

Klasifikasi Penyakit Diabetes Menggunakan Pendekatan Pembelajaran Mesin dengan Model Non-linier

Ilham Arif Kuncoro Adi*, Wahyu Aji Eko Prabowo

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

Email: ¹*111202113560@mhs.dinus.ac.id, ²prabowo@dsn.dinus.ac.id

Email Penulis Korespondensi: 111202113560@mhs.dinus.ac.id

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

Abstrak

Peningkatan prevalensi diabetes melitus menegaskan pentingnya pengembangan metode deteksi dini yang akurat. Penelitian ini mengusulkan model klasifikasi untuk prediksi diabetes menggunakan algoritma *machine learning* non-linier, yaitu support vector machine (SVM), random forest (RF), dan K-nearest neighbors (K-NN). Dataset yang digunakan diperoleh dari Kaggle dan mencakup fitur-fitur klinis seperti kadar glukosa, indeks massa tubuh (BMI), tekanan darah, dan kadar insulin. Metodologi penelitian mencakup tahap pra-pemrosesan data, pembagian data menjadi data latih dan data uji, serta evaluasi performa model menggunakan metrik akurasi, presisi, recall, dan F1-score. Hasil eksperimen menunjukkan bahwa algoritma random forest memberikan performa terbaik, diikuti oleh SVM dan K-NN. Kinerja unggul random forest dikaitkan dengan kemampuannya dalam menangani pola data yang kompleks dan meminimalkan *overfitting*. Penelitian ini diharapkan dapat berkontribusi dalam pengembangan alat deteksi dini diabetes yang praktis untuk mendukung pengambilan keputusan medis secara tepat waktu dan berbasis data.

Kata Kunci: Machine Learning; Diabetes Mellitus; Support Vector Machine; Random Forest; K-Nearest Neighbors

Abstract

The increasing prevalence of diabetes mellitus highlights the need for accurate early detection methods. This study proposes a classification model for diabetes prediction using non-linear machine learning algorithms, namely support vector machine (SVM), random forest (RF), and K-nearest neighbors (K-NN). The dataset, obtained from Kaggle, includes clinical features such as glucose levels, BMI, blood pressure, and insulin. The methodology comprises data preprocessing, partitioning the data into training and testing sets, and evaluating the model's using accuracy, precision, recall, and F1-score. Experimental results indicate that the random forest algorithm achieved the highest performance, followed by SVM and K-NN. We attribute random forest's superior performance to its robustness in handling complex patterns and minimizing overfitting. We expect this research to contribute to developing practical early detection tools for diabetes, thereby supporting timely and data-driven medical decision-making.

Keywords: Machine Learning; Diabetes Mellitus; Support Vector Machine; Random Forest; K-Nearest Neighbors

1. PENDAHULUAN

Diabetes melitus merupakan salah satu penyakit tidak menular yang menunjukkan peningkatan prevalensi secara global maupun nasional. Menurut data *World Health Organization* (WHO), sekitar 70% kematian di dunia disebabkan oleh penyakit tidak menular, di mana diabetes menempati posisi yang mengkhawatirkan sebagai penyebab kematian kedua tertinggi di Indonesia setelah Sri Lanka [1]. Pada tahun 2021, jumlah penderita diabetes di Indonesia mencapai 19,47 juta jiwa, meningkat drastis dari 7,29 juta pada tahun 2011 atau sekitar 267% [2]. Lebih dari dua pertiga penderita tidak menyadari kondisinya, sehingga risiko komplikasi serius seperti stroke, gagal ginjal, dan penyakit jantung menjadi sangat tinggi [3], [4].

