecanggihan teknologi dan intuisi klinis, agar pengambilan keputusan tetap mempertimbangkan konteks individu pasien.

Meskipun efektivitas *Recurvise Feature Elimination* (RFE) dalam meningkatkan performa model prediktif telah banyak dibuktikan pada studi sebelumnya, sebagian besar penelitian tersebut masih mengandalkan dataset publik seperti PIMA India, yang kurang merepresentasikan kondisi populasi lokal secara spesifik. Selain itu, keterbatasan juga tampak pada kurangnya kajian yang mengintegrasikan model prediktif ke dalam praktik klinis berbasis data nyata dari rumah sakit swasta, sehingga potensi penerapan model dalam konteks lokal belum sepenuhnya tergali. Pada konteks tersebut, pemanfaatan data klinis riil dari RS PKU Muhammadiyah Bima menjadi penting untuk menilai apakah pendekatan seleksi fitur berbasis RFE tetap efektif dalam meningkatkan akurasi dan efisiensi model, serta fitur-fitur klinis yang mana secara konsisten berkontribusi dalam klasifikasi risiko *diabetes mellitus*. Oleh karena itu, tujuan utama dari penelitian ini adalah untuk mengevaluasi dan menganalisis bagaimana metode RFE dapat meningkatkan performa klasifikasi model prediksi dini *diabetes mellitus* berbasis pembelajaran mesin (*Machine Learning*) dengan data klinis lokal di RS PKU Muhammadiyah Bima, serta mengidentifikasi fitur-fitur utama yang relevan secara klinis guna mendukung proses pengambilan keputusan medis yang lebih tepat sasaran.

2. Bahan dan Metode

Studi ini menggunakan pendekatan kuantitatif dalam menganalisis data pasien *diabetes mellitus*, dengan menerapkan metode klasifikasi berbasis pembelajaran mesin (*Machine Learning*). Kerangka kerja yang digunakan adalah CRISP-DM (*Cross-Industry Standard Process for Data Mining*), yang terdiri dari tahapan utama yaitu pemilihan data, pra-pemrosesan, transformasi, penambangan data, dan evaluasi model. Setiap tahapan disusun secara sistematis untuk mendukung pengembangan model prediksi dini yang akurat dan dapat diandalkan [20].



Gambar 1. Alur Penelitian

2.1. Seleksi Data

Pada tahap seleksi data, data diperoleh dari rekam medis pasien di RS PKU Muhammadiyah yang mencakup informasi tahun kunjungan, bulan kunjungan, tanggal lahir, alamat, usia, tinggi badan, berat badan, jenis kelamin, tekanan darah, kadar glukosa, insulin serta hasil diagnosa sebagai label target. Data dipilih berdasarkan kelengkapan rekam medis, dengan pengecualian untuk pasien yang memiliki kormobiditas kronis berat lainnya yang dapat mempengaruhi hasil analisis. Kriteria inklusi mencakup pasien dewasa berusia lebih dari 30 tahun dengan data yang lengkap dan konsisten. Dataset terdiri dari tahun 2022 sampai 2024 yang berjumlah 451 data. Data akan dijadikan data latih dan data uji dari penelitian. Dataset diperoleh dari sistem informasi elektronik rekam medis atau *Electronic Medical Record* (EMR) RS PKU Muhammadiyah Mataram.

2.2. Pra-Pemrosesan Data

Tahap pra-pemrosesan data melibatkan sejumlah teknik untuk memastikan kualitas data sebelum dianalisis lebih lanjut. Tahap ini dilakukan pembersihan data, seperti mendeteksi nilai yang hilang (*missing value*) dan data duplikasi. Untuk menangani nilai yang hilang, digunakan teknik imputasi berbasis regresi linier, yang memungkinkan pengisian nilai yang hilang secara prediktif berdasarkan hubungan antar fitur dalam *dataset* [21]. Selanjutnya, dilakukan normalisasi data menggunakan dua pendekatan, yaitu *min-max scaling* dan *z-score normalization*, untuk menyelaraskan skala fitur dan meningkatkan stabilitas algoritma pembelajaran mesin [22].

2.3. Transformasi Data

Pada tahap transformasi data, dilakukan pembentukan fitur baru dari atribut yang tersedia. Salah satu transformasi penting adalah perhitungan *Body Mass Index* (BMI) yang diperoleh dari rumus berat badan (kg) dibagi tinggi badan kuadrat (m^2) Mhd [23]. Transformasi ini merupakan bagian dari proses rakayasa fitur (*feature engineering*), di mana BMI digunakan sebagai indikator status gizi dan risiko obesitas yang berkorelasi dengan komplikasi diabetes [24]. Selain itu pada fitur jenis kelamin dan hasil diagnosa menggunakan teknik *one-hot-encoding* yaitu teknik merubah bentuk data kategorik menjadi data numerik 0 dan 1 [22]. Transformasi semacam ini bertujuan untuk meningkatkan daya informatif fitur dalam proses pemodelan.

