

Pemanfaatan Model Linier dalam Klasifikasi Penyakit Diabetes Berbasis Machine Learning

Faris Prasetya Ajisaputra *, Wahyu Aji Eko Prabowo

Fakultas Ilmu Komputer, Teknik Informatika, Universitas Dian Nuswantoro, Semarang, Indonesia

Email: ¹*111202113564@mhs.dinus.ac.id, ²prabowo@dsn.dinus.ac.id

Email Penulis Korespondensi: 111202113564@mhs.dinus.ac.id

Submitted 09-05-2025; Accepted 25-06-2025; Published 30-06-2025

Abstrak

Diabetes merupakan penyakit kronis yang dapat menyebabkan komplikasi kesehatan serius apabila tidak dideteksi dan ditangani sejak dini. Deteksi dini memegang peran penting dalam meminimalkan risiko jangka panjang. Penelitian ini bertujuan untuk mengklasifikasikan kasus diabetes dengan menggunakan pendekatan machine learning berbasis model linier. Model yang digunakan dalam penelitian ini meliputi logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier. Dataset telah melalui proses prapemrosesan untuk memastikan kualitas dan konsistensi data. Kinerja setiap model dievaluasi menggunakan lima metrik: akurasi, presisi, recall, F1-score, dan AUC-ROC. Hasil eksperimen menunjukkan bahwa ridge classifier memberikan performa terbaik, diikuti oleh LDA dan SVM linier dengan hasil yang sebanding. Logistic regression juga menunjukkan kinerja yang cukup baik meskipun sedikit lebih rendah. Temuan ini menunjukkan bahwa model linier mampu memberikan klasifikasi yang akurat dan handal dalam tugas prediksi diabetes serta berkontribusi dalam membuktikan bahwa model ini dapat digunakan sebagai dasar sistem pendukung keputusan untuk diagnosis dini diabetes dalam bidang kesehatan.

Kata Kunci: Diabetes; Klasifikasi; Model Linier; Pembelajaran Mesin; Evaluasi Kinerja

Abstract

Diabetes is a chronic disease that may lead to serious health complications if not detected and treated early. Early detection plays a crucial role in minimizing long-term risks. This study aims to classify diabetes cases using a machine-learning approach based on linear models. The models applied in this research include logistic regression, linear discriminant analysis (LDA), ridge classifier, and support vector machine (SVM) with a linear kernel. We preprocessed the dataset to ensure quality and consistency. We evaluated each model's performance using accuracy, precision, recall, F1-score, and AUC-ROC. Experimental results show that the ridge classifier achieved the highest performance, followed by LDA and linear SVM, with comparable results. Logistic regression also performed reasonably well, albeit with slightly lower metrics. These findings indicate that the linear model can provide accurate and reliable classification in the task of predicting diabetes, contributing to the proof that this model can serve as the basis for a decision support system for early diabetes diagnosis in the healthcare sector.

Keywords: Diabetes; Classification; Linear Model; Machine Learning; Performance Evaluation

1. PENDAHULUAN

Diabetes melitus merupakan penyakit metabolik kronis yang terjadi akibat ketidakmampuan tubuh dalam memproduksi atau memanfaatkan insulin secara efektif, sehingga kadar glukosa dalam darah tidak terkontrol [1]. Penyakit ini dapat menyerang siapa saja, tanpa memandang usia, mulai dari anak-anak hingga lanjut usia [2]. Berdasarkan data dari Federasi Diabetes Internasional tahun 2019, jumlah penderita diabetes di seluruh dunia telah mencapai 463 juta orang dan diperkirakan akan meningkat menjadi 700 juta pada tahun 2045. Di Indonesia, diabetes termasuk dalam kategori penyakit tidak menular yang paling umum, dengan prevalensi sebesar 6,9% pada tahun 2020, yang berarti lebih dari 16 juta penduduk terpengaruh [3]. Sistem kesehatan menghadapi tantangan besar dalam menangani penyakit ini karena tingginya tingkat morbiditas, mortalitas, serta biaya penanganan komplikasi jangka panjang. Peningkatan prevalensi ini sebagian besar disebabkan oleh gaya hidup tidak sehat, seperti pola makan tinggi gula, garam, dan lemak, serta kurangnya aktivitas fisik di masyarakat.

Dengan prevalensi yang terus meningkat, diabetes melitus telah menjadi salah satu masalah kesehatan masyarakat yang mendesak. Jumlah penderita terus mengalami peningkatan setiap tahunnya [4]. *World Health Organization* (WHO) memperkirakan bahwa jumlah penderita diabetes di Indonesia akan mencapai 21,3 juta jiwa pada tahun 2030 [5]. Salah satu faktor yang memicu peningkatan ini adalah keterlambatan dalam proses diagnosis. Jika tidak ditangani secara tepat, diabetes dapat menyebabkan komplikasi serius yang berpotensi mengancam jiwa [6]. Oleh karena itu, deteksi dini menjadi aspek yang sangat penting dalam upaya pengendalian penyakit ini. Pemanfaatan teknologi, khususnya machine learning, menjadi solusi yang potensial untuk mempercepat dan meningkatkan akurasi dalam proses deteksi dini diabetes melitus.

Penelitian ini menggunakan pendekatan machine learning berbasis model linier, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier. Pemilihan model linier didasarkan pada beberapa keunggulan, antara lain efisiensi komputasi, kemudahan interpretasi hasil, serta performa yang baik pada data berdimensi rendah hingga sedang. Logistic regression dikenal sebagai metode klasifikasi berbasis probabilitas yang efektif dalam berbagai kondisi data [7]. Linear discriminant analysis (LDA) unggul dalam melakukan proyeksi linier untuk memaksimalkan separasi antar kelas dan meminimalkan variasi dalam kelas [8]. Ridge classifier mampu mengatasi masalah multikolinearitas dan mencegah *overfitting* melalui regularisasi L2. Sementara itu, SVM

dengan kernel linier bekerja dengan membentuk *hyperplane* yang memiliki margin maksimum, sehingga mampu memisahkan kelas secara optimal, bahkan pada data yang cukup kompleks [9].