Kondisi ini diperburuk oleh keterlambatan dalam diagnosis dan rendahnya kesadaran masyarakat terhadap pentingnya deteksi dini. Oleh karena itu, diperlukan sistem deteksi awal yang akurat, cepat, dan dapat diandalkan. Salah satu pendekatan yang menjanjikan adalah penerapan algoritma *machine learning* dalam sistem prediksi berbasis data. Beberapa algoritma yang telah banyak digunakan dalam klasifikasi diabetes meliputi support vector machine (SVM), random forest (RF), naive bayes (NB), logistic regression (LR), dan K-nearest neighbors (K-NN) [5], [6]. Studi sebelumnya menunjukkan bahwa algoritma-algoritma ini memiliki potensi dalam memberikan akurasi prediksi yang tinggi. Misalnya, model random forest berhasil mencapai akurasi hingga 97% [7], sementara SVM dengan kernel Gaussian mampu mencapai 87,94% [4], [8]. Di sisi lain, efektivitas algoritma sangat bergantung pada proses pra-pemrosesan data, pemilihan fitur, serta metode evaluasi yang digunakan. Beberapa penelitian sebelumnya juga telah memanfaatkan teknik validasi seperti 10-fold *cross-validation* untuk meningkatkan keandalan model, serta melakukan normalisasi data menggunakan metode Min-Max Scaling [6], [9]. Selain itu, metrik evaluasi seperti akurasi, presisi, recall, dan F1-score menjadi indikator utama dalam mengukur performa model klasifikasi [10], [11], [12].

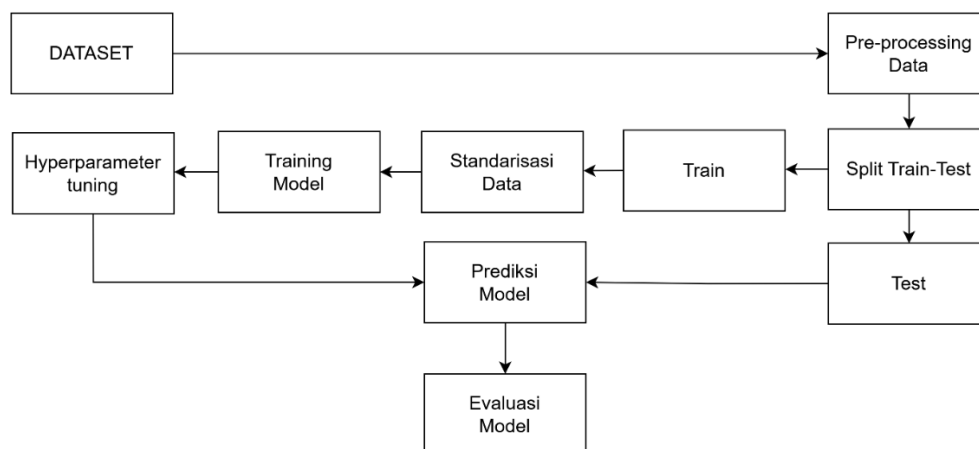
Penelitian terdahulu yang menjadi acuan dalam penelitian ini melibatkan enam studi terkait prediksi diabetes menggunakan algoritma pembelajaran mesin, masing-masing dengan pendekatan dan hasil evaluasi yang berbeda. Penelitian pertama dilakukan oleh Airlangga (2024) yang mengevaluasi secara komparatif berbagai algoritma pembelajaran mesin, termasuk random forest dan support vector machine, dalam memprediksi penyakit diabetes. Studi ini menyoroti manfaat metode ensemble dalam meningkatkan ketepatan model, mengindikasikan bahwa menggabungkan beberapa model bisa menghasilkan keluaran yang lebih konsisten dan tepat. Akan tetapi, penelitian ini tidak membahas penerapan model dalam praktik nyata seperti aplikasi klinis atau sistem untuk mendukung pengambilan keputusan [13]. Penelitian kedua yang dilakukan oleh Albert dan Leong (2024) membandingkan kinerja algoritma SVM dan random