2.4. Penambangan Data

Tahap penambangan data (*data mining*) atau dalam teknik CRISP-DM sebagai modeling phase ini melibatkan penerapan dan perbandingan beberapa algoritma klasifikasi, yaitu *Random Forest* (RF), *Support Vector Machine* (SVM), dan *Logistic Regression* (LR). Tahap ini juga dilakukan seleksi fitur menggunakan metode *Recursive Feature Elimination* (RFE), yang secara *iterative* menghapus fitur-fitur yang kurang signifikan untuk meningkatkan kinerja model dan mengurangi kompleksitas data [25]. Pemilihan algoritma ini didasarkan pada historinya dalam studi prediksi diabetes dan kemampuan mereka dalam menangani data multivariabel serta distribusi *non-linear*. Untuk menjamin generalisasi model, digunakan teknik validasi silang 5 fold *cross validation* yang secara luas digunakan untuk menilai performa model pada data yang tidak terlihat sebelumnya dan mengurangi bias evaluasi [26]. Selain itu algoritma *Random Forest* yang memiliki sifat *ensemble learning* dengan pendekatan bagging, yaitu membangun pohon keputusan pada subset data berbeda dan menggabungkan hasilnya untuk meningkatkan stabilitas dan akurasi, sekaligus mengurangi resiko *overfitting* [27]. Sementara pada algoritma SVM dan LR, kontrol terhadap *overfitting* dilakukan melalui penyesuaian parameter regularisasi. Pada SVM, parameter C diatur untuk mengendalikan margin kesalahan dan kompleksitas model [28], sedangkan pada LR digunakan regularisasi L2 untuk menghindari bobot koefisien yang terlalu besar [29].

2.5. Evaluasi Model

Evaluasi model dilakukan dengan mengukur metrik performa utama, yaitu akurasi, *presisi*, *recall*, dan *F1-score* [30]. Hal ini merupakan praktik standar dalam *machine learning* karena memungkinkan penilaian kualitas klasifikasi. Untuk mendalami hasil evaluasi secara visual, digunakan *confusion matrix* yang menggambarkan jumlah prediksi benar dan salah untuk kedua kelas dan membantu analisis kesalahan model.

3. Hasil

Bagian ini menyajikan hasil penelitian yang dilakukan untuk mengevaluasi pengaruh metode *Recursive Feature Elimination* (RFE) terhadap kinerja model prediksi dini diabetes mellitus. Studi ini membandingkan performa tiga algoritma klasifikasi yaitu, SVM, LR, dan RF sebelum dan sesudah penerapan RFE. Evaluasi dilakukan menggunakan metrik akurasi, *presisi*, *recall*, dan *F1-score*, serta divisualisasikan melalui *confusion matrix*. Penelitian ini bertujuan untuk menilai efektivitas RFE dalam meningkatkan efisiensi dan akurasi model prediktif berbasis data klinis di RS PKU Muhammadiyah Bima.

3.1. Pra-Pemrosesan Data

Data yang digunakan terdiri dari data rekam medis pasien rawat jalan di RS PKU Muhammadiyah Bima, dengan jumlah total 451 data dengan rasio pembagian 70% data uji dan 30% data latih. Dari total data sebanyak 229 data dengan pasien normal dan 222 data pasien *diabetes mellitus*. Fitur dalam dataset mencakup informasi klinis seperti tahun kunjungan, bulan kunjungan, tanggal lahir, alamat, usia, tinggi badan, berat badan, jenis kelamin, tekanan darah, kadar glukosa, insulin serta hasil diagnosa sebagai label target.

Tabel 1. Dataset Diabetes Mellitus

Tahun_Kunjungan	Bulan_Kunjungan	Tanggal_Lahir	Alamat	Usia	TB	BB	JK	Tekanan_Darah	Kadar_Glukosa	Insulin
2022	Januari	30/07/1970	Mande	51	0	0	P	110/80	419	1
2022	April	17/08/1958	Jatiwangi	65	0	0	P	110/80	256	1
2022	Juli	02/12/1978	Ambalawi	43	0	0	P	100/70	390	1
2022	Oktober	04/01/1957	Tonggondoa	65	0	0	P	120/80	262	1
2022	Januari	01/06/1975	Bolo	47	0	0	L	100/70	389	1
...	...	...	...	...	...	...	...	...	...	...
2024	Oktober	31/12/1963	Kumbe	58	158	65	P	110/70	196	1
2024	Januari	02/11/1969	Langgudu	52	146	40	P	110/70	164	1

Analisis awal terhadap dataset dilakukan untuk mengidentifikasi keberadaan nilai yang hilang (*missing value*) dan entri data yang duplikat. Pada Tabel 1 menunjukkan bahwa nilai yang hilang terdeteksi secara spesifik pada dua fitur, yaitu Tinggi Badan (TB) dan Berat Badan (BB). Keberadaan *missing value* pada fitur seperti TB dan BB sangat penting untuk diperhatikan karena keduanya merupakan komponen utama dalam perhitungan *Body Mass Index* (BMI), yang merupakan indikator status gizi dan faktor risiko penting dalam diagnosis *diabetes mellitus* [31], [32]. Keberadaan *missing value* dan data duplikasi berdampak signifikan terhadap performa dan *machine learning* [33], [34].

Tabel 2. Fitur Pada Dataset Diabetes Mellitus

Usia	Jenis_Kelamin	Tekanan_Darah	Kadar_Glukosa	Insulin	BMI	Hasil_Diagnosa
74	0	80	221	1	21.8	1
38	0	82	106	0	39.5	0
49	0	80	190	1	24.1	1
78	0	100	252	1	21.5	1
51	0	86	125	0	37.6	0
...	...	...	...	...	...	...