Evaluasi kinerja model dilakukan menggunakan confusion matrix dan lima metrik utama, yaitu: *accuracy*, *precision*, *recall*, *F1-score*, dan *AUC-ROC*. *Accuracy* mengukur tingkat keseluruhan akurasi prediksi, *precision* menilai ketepatan prediksi positif, sedangkan *recall* mengukur kemampuan model dalam mendeteksi kasus positif. *F1-score* merupakan rata-rata harmonik antara *precision* dan *recall*, yang memberikan gambaran menyeluruh mengenai keseimbangan antara keduanya [10]. Selain itu, analisis juga dilakukan melalui kurva ROC dan nilai AUC, di mana ROC membandingkan *true positive rate* (TPR) dan *false positive rate* (FPR), sementara AUC mengindikasikan kemampuan model dalam membedakan antar kelas. Nilai AUC yang mendekati 1 menunjukkan bahwa model memiliki performa klasifikasi yang sangat baik [11].

Beberapa studi sebelumnya mengevaluasi algoritma pembelajaran mesin pada kasus klasifikasi penyakit diabetes. Penelitian yang dilakukan oleh Kurniadi et al. (2021) membandingkan metode K-nearest neighbor (KNN) dan logistic regression dalam mengklasifikasikan penyakit diabetes dengan menggunakan dataset PIMA yang berasal dari UCI Repository. Algoritma KNN menghasilkan akurasi tertinggi, yaitu 85,06%, sedangkan logistic regression hanya mencatat 77,92%. Penelitian ini juga mengoptimalkan model menggunakan *Grid Search* dan melakukan validasi melalui 10 fold *cross-validation*. Meskipun metodenya cukup solid, penelitian ini terbatas pada dua algoritma dasar dan tidak mencakup metode berbasis *neural network*. Di samping itu, data yang digunakan lebih fokus pada diagnosis klinis daripada pada tahap risiko awal diabetes [12]. Penelitian yang sedang berlangsung saat ini berbeda karena membandingkan neural network dan LDA dalam konteks deteksi dini, menggunakan dataset terkait risiko awal, serta menyediakan pendekatan evaluasi yang lebih menyeluruh.

Penelitian lain yang dilakukan oleh Cahyani et al. (2019) menerapkan algoritma support vector machine (SVM) untuk klasifikasi pasien diabetes berdasarkan informasi klinis yang diperoleh dari Puskesmas Modopuro, Mojokerto. Dalam studi ini, tiga tipe kernel diuji, yaitu linear, polynomial, dan sigmoid, dengan akurasi tertinggi mencapai 64% saat menggunakan kernel polynomial. Evaluasi dilaksanakan dengan metode 5 fold *cross-validation*. Namun, penelitian ini terbatas pada penggunaan satu algoritma saja, yaitu SVM, tanpa melakukan perbandingan dengan metode lain, serta tidak memanfaatkan dataset umum yang dapat ditemukan di UCI Repository [13]. Berbeda dengan penelitian sebelumnya, penelitian ini menawarkan pendekatan yang lebih luas dalam membandingkan algoritma dan menggunakan dataset yang khusus dibuat untuk mendeteksi risiko awal terjadinya penyakit diabetes.

Penelitian yang dilakukan oleh Hana et al. (2020) di sisi lain mengevaluasi dua metode klasifikasi, yaitu neural network dan linear discriminant analysis (LDA), untuk mendeteksi diabetes dengan menggunakan dataset early stage diabetes risk prediction dari UCI repository. Temuan menunjukkan bahwa neural network mencapai tingkat akurasi 95,19%, yang lebih baik daripada LDA yang hanya meraih 90,38%. Meskipun penting dan menjadi dasar yang signifikan, penelitian ini tidak membahas secara mendalam proses validasi model maupun dampak distribusi data terhadap kinerja algoritma [14]. Penelitian ini melanjutkan dan memperdalam kajian tersebut dengan pendekatan evaluasi yang lebih sistematis yang melibatkan lima metrik utama dan analisis AUC-ROC, untuk mendapatkan pemahaman yang lebih komprehensif tentang kinerja masing-masing algoritma.