forest dengan menggunakan data pasien dari rumah sakit di Amerika Serikat. Temuan menunjukkan bahwa meskipun akurasi kedua algoritma tersebut hampir sama, random forest unggul dalam hal kecepatan komputasi. Salah satu kekurangan dari penelitian ini adalah tidak adanya pengujian pada kombinasi algoritma (*ensemble hybrid*) dan fokus yang terlalu sempit pada efisiensi tanpa menginvestigasi lebih jauh tentang interpretabilitas model [14]. Penelitian ketiga yang dilakukan oleh Alrasyid dkk. (2024) berfokus pada perbandingan antara algoritma SVM dan random forest dalam mendeteksi diabetes. Dalam studi ini, algoritma SVM mencatat tingkat akurasi sebesar 77%, sedikit lebih tinggi daripada random forest yang mendapatkan akurasi sebesar 75%. Kelemahan dari penelitian ini adalah tidak menggunakan teknik validasi silang yang lebih baik dan tidak adanya diskusi mengenai stabilitas model terhadap variasi dalam data [15]. Penelitian keempat oleh Maulidiyyah dkk. (2024) menganalisis dua pendekatan klasifikasi, yaitu decision tree dan random forest, dalam mengklasifikasikan penyakit diabetes mellitus. Temuan menunjukkan bahwa model random forest memiliki akurasi rata-rata 94%, sementara decision tree hanya mencapai 93%. Penelitian ini mengungkapkan bahwa metode ensemble lebih efektif dalam hal kinerja, meskipun hanya membandingkan dua algoritma tanpa mempertimbangkan algoritma lain seperti KNN atau SVM [16]. Penelitian kelima dilakukan oleh Alsadi dkk. (2024), yang menciptakan sebuah model yang menggunakan metode *ensemble learning* untuk memprediksi diabetes tipe 2 serta relevansinya dengan kesehatan tulang. Model ini menunjukkan kemungkinan yang besar untuk digunakan dalam praktik klinis sebagai sarana deteksi awal. Meski demikian, penelitian ini memiliki perhatian khusus terhadap hubungan dengan kesehatan tulang, sehingga penerapannya lebih terbatas dibandingkan dengan prediksi umum tentang diabetes tanpa mempertimbangkan komorbiditas [17]. Acuan terakhir dari penelitian yang dilakukan oleh Aditi Apurva (2024) yang membahas tentang analisis berbagai algoritma mesin belajar seperti KNN, SVM, decision tree, dan random forest dalam hal prediksi risiko diabetes. Hasil dari studi ini menunjukkan bahwa random forest menghasilkan tingkat akurasi tertinggi sebesar 99,03%, sehingga dianggap sebagai metode terbaik dalam penelitian ini. Akan tetapi, penelitian ini tidak mencakup analisis mengenai waktu komputasi atau kinerja model dalam situasi data nyata yang mengandung noise atau ketidakseimbangan kelas [18]. Berdasarkan beberapa studi yang menjadi acuan di atas, studi ini berupaya menggabungkan keunggulan dari berbagai pendekatan sekaligus, tidak hanya terfokus pada satu atau dua algoritma, namun beberapa model berbasis model non-linier. Selain itu, penelitian ini juga bertujuan untuk tidak hanya menilai kinerja model dari aspek akurasi, tetapi juga memperhatikan stabilitas, efisiensi dalam perhitungan, serta kemungkinan penerapannya dalam sistem berbasis teknologi informasi yang dapat dimanfaatkan oleh lembaga medis atau profesional kesehatan dalam proses diagnosis awal diabetes.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Tahapan penelitian yang dilakukan dalam penelitian ini dijelaskan secara umum pada Gambar 1. Proses ini mencakup serangkaian langkah yang dimulai dari pengumpulan dan pra-pemrosesan data, pembagian data menjadi data latih dan data uji. Data latih digunakan sebagai pelatihan model, sementara data uji dikondisikan dalam keadaan terjaga (*unseen*) sebagai evaluasi performa model yang dibangun. Gambar 1 menunjukkan bahwa proses dimulai dengan penyediaan dan pra-pemrosesan dataset, dilanjutkan dengan pembagian data menjadi data latih dan data uji. Selanjutnya, dilakukan pelatihan dan *hyperparameter tuning* untuk menentukan parameter-parameter yang penting dan kemudian diterapkan pada algoritma pembelajaran mesin yang telah dipilih. Langkah akhir adalah prediksi dan evaluasi performa model menggunakan berbagai metrik pengukuran.