Usia	Jenis_Kelamin	Tekanan_Darah	Kadar_Glukosa	Insulin	BMI	Hasil_Diagnosa
71	1	70	279	1	22.0	1
46	0	82	61	0	34.4	0

Pada penelitian ini, nilai yang hilang pada fitur TB dan BB diimputasi menggunakan regresi linier berbasis fitur jenis kelamin dan usia. Pendekatan ini dipilih karena regresi linier memungkinkan estimasi nilai numerik yang hilang berdasarkan hubungan linear antar fitur, yang telah terbukti efektif dalam mengatasi *missing value* pada data kesehatan [35]. Selanjutnya, dilakukan transformasi data kategorikal menggunakan teknik *one-hot encoding* pada fitur jenis kelamin dan insulin. Pada proses ini, nilai 0 diberikan untuk Perempuan dan 1 untuk laki-laki, sedangkan untuk fitur insulin, nilai 0 menunjukkan kondisi non-diabetes dan 1 menunjukkan pasien dengan diabetes. *One-hot encoding* dipilih karena mampu menghindari asumsi ordinal pada data kategorikal dan telah banyak digunakan dalam praktek *machine learning* medis untuk meningkatkan akurasi model klasifikasi [22].

Tabel 3. Hasil *Min-Max Normalization*

Usia	Jenis_Kelamin	Tekanan_Darah	Kadar_Glukosa	Insulin	BMI	Hasil_Diagnosa
0.7931	0	0.5714	0.3790	1	0.4658	1
0.1724	0	0.5857	0.1327	0	0.8440	0
0.3620	0	0.5714	0.3126	1	0.5149	1
0.8620	0	0.7142	0.4453	1	0.4594	1
0.3965	0	0.6142	0.1734	0	0.8034	0
...	...	...	...	...	...	...
0.7413	1	0.5	0.5032	1	0.4700	1
0.3103	0	0.5857	0.0364	0	0.7350	0

Tabel 4. Hasil *Z-Score Normalization*

Usia	Jenis_Kelamin	Tekanan_Darah	Kadar_Glukosa	Insulin	BMI	Hasil_Diagnosa
1.7987	0	0.2597	0.1045	1.0156	-0.8826	1
-0.8053	0	0.38617	-0.8967	-0.9845	1.9122	0
-0.0096	0	0.2597	-0.1653	1.0156	-0.5194	1
2.0881	0	1.5236	0.3744	1.0156	-0.9300	1
0.1350	0	0.6389	-0.7313	-0.9845	1.6122	0
...	...	...	...	...	...	...
1.5817	1	-0.3721	0.6095	1.0156	-0.8510	1
-0.2266	0	0.3861	-1.2885	-0.9845	1.1069	0

Penelitian ini menggunakan normalisasi data secara khusus pada model SVM dan LR menggunakan dua pendekatan standar *Z-Score Standarization* dan *Min-Max Scaling*. Pilihan ini didasarkan pada fakta bahwa kedua algoritma tersebut sensitif terhadap skala fitur bahkan perbedaan satuan atau rentang dapat mempengaruhi konvergensi parameter dan performa prediksi [36].

3.2. Implementasi Support Vector Machine (SVM)

Tabel 5. Confusion Matrix Hasil Pengujian Diabetes Mellitus dengan Metode SVM-RFE

Confusion Matrix	Prediksi	
	0	1
Aktual		
0	69	0
1	0	67

Tabel 6. Hasil Evaluasi Model SVM

No.	Hasil Diagnosa	Nilai Keberhasilan	
		SVM	SVM-RFE
1.	0	97.43	100
2.	1	98.27	100
Akurasi Keseluruhan		97.85%	100%

Berdasarkan tabel 5, hasil *confusion matrix* menunjukkan bahwa model SVM-RFE berhasil mengklasifikasikan seluruh data uji dengan benar, tanpa kesalahan klasifikasi. Selanjutnya pada tabel 6 menunjukkan akurasi sebelum seleksi fitur adalah 97,85% menjadi 100% setelah seleksi fitur dengan metode RFE.

3.3. Implementasi Logistic Regression (LR)

Tabel 7. Confusion Matrix Hasil Pengujian Diabetes Mellitus dengan Metode LR-RFE

Confusion Matrix Aktual	Prediksi	
	0	1
0	69	0
1	0	67

Pada tabel 7, model LR dengan penerapan metode RFE berhasil mengklasifikasikan seluruh data uji dengan benar, yang ditunjukkan oleh tidak adanya kesalahan prediksi baik pada kelas *non-diabetes* (0) dan kelas *diabetes* (1). Jumlah 69 data dengan label 0 diprediksi dengan benar, sama dengan data label 1. Hal ini menghasilkan *confusion matrix* yang ideal dengan *True Positive* dan *True Negative* sempurna.

Tabel 8. Hasil Evaluasi Model LR

No.	Hasil Diagnosa	Nilai Keberhasilan	
		LR	LR-RFE
1.	0	98.71	100
2.	1	98.27	100
Akurasi Keseluruhan		98.49%	100%

Selanjutnya, Tabel 8 memperlihatkan performa model setelah dilakukan seleksi fitur menggunakan RFE. Tanpa seleksi fitur, akurasi model LR mencapai 98.49%, dengan nilai keberhasilan untuk kelas 0 sebesar 98.71% dan untuk kelas 1 sebesar 98.27%. Setelah seleksi fitur, seluruh nilai keberhasilan naik menjadi 100%. Hal ini mengindikasikan bahwa RFE mampu meningkatkan model LR dalam mengklasifikasikan data, dan menyederhanakan model dengan memilih fitur-fitur paling relevan.