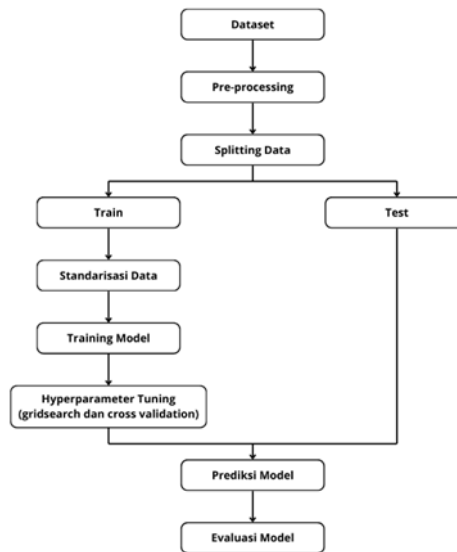
Berdasarkan latar belakang tersebut, penelitian ini bertujuan untuk menerapkan dan membandingkan kinerja beberapa algoritma model linier dalam klasifikasi diabetes berdasarkan data medis, yaitu: logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier. Dataset yang digunakan adalah Pima Indians Diabetes Database yang diambil dari platform Kaggle [15]. Berbeda dengan penelitian sebelumnya yang biasanya hanya membandingkan dua algoritma dasar atau menggunakan pendekatan tunggal, studi ini mencakup perbandingan menyeluruh di antara beberapa model linier sekaligus. Evaluasi dilakukan dengan menggunakan lima metrik utama, yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC AUC*, untuk memberikan gambaran menyeluruh tentang performa klasifikasi. Dengan pendekatan ini, diharapkan penelitian dapat menjembatani kekurangan dari studi studi sebelumnya yang belum menguji berbagai algoritma linier secara bersamaan dan belum menerapkan evaluasi multi metrik pada dataset yang sama, sehingga dapat mendukung keputusan medis yang lebih akurat dan efisien dalam deteksi dini penyakit diabetes.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini mengadopsi pendekatan berbasis model linier, yang dirancang secara sistematis mulai dari tahap pengumpulan data hingga evaluasi kinerja model. Proses pengembangan model machine learning secara umum mengikuti alur yang sistematis untuk memastikan hasil prediksi yang optimal dan valid. Gambar 1 menyajikan sebuah diagram alir mengenai langkah-langkah dalam pengembangan model machine learning yang terdiri dari beberapa langkah utama. Langkah pertama adalah pengumpulan data, dilanjutkan dengan tahap pra-pemrosesan guna membersihkan dan mempersiapkan data. Kemudian, data dibagi menjadi dua segmen, yaitu data untuk pelatihan (*train*) dan data untuk pengujian (*test*). Pada segmen data pelatihan, dilakukan proses standarisasi data, yang kemudian diikuti dengan pelatihan model. Setelah model selesai dilatih, dilakukan penyesuaian parameter model menggunakan teknik *GridSearch* dan *Cross-validation* untuk mendapatkan hasil yang paling baik. Model yang telah disesuaikan tersebut kemudian digunakan untuk melakukan prediksi terhadap data pengujian. Data pengujian dibuat benar-benar terjaga (*unseen*) sehingga

mencegah terjadinya kebocoran data (*data leakage*) dan data pengujian tersebut pada tahap akhir digunakan pada proses evaluasi model. Diagram ini menggambarkan urutan langkah dengan jelas, serta menunjukkan hubungan antar setiap langkah dalam alur pengembangan sistem pembelajaran mesin yang terstruktur.



Gambar 1. Pengembangan model machine learning

2.2 Dataset

Penelitian ini memanfaatkan dataset yang diambil dari platform Kaggle [15], yaitu Pima Indians Diabetes Dataset. Dataset ini banyak digunakan dalam penelitian klasifikasi untuk mendiagnosis diabetes berdasarkan informasi medis dan demografis. Terdiri dari 768 entri (baris) dan 9 fitur (kolom), dataset ini menggambarkan berbagai faktor yang berpengaruh terhadap kemungkinan seseorang menderita diabetes. Beberapa faktor tersebut meliputi kadar glukosa, tekanan darah, indeks massa tubuh, serta riwayat keluarga. Nilai target dalam dataset ini adalah variabel Outcome, yang menunjukkan apakah seseorang terdiagnosis diabetes (1) atau tidak (0). Dalam penelitian ini, terdapat sembilan atribut utama yang dipakai sebagai masukan dalam proses klasifikasi diabetes. Atribut-atribut tersebut mencerminkan informasi klinis pasien dan berasal dari pengukuran medis yang berhubungan dengan potensi diabetes. Rincian lebih lanjut mengenai atribut-atribut tersebut dapat ditemukan dalam Tabel 1 di bawah ini.

Tabel 1. Deskriptor Fitur yang Digunakan Pada Penelitian Ini

No	Fitur	Deskripsi
1	Pregnancies	Jumlah kehamilan pasien
2	Glucose	Konsentrasi glukosa dalam plasma setelah puasa
3	Blood Pressure	Tekanan darah diastolik (mm Hg)
4	Skin Thickness	Ketebalan lipatan kulit trisep (mm)
5	Insulin	Jumlah insulin serum (mu U/ml)
6	BMI	Indeks massa tubuh (kg/m ²)
7	Diabetes Pedigree Function	Fungsi silsilah diabetes, menunjukkan kemungkinan keturunan
8	Age	Usia pasien (dalam tahun)
9	Outcome	Hasil diagnosis (1 = positif diabetes, 0 = negatif diabetes)

Tabel 1 menunjukkan sembilan fitur medis yang berfungsi sebagai variabel input dalam proses klasifikasi diabetes. Fitur-fitur tersebut mencakup *pregnancies* (jumlah kehamilan), *glucose* (tingkat glukosa puasa), *blood pressure* (tekanan darah diastolik), *skin thickness* (ketebalan lipatan kulit di bagian trisep), *insulin* (kadar insulin dalam serum), BMI, *diabetes pedigree function* (fungsi silsilah diabetes), *age* (umur pasien), dan *outcome* (diagnosis: 1 untuk positif, 0 untuk negatif). Semua fitur ini menggambarkan kondisi fisik, latar belakang genetik, serta hasil tes laboratorium pasien yang secara penting untuk mendeteksi dan memprediksi risiko diabetes dengan tepat dan lebih awal.

2.3 Pre-processing

Tahapan ini dilakukan untuk memastikan bahwa data yang digunakan dalam kondisi bersih dan siap untuk pemodelan. Proses ini meliputi beberapa langkah, seperti penghapusan data duplikat, penanganan nilai yang hilang, pengkodean fitur kategorikal, serta normalisasi nilai numerik. Tahap ini sangat krusial, karena kualitas data yang baik akan langsung memengaruhi performa model yang dihasilkan [16].