Gambar 1. Tahapan Penelitian

2.2 Dataset

Dataset yang digunakan dalam penelitian ini adalah dataset Diabetes yang terdiri dari 768 entri data dan 9 fitur, termasuk fitur target (*Outcome*). Dataset ini memuat informasi medis tentang pasien perempuan keturunan Pima Indian yang berusia di atas 21 tahun, dengan tujuan untuk memprediksi apakah seorang pasien menderita diabetes. Setiap fitur

menggambarkan parameter diagnostik, seperti jumlah kehamilan, kadar glukosa, tekanan darah, ketebalan kulit, kadar insulin, indeks massa tubuh (BMI), riwayat keluarga diabetes, dan usia. Fitur target (Outcome) bernilai 0 untuk pasien yang tidak menderita diabetes dan 1 untuk yang menderita diabetes. Dataset ini sering dijadikan acuan (benchmark) dalam tugas klasifikasi di bidang data science dan kesehatan. Tabel 1 di bawah ini menunjukkan semua fitur yang digunakan dalam penelitian ini.

Tabel 1. Deskripsi Fitur Dataset

No	Nama Fitur	Tipe Data	Deskripsi
1	Pregnancies	Integer	Jumlah kehamilan yang pernah dialami
2	Glucose	Integer	Kadar glukosa dalam plasma (mg/dL)
3	BloodPressure	Integer	Tekanan darah diastolik (mm Hg)
4	SkinThickness	Integer	Ketebalan lipatan kulit trisep (mm)
5	Insulin	Integer	Jumlah insulin serum 2 jam (mu U/ml)
6	BMI	Float	Indeks massa tubuh (berat badan (kg) / (tinggi badan (m)) ²)
7	DiabetesPedigreeFunction	Float	Fungsi riwayat keluarga diabetes (faktor genetik)
8	Age	Integer	Usia pasien (tahun)
9	Outcome	Integer	Kelas hasil: 0 = Tidak diabetes, 1 = Diabetes

Tabel 1 menunjukkan bahwa semua fitur yang digunakan dalam penelitian ini bersifat numerik, yang memudahkan dalam proses normalisasi dan pengolahan data menggunakan algoritma *machine learning*. Fitur-fitur ini mencerminkan kondisi klinis yang relevan dengan diagnosis diabetes.

2.3 Pre-processing

Pra-pemrosesan data adalah langkah awal yang krusial untuk memastikan bahwa data dalam kondisi optimal sebelum digunakan dalam proses pemodelan. Tahapan ini mencakup berbagai aktivitas, seperti menghapus data duplikat, menangani data yang tidak lengkap, mengonversi variabel kategorikal menjadi format numerik, serta melakukan normalisasi pada data numerik. Kualitas hasil dari model sangat dipengaruhi oleh kebersihan dan kelengkapan data, sehingga tahap ini memegang peranan yang sangat penting dalam keseluruhan proses analisis data.

2.4 Splitting Data

Setelah dilakukan pra-pemrosesan, data kemudian dibagi menjadi dua bagian: 80% digunakan sebagai data latih (training set), dan 20% sebagai data uji (testing set). Pembagian ini dilakukan secara acak untuk memastikan evaluasi model lebih objektif. Rincian jumlah data dalam masing-masing set ditunjukkan pada tabel di bawah ini. Tabel 2 menunjukkan bahwa dari total 768 data, sebanyak 614 digunakan untuk melatih model dan 154 digunakan untuk menguji model. Perbandingan rasio ini dipilih dengan dasar agar model memperoleh cukup data untuk belajar dan diuji pada data yang benar-benar baru (*unseen*).