3.4. Implementasi Random Forest (RF)

Tabel 9. Confusion Matrix Hasil Pengujian Diabetes Mellitus dengan Metode RF-RFE

Confusion Matrix Aktual	Prediksi	
	0	1
0	69	0
1	0	67

Tabel 9 memperlihatkan confusion matrix hasil pengujian model RF setelah dilakukan seleksi fitur menggunakan RFE. Dari hasil pengujian terhadap data uji, model RF-RFE berhasil mengklasifikasikan seluruh data dengan sempurna. Sebanyak 69 data pada kelas 0 dan pada kelas 1 diprediksi dengan benar.

Tabel 10. Hasil Evaluasi Model RF

No.	Hasil Diagnosa	Nilai Keberhasilan	
		RF	RF-RFE
1.	0	100	100
2.	1	100	100
Akurasi Keseluruhan		100%	100%

Pada Tabel 10 memperkuat hasil tersebut dengan menunjukkan nilai keberhasilan model baik sebelum maupun setelah RFE. Sebelum RFE, model RF telah menunjukkan performa sempurna dengan tingkat keberhasilan 100% untuk kedua kelas 0 dan 1. Setelah dilakukan RFE, hasil yang diperoleh tetap konsisten, menunjukkan bahwa pemilihan fitur yang dilakukan tidak mengurangi akurasi model, bahkan dapat menyederhanakan kompleksitas model tanpa kehilangan performa. Capaian ini menunjukkan bahwa RF memiliki ketahanan yang sangat baik terhadap fitur-fitur yang tidak relevan (noise), dan tetap mampu mempertahankan akurasi tinggi bahkan setelah dilakukan pengurangan fitur. Performa baik yang ditunjukkan oleh RF dengan RFE maupun tanpa RFE mencerminkan kekuatan model ensemble dalam menangani data klasifikasi biner, khususnya dalam konteks diagnosis *diabetes mellitus*.

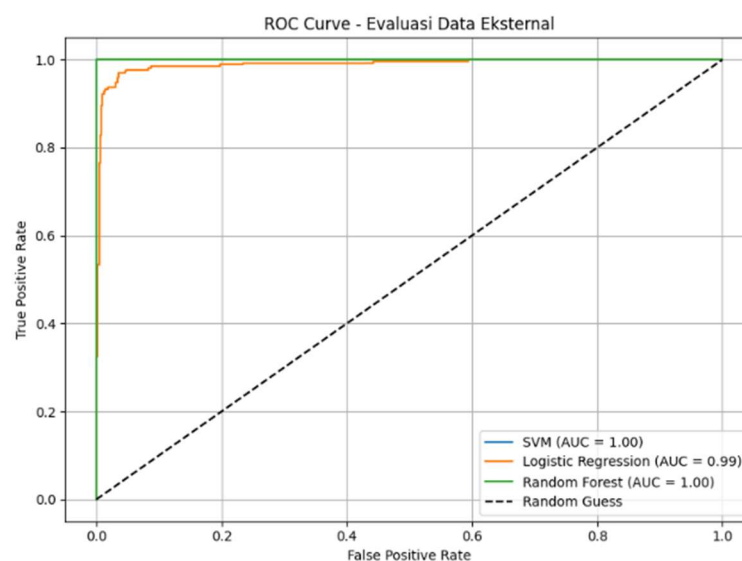
Tabel 11. Perbandingan Akurasi SVM, Logistic Regression dan Random Forest

Matrik Evaluasi	Metode								
	SVM	SVM-RFE	Fitur Terpilih SVM-RFE	LR	LR-RFE	Fitur Terpilih LR-RFE	RF	RF-RFE	Fitur Terpilih RF-RFE
Akurasi	97%	100%	Jenis_Kelamin	98%	100%	Usia	100%	100%	Usia
Presisi	96%	100%	Tekanan_Darah	98%	100%	Jenis_Kelamin	100%	100%	Kadar_Glukosa
Recall	97%	100%	Insulin	98%	100%	Insulin	100%	100%	Insulin
F1-Score	96%	100%	BMI	98%	100%	BMI	100%	100%	BMI

Tabel 11 menyajikan hasil evaluasi performa dari tiga algoritma klasifikasi SVM, RF dan LR, baik sebelum dan setelah penerapan metode seleksi fitur RFE. Semua model menunjukkan peningkatan performa setelah dilakukan RFE. SVM meningkat dari 97% ke 100% akurasi, LR meningkat dari 98% ke 100% dan RF konsisten di angka 100%, baik sebelum dan setelah penerapan RFE. Pada metode RF hasil menunjukkan sama sebelum dan setelah penerapan RFE hal ini tetap berguna walaupun hasil akurasi sama. RFE mampu menyederhanakan model, seperti terlihat pada fitur terpilih (Usia, Kadar_Glukosa, Insulin dan BMI) yang lebih ringkas dan interpretatif. Pada tabel 11 juga terlihat perbedaan yang menggambarkan bahwa algoritma yang berbeda memiliki strategi dan sensitivitas berbeda dalam memilih fitur yang optimal, seperti RF menilai Kadar_Glukosa lebih penting dibanding Jenis_Kelamin, sedangkan SVM menganggap Jenis_kelamin penting untuk klasifikasi.