2.4 Splitting Data

Setelah menyelesaikan proses pra-pemrosesan, dataset dibagi menjadi dua bagian: 80% untuk data latih (*training data*) dan 20% untuk data uji (*testing data*). Pembagian ini dilakukan secara acak, yang bertujuan untuk memastikan hasil yang konsisten dan dapat direproduksi. Langkah ini penting untuk menghindari overfitting serta memastikan bahwa evaluasi kinerja model dilakukan pada data yang belum pernah dilihat sebelumnya [17].

2.5 Standarisasi Data

Standarisasi diterapkan pada semua fitur numerik dalam dataset dengan menggunakan metode *StandardScaler* dari *Scikit-learn*. Proses ini sangat penting karena banyak algoritma linier sangat peka terhadap skala data [18]. Dengan melakukan standarisasi, setiap fitur akan memiliki nilai rata-rata sebesar 0 dan standar deviasi sebesar 1. Hal ini tidak hanya mempercepat proses pelatihan, tetapi juga meningkatkan akurasi model yang dihasilkan.

2.6 Pelatihan Model

Penelitian ini mengadopsi empat jenis model linier, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier. Pemilihan model-model ini didasarkan pada efisiensinya untuk dataset yang terstandarisasi serta kemampuannya yang baik dalam melakukan generalisasi. Proses pelatihan model dilaksanakan pada data pelatihan yang telah melalui tahap pengolahan dan distandarisasi.

2.7 Hyperparameter Tuning

Untuk mencapai performa model yang optimal, dilakukan pencarian parameter terbaik menggunakan *GridSearch cross-validation* dengan pengaturan *cross-validation fold* sebanyak 5. Ini berarti bahwa data latih akan dibagi secara otomatis menjadi lima bagian selama proses validasi silang. Dengan pendekatan ini, model akan diuji secara komprehensif terhadap berbagai kombinasi parameter, sehingga dapat ditemukan konfigurasi yang secara konsisten memberikan performa terbaik.

2.8 Prediksi Model

Setelah model dilatih dan dioptimalkan, versi terbaik dari model tersebut akan digunakan untuk melakukan prediksi terhadap data pengujian. Tujuan dari proses ini adalah untuk menguji sejauh mana kemampuan model dalam menggeneralisasi pola dari data yang sebelumnya belum pernah dihadapi. Hasil prediksi ini nantinya akan menjadi acuan dalam mengevaluasi performa model.

2.9 Evaluasi Model

Dalam penelitian ini, terdapat lima parameter evaluasi utama yang digunakan untuk menilai kinerja model klasifikasi, yaitu accuracy, presision, recall, F1-score, dan AUC-ROC. Akurasi mengukur proporsi prediksi yang benar terhadap seluruh data uji, memberikan gambaran umum mengenai kinerja model dalam klasifikasi data [19]. Presisi berfungsi untuk menilai ketepatan dari prediksi positif yang dihasilkan oleh model, yaitu seberapa besar proporsi prediksi positif yang memang benar-benar akurat [20]. Recall, yang juga dikenal sebagai sensitivitas, mengukur kemampuan model dalam mendeteksi semua kasus positif yang sebenarnya, sehingga sangat penting untuk memastikan bahwa tidak ada data positif yang terlewat [19]. F1-Score merupakan rata-rata harmonis antara presisi dan recall, yang digunakan untuk menyeimbangkan keduanya, terutama ketika data yang digunakan tidak seimbang [20]. Di sisi lain, AUC-ROC mengukur kemampuan model dalam membedakan antara kelas positif dan negatif pada berbagai ambang batas klasifikasi; semakin tinggi nilai AUC, semakin baik pula kemampuan diskriminatif model tersebut [19].

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

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

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

$$F1\ Score = \frac{2(Recall \cdot Precision)}{(Recall + Precision)} \quad (4)$$

$$AUC = \sum_{i=1}^{n-1} (FPR_{i+1} - FPR_i) \left(\frac{TPR_{i+1} + TPR_i}{2} \right) \quad (5)$$

Keterangan :

TP (True Positive) : Jumlah prediksi positif yang benar

TN (True Negative) : Jumlah prediksi negatif yang benar

FP (False Positive) : Jumlah prediksi positif yang salah

FN (False Negative) : Jumlah prediksi negatif yang salah

$FPR = \frac{FP}{FP+TN}$: False Positive Rate (FPR) = Rasio kesalahan prediksi positif.

$TPR = \frac{TP}{TP+TN}$: True Positive Rate (TPR) = Rasio prediksi positif yang benar (sensitivitas).

i : Indeks titik-titik pada ROC Curve.

3. HASIL DAN PEMBAHASAN

Penelitian ini bertujuan untuk investigasi klasifikasi penyakit diabetes dengan menggunakan empat algoritma pembelajaran mesin yang berbasis linear, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linear. Proses dimulai dengan pre-processing data, yang mencakup pembersihan dataset, normalisasi atau transformasi fitur jika diperlukan, serta pemisahan dataset menjadi set data pelatihan dan set data pengujian. Selanjutnya, model dilatih dengan data pelatihan menggunakan parameter bawaan untuk menjaga konsistensi dalam perbandingan kinerja antar algoritma. Evaluasi dilakukan dengan lima metrik utama, yaitu akurasi, presisi, recall, F1-score, dan ROC AUC, untuk memberikan gambaran menyeluruh tentang kinerja model. Setelah tahap pelatihan selesai, model kemudian diuji menggunakan set data pengujian untuk mengevaluasi kemampuan generalisasi dalam mengklasifikasikan data baru yang belum pernah ditemui sebelumnya.

3.1 Hasil Training

Pada fase ini, proses pelatihan dilakukan terhadap empat tipe model yang berbasis linier, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier. Data yang digunakan merupakan informasi mengenai diagnosis penyakit diabetes, dengan tujuan klasifikasi yang terbagi menjadi dua kategori, yaitu positif dan negatif diabetes. Penilaian model dilakukan dengan memanfaatkan lima metrik kinerja, yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC AUC*. Berikut merupakan hasil performa model pada tahap pelatihan berdasarkan beberapa metrik evaluasi, seperti yang ditunjukkan pada Tabel 2.

Tabel 2. Hasil Performa Pada Tahap *Training*

Model	Accuracy	Precision	Recall	F1-Score	ROC AUC
Logistic Regression	0.770358	0.711765	0.568075	0.631854	0.722940
Linear Discriminant Analysis	0.768730	0.712575	0.558685	0.626316	0.719492
Ridge Classifier	0.767101	0.716049	0.544601	0.618667	0.714944
SVM dengan kernel linier	0.770358	0.714286	0.563380	0.629921	0.721840

Tabel 2 menunjukkan evaluasi performa empat model klasifikasi pada tahap pelatihan menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *ROC AUC*. Keempat metrik ini dipilih untuk memberikan gambaran yang komprehensif mengenai kemampuan setiap model dalam mengenali pola pada data latih. Dengan melihat nilai pada tabel, dapat dianalisis kelebihan dan kelemahan masing-masing model sebelum dilakukan pengujian lebih lanjut pada data uji.

Berdasarkan Tabel 2, dapat disimpulkan bahwa performa keempat model linier yang diuji menunjukkan tingkat kompetisi yang cukup signifikan dengan perbedaan nilai evaluasi yang relatif kecil. Logistic regression dan support vector machine (SVM) dengan kernel linier mencatatkan nilai akurasi tertinggi sebesar 77,04%. Hal ini mengindikasikan bahwa kedua model tersebut mampu melakukan klasifikasi dengan benar terhadap 77% dari keseluruhan data yang tersedia. Akurasi yang tinggi ini menjadi sinyal bahwa logistic regression dan SVM memiliki kinerja yang seimbang dalam mendeteksi data positif (penderita diabetes) dan data negatif (bukan penderita diabetes). Dari perspektif *precision*, ridge classifier menduduki peringkat tertinggi dengan nilai 71,60%. Tingginya nilai *precision* menunjukkan bahwa model ini cenderung lebih berhati-hati dalam memberikan prediksi positif, sehingga kesalahan dalam mengklasifikasikan pasien non-diabetes sebagai diabetes (*false positive*) dapat diminimalkan. Namun, perlu dicatat bahwa peningkatan *precision* sering kali berbanding terbalik dengan *recall*, sebagaimana yang terlihat pada ridge classifier, di mana nilai *recall* justru lebih rendah dibandingkan dengan model lainnya. *Recall* merupakan metrik yang sangat penting dalam konteks diagnosis medis, karena tujuan utamanya adalah untuk mengidentifikasi sebanyak mungkin pasien yang benar-benar menderita diabetes.

Dalam analisis metrik ini, model logistic regression menunjukkan performa terbaik dengan nilai *recall* sebesar 56,81%, diikuti oleh support vector machine (SVM) yang memiliki nilai *recall* sebesar 56,34%. Hal ini mengindikasikan bahwa model logistic regression lebih sensitif dalam mendeteksi pasien yang benar-benar positif, meskipun terdapat kemungkinan pengorbanan terhadap nilai *precision*. Dalam konteks diagnosis penyakit, *recall* yang tinggi sering kali diprioritaskan untuk meminimalkan risiko tidak terdeteksinya pasien yang sebenarnya positif. Selanjutnya, *F1-score*, yang merupakan harmonisasi antara *precision* dan *recall*, menunjukkan bahwa model logistic regression (*F1-Score* 0,631854) tetap mempertahankan performa terbaik. *F1-Score* sangat penting dalam situasi di mana keseimbangan antara kesalahan positif dan kesalahan negatif perlu dipertahankan, khususnya dalam kasus penyakit kronis seperti diabetes, yang dapat memiliki konsekuensi serius jika tidak terdiagnosis secara tepat. Selain itu, metrik ROC AUC juga menegaskan keunggulan logistic regression, di mana model ini memperoleh nilai tertinggi sebesar 0,722940. ROC AUC mengukur kemampuan model dalam membedakan antara kelas positif dan negatif secara keseluruhan. Nilai di atas 0,7 menunjukkan bahwa model memiliki performa klasifikasi yang cukup baik dalam berbagai skenario ambang batas (*threshold*).

Secara umum, berdasarkan keseluruhan metrik yang dianalisis, dapat disimpulkan bahwa logistic regression merupakan model yang paling konsisten dan seimbang untuk digunakan dalam klasifikasi penyakit diabetes berdasarkan data ini. Model support vector machine (SVM) juga menunjukkan performa yang serupa dengan logistic regression, menjadikannya sebagai pilihan alternatif yang layak. Sementara itu, ridge classifier menawarkan keunggulan dalam hal *precision* untuk aplikasi yang memerlukan akurasi prediksi positif yang tinggi, meskipun memiliki nilai *recall* yang lebih

rendah. Di sisi lain, linear discriminant analysis (LDA) menunjukkan kinerja yang kompetitif, namun secara konsisten berada sedikit di bawah ketiga model lainnya dalam semua metrik yang dievaluasi.

3.2 Hasil Testing

Untuk mengevaluasi kinerja model yang telah dibangun, kami melakukan pengujian dengan menggunakan berbagai algoritma klasifikasi. Hasil dari pengujian masing-masing model yang digunakan di tampilkan didalam Tabel 3 berikut.