3.5. Pengujian dengan data eksternal

Evaluasi model menggunakan dataset eksternal dilakukan untuk menguji generalisasi model terhadap data yang belum pernah dilihat sebelumnya. Berikut hasil pengujian dengan data eksternal dapat dilihat pada Gambar 2 dan Tabel 12.



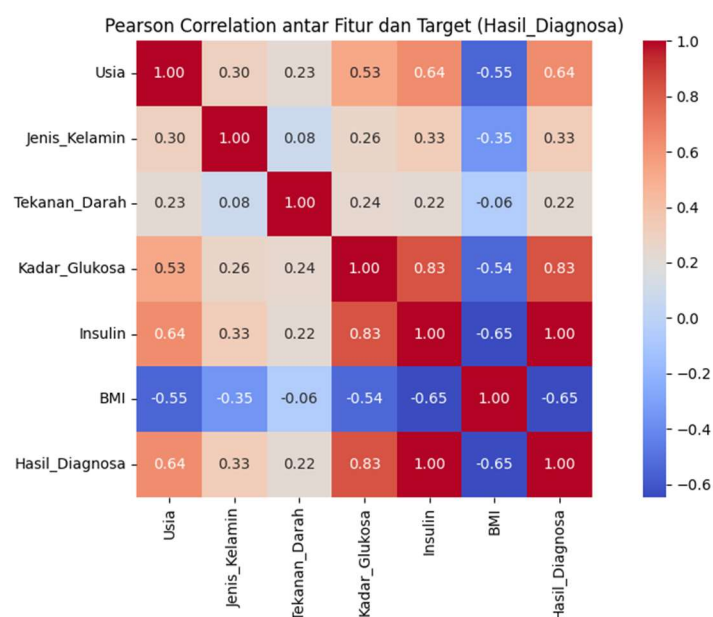
Gambar 2. Kurva ROC-AUC

Analisis *Receiver Operating Characteristic* (ROC) digunakan dalam penelitian ini untuk mengevaluasi kemampuan model dalam membedakan antara dua kelas, yaitu penderita diabetes (positif) dan non-dibetes (negatif). Gambar 1 ROC menunjukkan performa dari tiga model klasifikasi yang digunakan, yaitu SVM, RF dan LR. Berdasarkan hasil yang diperoleh, dua model menunjukkan nilai *Area Under the Curve* (AUC) sebesar 1.00 untuk SVM dan RF sedangkan 0.99 untuk LR, yang mengindikasikan performa prediksi baik.

Tabel 12. Perbandingan Akurasi SVM, *Logistic Regression* dan *Random Forest* dengan data eksternal PIMA Indians

Matrik Evaluasi	Metode		
	SVM	RF	LR
Akurasi	100%	100%	77%
Presisi	100%	100%	98%
Recall	100%	100%	36%
F1-Score	100%	100%	53%
ROC-AUC	100%	100%	98%

Berdasarkan hasil pengujian dengan data eksternal menggunakan data dari PIMA Indians dengan jumlah data 768 menunjukkan hasil akurasi 100% untuk SVM, 100% RF dan 77% untuk LR. Hal ini menunjukkan bahwa baik SVM dan RF mampu mengklasifikasikan data eksternal dengan baik. Sebaliknya, kinerja LR berada di bawah dua model lainnya. Hasil menunjukkan akurasi sebesar 77% dan ROC-AUC 98%, namun model ini menunjukkan recall 36% dan F1-score sebesar 53%. Nilai recall rendah menunjukkan bahwa LR gagal mendeteksi sebagian besar kasus positif (diabetes), walaupun nilai presisi tinggi 98%. Hal ini mengindikasikan kecenderungan model untuk sering memprediksi sebagai negatif, sehingga berdampak pada F1-score yang relatif rendah. Perbedaan kinerja ini memberikan gambaran bahwa model RF dan SVM lebih adaptif terhadap kompleksitas data dan relasi non-linear antar fitur dibandingkan model LR.



Gambar 3. Korelasi antar Fitur

Analisis korelasi Pearson dilakukan untuk mengevaluasi kekuatan dan arah hubungan linear antar fitur-fitur dalam dataset dan variabel target Hasil_Diagnosa. Hasil korelasi menunjukkan pola yang bermakna dan relevan klinis seperti fitur Kadar_Glukosa ($r=0.83$) yang memiliki korelasi positif kuat. Hal ini menunjukkan bahwa peningkatan kadar glukosa berasosiasi dengan peningkatan kemungkinan diabetes terdiagnosis. Hal ini sejalan dengan observasi klinis bahwa hiperglikemia adalah indikator utama dalam diagnosis diabetes tipe 2 [37]. Selain itu pada fitur usia ($r=0.64$) memiliki korelasi positif sedang hingga kuat, menunjukkan bahwa tipe diabetes terdiagnosis lebih sering ditemukan pada populasi lanjut usia. Secara fisiologis, hal ini dijelaskan oleh kemunduran sensitivitas insulin serta peningkatan resistensi insulin seiring bertambahnya usia [38].