Tabel 3. Hasil Performa Pada Tahap Testing

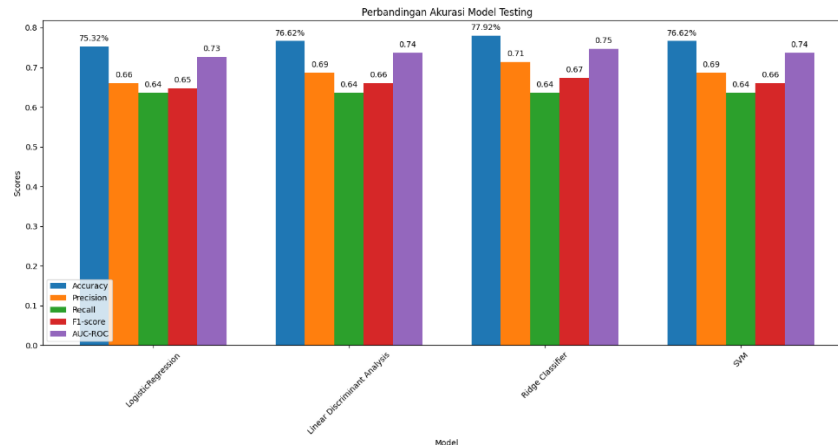
Model	Accuracy	Precision	Recall	F1-Score	ROC AUC
Logistic Regression	0.753247	0.660377	0.636364	0.648148	0.727273
Linear Discriminant Analysis	0.766234	0.686275	0.636364	0.660377	0.737374
Ridge Classifier	0.779221	0.714286	0.636364	0.673077	0.747475
SVM dengan kernel linier	0.766234	0.686275	0.636364	0.660377	0.737374

Tabel 3 menyajikan kinerja empat model klasifikasi pada data uji dengan menggunakan lima metrik evaluasi. Model yang dibandingkan meliputi logistic regression, linear discriminant analysis, ridge classifier, dan SVM dengan kernel linier. Metrik yang digunakan mencakup *accuracy*, *precision*, *recall*, *F1-score*, dan ROC AUC, yang masing-masing mewakili aspek berbeda dalam mengukur performa prediksi model. Data dalam tabel ini digunakan untuk menilai sejauh mana setiap model mampu mempertahankan kualitas klasifikasi saat diterapkan pada data yang belum pernah dilatih sebelumnya.

Berdasarkan evaluasi yang tertuang dalam Tabel 3, terlihat bahwa performa keempat model yang diuji, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier, menunjukkan tingkat persaingan yang cukup ketat. Metrik evaluasi yang dihasilkan oleh masing-masing model memiliki nilai yang cukup berdekatan, dengan *accuracy* berada di kisaran 75% hingga 78%. Model logistic regression tercatat menghasilkan *accuracy* sebesar 75,32%, dengan *precision* sebesar 66,03%, *recall* sebesar 63,64%, *F1-Score* sebesar 64,81%, dan ROC AUC sebesar 0,727273. Dari keseluruhan model yang diuji, logistic regression menunjukkan performa terendah dalam hal ketepatan, sensitivitas, dan kemampuan membedakan kelas. Hal ini menandakan bahwa meskipun logistic regression merupakan model yang sederhana dan cepat dalam proses pelatihannya, kinerjanya dalam melakukan klasifikasi pada dataset ini belum optimal jika dibandingkan dengan model lainnya. Linear discriminant analysis (LDA) menunjukkan hasil yang menjanjikan dengan *accuracy* mencapai 76,62%, *precision* sebesar 68,63%, *recall* sebesar 63,64%, *F1-score* 66,04%, dan ROC AUC 0,737374. LDA mampu menyajikan klasifikasi yang lebih baik karena metodologinya mencari kombinasi linier yang memisahkan kelas-kelas dengan optimal, terutama ketika distribusi data cenderung normal dan homogen antar kelas. Dalam konteks dataset ini, LDA membuktikan bahwa dengan memanfaatkan asumsi distribusi data, model dapat memberikan prediksi yang lebih seimbang antara *precision* dan *recall*.

Sementara itu, model ridge classifier menonjol sebagai yang terbaik di antara semua model yang diuji. Dengan *accuracy* sebesar 77,92%, *precision* sebesar 71,43%, *recall* sebesar 63,64%, *F1-score* sebesar 67,31%, dan ROC AUC sebesar 0,747475. Ridge classifier tidak hanya unggul dalam prediksi keseluruhan, tetapi juga lebih konsisten dalam menjaga keseimbangan antara ketepatan dan sensitivitas. Ridge classifier menerapkan regularisasi L2, yang efektif mengurangi overfitting dengan membatasi nilai koefisien agar tetap dalam rentang yang wajar. Pendekatan ini membuat model menjadi lebih sederhana namun tetap akurat, terutama saat menghadapi variasi data baru yang mungkin berbeda dari data pelatihan. Support vector machine (SVM) dengan kernel linier menunjukkan hasil yang sebanding dengan linear discriminant analysis (LDA), dengan *accuracy* mencapai 76,62%, *precision* sebesar 68,63%, *recall* 63,64%, *F1-score* sebesar 66,04%, dan ROC AUC 0,737374. SVM linier dirancang untuk menemukan *hyperplane* optimal yang memisahkan kelas data dengan margin maksimum. Pendekatan ini sangat sesuai digunakan untuk data yang dapat dipisahkan secara linier, yang sejalan dengan hasil evaluasi yang menunjukkan performa yang setara dengan LDA. Keduanya beroperasi dengan asumsi bahwa data dapat dibedakan dengan batas linier.

Dari segi metrik evaluasi, ridge classifier muncul sebagai model terbaik, dengan nilai *accuracy*, *precision*, *F1-score*, dan ROC AUC tertinggi dibandingkan dengan semua model lainnya. Meskipun seluruh model memiliki nilai *recall* yang sama, yaitu 63,64%, model ridge classifier menunjukkan nilai *precision* yang lebih tinggi, yang berimplikasi pada *F1-score* yang juga lebih besar. Hal ini menunjukkan adanya keseimbangan yang lebih baik antara jumlah prediksi positif yang akurat dan kemampuan untuk menangkap semua contoh positif. Secara keseluruhan, meskipun terdapat perbedaan yang tidak signifikan antar model, ridge classifier lebih unggul dalam memberikan hasil yang akurat, stabil, dan tahan terhadap *overfitting*. Meskipun logistic regression memiliki kecepatan dan kesederhanaan, performanya tidak optimal. Baik LDA maupun SVM linier sama-sama efektif untuk data yang mengikuti pola linier, namun ridge classifier menawarkan keunggulan tambahan berkat teknik regularisasi yang digunakannya. Oleh karena itu, dalam penelitian ini, ridge classifier direkomendasikan sebagai model terbaik untuk proses klasifikasi. Untuk mengetahui kinerja masing-masing model dalam mengklasifikasikan data, dilakukan pengujian dan evaluasi menggunakan beberapa metrik. Hasil evaluasi model tersebut disajikan pada Gambar 2.



Gambar 2. Perbandingan Akurasi Model

Gambar 2 menunjukkan sebuah grafik batang yang mengilustrasikan perbandingan kinerja empat jenis model klasifikasi logistic regression, linear discriminant analysis (LDA), ridge classifier, dan SVM dengan kernel linier berdasarkan lima ukuran evaluasi yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan AUC-ROC. Setiap model ditandai dengan lima batang berwarna yang berbeda: biru untuk *accuracy*, oranye untuk *precision*, hijau untuk *recall*, merah untuk *F1-score*, dan ungu untuk AUC-ROC. Sumbu horizontal menampilkan nama model yang berbeda, sementara sumbu vertikal menunjukkan skala nilai metrik dari 0 sampai 0,8. Persentase yang tertera di atas batang merupakan akurasi memberikan informasi tambahan tentang nilai akurasi dalam format persentase. Visualisasi ini dibuat untuk mempermudah perbandingan kinerja di antara model-model tersebut dan juga antar metrik dalam satu model, sehingga memberikan pandangan menyeluruh mengenai evaluasi kinerja model secara visual dan informatif. Dari hasil tersebut, dapat dilihat bahwa ridge classifier secara konsisten meraih nilai tertinggi di semua metrik utama, khususnya pada akurasi dan *precision* yang masing-masing mencapai 77,92% dan 71%. Temuan ini konsisten dengan hasil penelitian Maksabedian dkk. [21], yang juga menunjukkan bahwa ridge classifier merupakan model linier dengan kinerja terbaik dibandingkan model lainnya. Sementara itu, linear discriminant analysis dan SVM dengan kernel linier menunjukkan performa yang sangat mirip, keduanya lebih unggul dibandingkan logistic regression, meskipun masih sedikit di bawah model ridge classifier. Logistic regression berada di posisi paling rendah dalam semua metrik evaluasi, menunjukkan bahwa model ini kurang optimal untuk data yang diuji dibandingkan dengan model-model lainnya. Pola yang terlihat pada grafik juga menegaskan bahwa semua model mencapai nilai recall yang sama, yaitu 64%. Ini menunjukkan bahwa perbedaan kinerja antar model lebih banyak ditentukan oleh *precision*, *F1-score*, dan AUC-ROC. Visualisasi ini mendukung kesimpulan bahwa model ridge classifier adalah pilihan terbaik untuk klasifikasi pada dataset ini.

4. KESIMPULAN

Berdasarkan hasil pengujian terhadap empat model klasifikasi, yaitu logistic regression, linear discriminant analysis (LDA), ridge classifier, dan support vector machine (SVM) dengan kernel linier, dapat disimpulkan bahwa ridge classifier tampil sebagai model terbaik. Model ini mencatatkan *accuracy* sebesar 77,92%, *precision* sebesar 71,43%, *F1-score* sebesar 67,31%, serta ROC AUC sebesar 0,7475. Kelebihannya terletak pada kemampuannya dalam menjaga keseimbangan antara ketepatan dan sensitivitas, jika dibandingkan dengan model-model lainnya. Sementara itu, baik LDA maupun SVM linier menghasilkan performa yang cukup serupa dengan akurasi 76,62%. Di sisi lain, logistic regression menunjukkan performa terendah dengan *accuracy* sebesar 75,32%. Temuan ini menunjukkan bahwa penerapan regularisasi pada model ridge classifier berhasil meningkatkan stabilitas dan akurasi model, sehingga model ini direkomendasikan sebagai yang terbaik untuk klasifikasi penyakit diabetes pada dataset ini.

REFERENCES

- [1] N. M. Putry, "Komparasi Algoritma Knn Dan Naïve Bayes Untuk Klasifikasi Diagnosis Penyakit Diabetes Mellitus," *EVOLUSI J. Sains dan Manaj.*, vol. 10, no. 1, 2022, doi: 10.31294/evolusi.v10i1.12514.
- [2] S. Nurvita, "Diabetes Mellitus Tipe 1 Pada Anak Di Indonesia," *PREPOTIF J. Kesehat. Masy.*, vol. 7, no. 1, pp. 635–639, 2023.
- [3] A. E. Satriatama *et al.*, "Analisis Klaster Data Pasien Diabetes untuk Identifikasi Pola dan Karakteristik Pasien," *J. Teknol. Dan Sist. Inf. Bisnis*, vol. 5, no. 3, pp. 172–182, 2023, doi: 10.47233/jteksis.v5i3.828.
- [4] M. Saputra, J. P. Sidabuke, R. P. Sinulingga, and R. B. Tamba, "Analisis Metode Algoritma K-Nearest Neighbor (KNN) dan Naive Bayes Untuk Klasifikasi Diabetes Mellitus," *J. TEKINKOM*, vol. 6, no. 2, pp. 723–729, 2023, doi: 10.37600/tekinkom.v6i2.942.
- [5] N. S. Niko, A. Rahman, D. Marini Umi Atmaja, and A. Basri, "Prediksi Penyakit Diabetes Untuk Pencegahan Dini Dengan Metode Regresi Linear," *Bull. Inf. Technol.*, vol. 4, no. 3, pp. 313–219, 2023, doi: 10.47065/bit.v4i3.739.
- [6] N. L. Sabili and F. R. Umbara, "KLASIFIKASI PENYAKIT DIABETES MENGGUNAKAN ALGORITMA CATEGORICAL BOOSTING DENGAN FAKTOR RISIKO DIABETES," vol. 8, no. 6, pp. 11391–11398, 2024.