4. Pembahasan

4.1. Evaluasi Model Support Vector Machine (SVM)

Model SVM sebelum penggunaan RFE menghasilkan akurasi 97%, presisi 96%, recall 97% dan F1-score 96%. Nilai-nilai ini menunjukkan performa prediksi yang sangat baik, dengan keseimbangan antara kemampuan model dalam mengenali kasus diabetes (recall tinggi) dan ketepatan prediksi (presisi tinggi). Setelah penerapan RFE, semua matrik evaluasi meningkat menjadi 100%, yang menandakan peningkatan signifikan. Hal ini menunjukkan bahwa penghapusan fitur yang tidak relevan atau reduksi secara sistematis melalui RFE dapat meningkatkan kinerja model secara drastis dengan menyederhanakan kompleksitas model tanpa mengurangi informasi penting [39].

4.2. Evaluasi Model Logistic Regression (LR)

Model *Logistic Regression* mencatatkan performa tinggi sebelum RFE akurasi, presisi, recall dan F1-score masing-masing sebesar 98%. Ini menunjukkan bahwa LR merupakan model yang handal dan stabil untuk klasifikasi biner seperti prediksi diabetes. Setelah RFE diterapkan, seluruh matrik evaluasi meningkat menjadi 100%, menunjukkan bahwa RFE mampu meningkatkan efisiensi model bahkan pada model yang sebelumnya telah optimal. Hal ini memperkuat temuan bahwa reduksi fitur berdampak positif terhadap stabilitas dan generalisasi model klasifikasi linier dalam data kesehatan.

4.3. Evaluasi Model Random Forest (RF)

Model *Random Forest* menunjukkan performa sempurna 100% baik sebelum maupun sesudah RFE. Hal ini mencerminkan keunggulan model ensemble dalam menangani kompleksitas dan ketidakseimbangan data. Meskipun RFE tidak meningkatkan metrik lebih lanjut, penerapannya tetap relevan untuk meningkatkan efisiensi komputasi dan mengurangi dimensi data.

Meskipun hasil menunjukkan kinerja model yang sangat baik, namun ada beberapa keterbatasan yang perlu dicermati. Pertama, ukuran sample kecil atau terbatas dapat menyebabkan bias sehingga data yang semakin banyak memungkinkan model yang lebih stabil dan realistis. Kedua, penambahan fitur klinis lanjutan seperti *HbA1c* (*Hemoglobin A1c*), riwayat keluarga, aktivitas fisik dan pola makan yang dapat memperkaya model dan meningkatkan akurasi diagnosis dini. Keterbatasan ini menjadi peluang untuk pengembangan lebih lanjut di masa depan.

5. Kesimpulan

Penelitian ini berhasil menunjukkan bahwa penerapan teknik *Recursive Feature Elimination* (RFE) secara signifikan meningkatkan kinerja model klasifikasi dalam mendeteksi dini risiko *diabetes mellitus*. Tiga algoritma pembelajaran mesin yang digunakan, yaitu *Support Vector Machine* (SVM), *Logistic Regression* (LR), dan *Random Forest* (RF), menunjukkan peningkatan performa setelah dioptimasi dengan RFE, dengan seluruh matrik evaluasi (akurasi, presisi, *recall* dan *F1-score*) mencapai 100%. Kesamaan pemilihan fitur insulin dan BMI pada ketiga model menegaskan pentingnya kedua variabel tersebut sebagai indikator kuat dalam deteksi awal *diabetes mellitus*. Secara keseluruhan, hasil penelitian ini menegaskan pentingnya seleksi fitur RFE tidak sebatas dalam peningkatan akurasi dan efisiensi prediksi, tetapi juga penyederhanaan kompleksitas model, yang sangat berguna untuk implementasi di lingkungan klinis dengan sumber daya terbatas.

Ucapan Terima Kasih: Penulis mengucapkan puji syukur kepada Allah SWT atas berkah dan karunia-Nya, sehingga penelitian ini dapat terlaksana. Tidak lupa juga penulis mengucapkan terima kasih sebesar-besarnya kepada Universitas Muhammadiyah Mataram dan LPPM yang telah membantu dalam pembiayaan penelitian ini, serta ucapan terima kasih juga disampaikan kepada RS PKU Muhammadiyah Bima atas bantuan dan kerjasamanya dalam penyediaan data serta fasilitas dalam proses penelitian.

Referensi

- [1] L. Fregoso-Aparicio, J. Noguez, L. Montesinos, and J. A. García-García, "Machine learning and deep learning predictive models for type 2 diabetes: a systematic review," *Diabetol Metab Syndr*, vol. 13, no. 1, Dec. 2021, <https://doi.org/10.1186/s13098-021-00767-9>.
- [2] O. O. Oladimeji, A. Oladimeji, and O. Oladimeji, "Classification models for likelihood prediction of diabetes at early stage using feature selection," *Applied Computing and Informatics*, vol. 20, no. 3–4, pp. 279–286, Jun. 2024, <https://doi.org/10.1108/ACI-01-2021-0022>.
- [3] M. J. Uddin *et al.*, "A Comparison of Machine Learning Techniques for the Detection of Type-2 Diabetes Mellitus: Experiences from Bangladesh," *Information (Switzerland)*, vol. 14, no. 7, Jul. 2023, <https://doi.org/10.3390/info14070376>.
- [4] E. Sreehari and L. D. D. Babu, "Critical Factor Analysis for prediction of Diabetes Mellitus using an Inclusive Feature Selection Strategy," *Applied Artificial Intelligence*, vol. 38, no. 1, 2024, <https://doi.org/10.1080/08839514.2024.2331919>.
- [5] I. J. Kakoly, M. R. Hoque, and N. Hasan, "Data-Driven Diabetes Risk Factor Prediction Using Machine Learning Algorithms with Feature Selection Technique," *Sustainability (Switzerland)*, vol. 15, no. 6, Mar. 2023, <https://doi.org/10.3390/su15064930>.
- [6] Kementerian Kesehatan Republik Indonesia, "PERATURAN MENTERI KESEHATAN REPUBLIK INDONESIA No.30 Tahun 2013," 2013. Accessed: Sep. 30, 2024. <https://peraturan.bpk.go.id/Details/172111/permenkes-no-30-tahun-2013>