- [7] Erlin, Yulvia Nora Marlim, Junadhi, Laili Suryati, and Nova Agustina, "Deteksi Dini Penyakit Diabetes Menggunakan Machine Learning dengan Algoritma Logistic Regression," *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 11, no. 2, pp. 88–96, 2022, doi: 10.22146/jnteti.v11i2.3586.
- [8] D. S. Negara and K. U. Haqiqi, "Perbandingan Algoritma Neural Network dengan Linear Discriminant Analysis (LDA) pada Klasifikasi Penyakit Diabetes di RSUD Kudus," *J. Ilmu Komput. dan atematika*, vol. 1, no. 1, pp. 47–55, 2020.
- [9] S. Rabbani, D. Safitri, N. Rahmadhani, A. A. F. Sani, and M. K. Anam, "Perbandingan Evaluasi Kernel SVM untuk Klasifikasi Sentimen dalam Analisis Kenaikan Harga BBM," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 3, no. 2, pp. 153–160, 2023, doi: 10.57152/malcom.v3i2.897.
- [10] Muhammad Nur Ihsan Muhlashin and A. Stefanie, "Klasifikasi Penyakit Mata Berdasarkan Citra Fundus Menggunakan YOLO V8," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 2, pp. 1363–1368, 2023, doi: 10.36040/jati.v7i2.6927.
- [11] I. Maulana, A. Mutoi Siregar, and A. Fauzi, "Optimization of Machine Learning Model Accuracy for Brain Tumor Classification With Principal Component Analysis," *J. Tek. Inform.*, vol. 5, no. 3, pp. 903–915, 2024.
- [12] R. P. Kurniadi, R. R. Saedudin, and V. P. Widartha, "Perbandingan Akurasi Algoritma K-Nearest Neighbor Dan Logistic Regression Untuk Klasifikasi Penyakit Diabetes," *e-Proceeding Eng.*, vol. 8, no. 5, pp. 9757–9764, 2021.
- [13] A. D. Cahyani and A. Basuki, "Klasifikasi Diabetes Mellitus Menggunakan Support Vector Machine (Studi Kasus: Puskesmas Modopuro, Mojokerto)," *Rekayasa*, vol. 12, no. 2, pp. 174–182, 2019, doi: 10.21107/rekayasa.v12i2.19763.
- [14] F. M. Hana, N. D. Setia, and K. U. Khaqiqi, "PERBANDINGAN ALGORITMA NEURAL NETWORK DENGAN LINIER DISCRIMINANT ANALYSIS (LDA) PADA KLASIFIKASI PENYAKIT DIABETES," vol. 1, pp. 1541–1541, 2020.
- [15] UCI Machine Learning, "Pima Indians Diabetes Database," 2019. [Online]. Available: <https://www.kaggle.com/datasets/uciml/pima-indians-diabetes-database>
- [16] I. Maulana and S. Ernawati, "Meningkatkan Klasifikasi Penyakit Diabetes Menggunakan Metode Ensemble Softvoting Dengan SMOTE-ENN dan Optimasi Bayesian," vol. 13, no. 1, pp. 71–86, 2025.
- [17] A. N. W. Zain, M. Muafi, and A. Tholib, "Klasifikasi Data Mining Di Tingkat Kepuasan Mahasiswa Terhadap Pelayanan Sistem Informasi Fakultas Teknik Universitas Nurul Jadid," *SKANIKA Sist. Komput. dan Tek. Inform.*, vol. 7, no. 2, pp. 204–213, 2024, doi: 10.36080/skanika.v7i2.3200.
- [18] I. Hakim and T. Afriliansyah, "Implementasi Algoritma Komputasi Linear Regression untuk Optimasi Prediksi Hasil Pertanian," vol. 5, no. 3, pp. 1423–1434, 2024.
- [19] A. Pramudyantoro, E. Utami, and D. Ariatmanto, "Penggabungan K-Nearest Neighbors dan Lightgbm Untuk prediksi Diabetes Pada Dataset Pima Indians: Menggunakan pendekatan Exploratory Data Analysis," *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.)*, vol. 9, no. 3, pp. 1133–1144, 2024, [Online]. Available: <https://jurnal.stkipppgritlungagung.ac.id/index.php/jupi/article/view/4966/2114>
- [20] D. Kurniadi, F. Nuraeni, N. Faturrohman, and A. Mulyani, "Klasifikasi Perputaran Karyawan Perusahaan Menggunakan Algoritma Random Forest dan Random Over-sampling," *Edu Komputika J.*, vol. 10, no. 2, pp. 93–103, 2024, doi: 10.15294/edukomputika.v10i2.73782.
- [21] E. J. Maksabedian Hernandez, I. Tingzon, L. Ampil, and J. Tiu, "Identifying chronic disease patients using predictive algorithms in pharmacy administrative claims: an application in rheumatoid arthritis," *J. Med. Econ.*, vol. 24, no. 1, pp. 1272–1279, 2021, doi: 10.1080/13696998.2021.1999132.