- [7] A. Z. Arrayyan, H. Setiawan, and K. T. Putra, "Naive Bayes for Diabetes Prediction: Developing a Classification Model for Risk Identification in Specific Populations," *Semesta Teknika*, vol. 27, no. 1, pp. 28–36, Apr. 2024, <https://doi.org/10.18196/st.v27i1.21008>.
- [8] Dinas Kesehatan Provinsi NTB, "Cakupan Pelayanan Kesehatan Penderita Diabetes Melitus Prov NTB - SMT I 2023," 2023. Accessed: Sep. 30, 2024. <https://data.ntbprov.go.id/dataset/pelayanan-kesehatan-penderita-diabetes-melitus-dm-di-provinsi-ntb>
- [9] E. Sabitha and M. Durgadevi, "Improving the Diabetes Diagnosis Prediction Rate Using Data Preprocessing, Data Augmentation and Recursive Feature Elimination Method," *IJACSA International Journal of Advanced Computer Science and Applications*, vol. 13, no. 9, 2022, <https://dx.doi.org/10.14569/IJACSA.2022.01309107>
- [10] H. Jeon and S. Oh, "Hybrid-recursive feature elimination for efficient feature selection," *Applied Sciences (Switzerland)*, vol. 10, no. 9, May 2020, <https://doi.org/10.3390/app10093211>.
- [11] Y. Han, L. Huang, and F. Zhou, "A dynamic recursive feature elimination framework (dRFE) to further refine a set ofOMIC biomarkers," *Bioinformatics*, vol. 37, no. 15, pp. 2183–2189, Aug. 2021, <https://doi.org/10.1093/bioinformatics/btab055>.
- [12] S. Matharaarachchi, M. Domaratzki, and S. Muthukumarana, "Minimizing features while maintaining performance in data classification problems," *PeerJ Comput Sci*, vol. 8, 2022, <https://doi.org/10.7717/PEERJ-CS.1081>.
- [13] M. DEMİR and İ. KILIÇ, "An Application of Feature Selection Methods to Compare the Performances of Classification Algorithms," *Afyon Kocatepe University Journal of Sciences and Engineering*, vol. 22, no. 6, pp. 1307–1313, Dec. 2022, <https://doi.org/10.35414/akufemubid.1153610>.
- [14] E. Helmud, E. Helmud, F. Fitriyani, and P. Romadiana, "Classification Comparison Performance of Supervised Machine Learning Random Forest and Decision Tree Algorithms Using Confusion Matrix," *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, vol. 13, no. 1, pp. 92–97, Feb. 2024, <https://doi.org/10.32736/sisfokom.v13i1.1985>.
- [15] A. Theissler, M. Thomas, M. Burch, and F. Gerschner, "ConfusionVis: Comparative evaluation and selection of multi-class classifiers based on confusion matrices," *Knowl Based Syst*, vol. 247, p. 108651, Jul. 2022, <https://doi.org/10.1016/J.KNOSYS.2022.108651>.
- [16] A. Deyol, . A., O. Shandilya, and Y. Nayak, "Comparative Study of Different Machine Learning Classifiers Using Multiple Feature Selection Techniques for Breast Cancer Classification," *Int J Res Appl Sci Eng Technol*, vol. 10, no. 12, pp. 614–620, Dec. 2022, <https://doi.org/10.22214/ijraset.2022.47953>.
- [17] M. Zhang, "Research on Diabetes Risk Prediction Using Multiple Machine Learning Models," *Transactions on Materials, Biotechnology and Life Sciences*, vol. 5, 2024.
- [18] U. Das and B. Ahmed, "An Explainable AI-based Machine Learning Approach for Predicting Diabetes in the Early Stage Using the Influential Features," Jun. 2024, <https://doi.org/10.20944/preprints202406.0364.v1>.
- [19] S. A. Hamzah, "Association between Lipid Profiles and Renal Functions among Adults with Type 2 Diabetes," *Dubai Diabetes and Endocrinology Journal*, vol. 25, no. 3–4, pp. 134–138, 2019, <https://doi.org/10.1159/000502005>.
- [20] B. Sutara, F. Vulture, and R. Novianti, "Application of K-Means algorithm with CRISP-DM method in student data analysis as a support for promotion strategy SIDE: Scientiftic Development Journal," *SIDE: Scientiftic Development Journal*, vol. 1, no. 1, 2024, <https://ojs.arbain.co.id/index.php/side/article/view/6>
- [21] N. Solomon, Y. Lokhnygina, and S. Halabi, "Comparison of regression imputation methods of baseline covariates that predict survival outcomes," *J Clin Transl Sci*, vol. 5, no. 1, 2021, <https://doi.org/10.1017/cts.2020.533>.
- [22] I. Saputra, *Belajar Mudah Data Mining Untuk Pemula*. Bandung: Informatika Bandung, 2023.
- [23] M. A. Usni Zamzami Hasibuan dan Palmizal and M. Usni Zamzami Hasibuan, "Sosialisasi Penerapan Indeks Massa Tubuh (IMT) di Suta Club," *Jurnal Cerdas Sifa Pendidikan*, vol. 10, no.2, pp. 84–89, 2021, <https://doi.org/10.22437/csp.v10i2.15585>
- [24] N. Gray, G. Picone, F. Sloan, and A. Yashkin, "Relation between BMI and diabetes mellitus and its complications among US older adults," *South Med J*, vol. 108, no. 1, pp. 29–36, 2015, <https://doi.org/10.14423/SMJ.0000000000000214>.
- [25] A. S. Sandhu, G. Sharma, and P. K. Mishra, "Diabetes Prediction Using Machine Learning: A Comparative Analyses of Classification Algorithms," in *2024 International Conference on Signal Processing and Advance Research in Computing (SPARC)*, 2024, pp. 1–6. <https://doi.org/10.1109/SPARC61891.2024.10828692>.
- [26] D. Wilimitis and C. G. Walsh, "Practical Considerations and Applied Examples of Cross-Validation for Model Development and Evaluation in Health Care: Tutorial," *JMIR AI*, vol. 2, p. e49023, Dec. 2023, <https://doi.org/10.2196/49023>.
- [27] M. Rose and H. R. Hassen, "A Survey of Random Forest Pruning Techniques," *Academy and Industry Research Collaboration Center (AIRCC)*, Dec. 2019, pp. 99–109. <https://doi.org/10.5121/csit.2019.91808>.
- [28] A. Tsigler and P. L. Bartlett, "Benign overfitting in ridge regression," 2023. <http://jmlr.org/papers/v24/22-1398.html>.
- [29] A. Arafa, M. Radad, M. Badawy, and N. El - Fishawy, "Regularized Logistic Regression Model for Cancer Classification," in *2021 38th National Radio Science Conference (NRSC)*, IEEE, Jul. 2021, pp. 251–261. <https://doi.org/10.1109/NRSC52299.2021.9509831>.
- [30] D. Gunawan and H. Setiawan, "Convolutional Neural Network dalam Analisis Citra Medis," 2022.

- [31] Y. I. Putri and A. Rosidi, "Sensitifitas dan Spesifisitas IMT dan Lingkar Perut Sebagai Indikator Risiko Diabetes Mellitus Sensitivity and Specificity of BMI and Abdominal Circumference as Indicators of Diabetes Risks," *Jurnal Kesehatan Medika Saintika*, vol.9, no.1, 68-77, 2018, <http://dx.doi.org/10.30633/jkms.v9i1.113>
- [32] H. Sanada *et al.*, "High Body Mass Index is an Important Risk Factor for the Development of Type 2 Diabetes NIH Public Access \$watermark-text \$watermark-text \$watermark-text," 2012.
- [33] P.-O. Côté, A. Nikanjam, N. Ahmed, D. Humeniuk, and F. Khomh, "Data Cleaning and Machine Learning: A Systematic Literature Review," Oct. 2023, <http://arxiv.org/abs/2310.01765>
- [34] T. Emmanuel, T. Maupong, D. Mpoeleng, T. Semong, B. Mphago, and O. Tabona, "A survey on missing data in machine learning," *J Big Data*, vol. 8, no. 1, p. 140, 2021, <https://doi.org/10.1186/s40537-021-00516-9>.
- [35] G. S. and L. X. I. and P. N. M. and S. J. L. Ehrig Molly and Bullock, "Imputation and Missing Indicators for Handling Missing Longitudinal Data: Data Simulation Analysis Based on Electronic Health Record Data," *JMIR Med Inform*, vol. 13, p. e64354, Mar. 2025, <https://doi.org/10.2196/64354>.
- [36] X. H. Cao, I. Stojkovic, and Z. Obradovic, "A robust data scaling algorithm to improve classification accuracies in biomedical data," *BMC Bioinformatics*, vol. 17, no. 1, Sep. 2016, <https://doi.org/10.1186/s12859-016-1236-x>.
- [37] D. K. Hestiani and A. Mappa Oudang, "Tentang Waspada Hipertensi, Hiperglikemia, Hiperurisemia, Dan Hiperkolesterolemia Health Services: Health Check-Ups & Health Education On Hypertension, Hyperglycaemia, Hyperuricaemia, And Hypercholesterolemia Awareness," *Jurnal PEDAMAS (Pengabdian Kepada Masyarakat)*, vol. 2, no. 5, 1362-1371, 2024, <https://pekatpkm.my.id/index.php/JP/article/view/430>.
- [38] S. Supadmi *et al.*, *Sindrom Metabolik Pada Lanjut Usia (Lansia)*. 2024.
- [39] B. Satria, T. A. Y. Siswa, and W. J. Pranoto, "Optimasi Random Forest dengan Genetic Algorithm dan Recursive Feature Elimination pada High Dimensional Data Stunting Samarinda," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 8, no. 3, p. 1778, Jul. 2024, <https://doi.org/10.30865/mib.v8i3.7883>.