Tabel 2. Splitting Data

Jenis Data	Jumlah Data
Data Latih	614
Data Uji	154
Total	768

2.5 Standarisasi Data

Proses standarisasi diterapkan pada seluruh fitur numerik dalam dataset menggunakan metode StandardScaler dari Scikit-learn. Langkah ini sangat krusial karena banyak algoritma linier memiliki sensitivitas yang tinggi terhadap skala data. Dengan melakukan standarisasi, setiap fitur akan memiliki rata-rata sebesar 0 dan standar deviasi sebesar 1. Hal ini tidak hanya mempercepat proses pelatihan, tetapi juga meningkatkan akurasi model yang dihasilkan [19].

2.6 Training Model

Penelitian ini menggunakan lima jenis model algoritma pembelajaran mesin non-linier, yaitu: K-nearest neighbors (KNN), random forest, decision tree, support vector machine (SVM) menggunakan radial basis function (RBF), serta SVM dengan Polynomial. Pemilihan model-model ini didasarkan pada efisiensinya terhadap dataset yang telah distandarisasi serta kemampuannya yang baik dalam melakukan generalisasi. Proses pelatihan model dilakukan pada data pelatihan yang telah melalui tahap pengolahan dan distandarisasi.

2.7 Hyperparameter Tuning

Untuk mencapai performa model yang optimal, dilakukan pencarian kombinasi parameter terbaik dengan menggunakan *GridSearchCV* dan pengaturan cross-validation = 5. Hal ini berarti bahwa data pelatihan akan secara otomatis dibagi menjadi lima bagian selama proses validasi silang. Model akan diuji secara menyeluruh terhadap berbagai kombinasi parameter, sehingga memungkinkan kita menemukan konfigurasi yang secara konsisten memberikan hasil terbaik [20].

2.8 Prediksi Model

Setelah proses pelatihan dan optimasi model selesai, versi terbaik dari model tersebut akan diaplikasikan untuk prediksi terhadap data pengujian. Tujuan dari langkah ini adalah untuk mengukur sejauh mana kemampuan model dalam menggeneralisasi pola dari data yang belum pernah ditemui sebelumnya. Hasil prediksi tersebut akan menjadi dasar dalam mengevaluasi kinerja model.

2.9 Evaluasi Model

Kinerja model akan dinilai dengan menggunakan lima metrik utama, yaitu akurasi, presisi, recall, F1-score, dan AUC-ROC. Setiap metrik ini memberikan perspektif yang berbeda dalam menilai keberhasilan model, baik dari segi akurasi keseluruhan maupun ketepatan dalam memprediksi kelas positif. Kombinasi dari berbagai metrik ini bertujuan untuk menyajikan gambaran yang komprehensif mengenai performa dan efektivitas model klasifikasi yang telah dikembangkan.

2.10 Metrik Pengukuran Evaluasi

Kinerja model akan dievaluasi menggunakan lima metrik utama, yaitu akurasi, presisi, recall, F1-score, dan AUC-ROC. Kinerja model dari berbagai metrik ini bertujuan untuk menyajikan gambaran yang komprehensif mengenai performa dan efektivitas model klasifikasi yang telah dikembangkan.

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

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

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

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

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

$$FPR = \frac{FP}{FP+TN} \quad (6)$$

$$TPR = \frac{TP}{TP+FN} \quad (7)$$

Penjelasan:

TP = Prediksi benar positif (*True Positive*)

TN = Prediksi benar negatif (*True Negative*)

FP = Prediksi Salah (*False Positive*, seharusnya negatif tapi diprediksi positif)

FN = Prediksi Salah (*False Negative*, seharusnya positif tapi diprediksi negative)

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

$TPR = \frac{TP}{TP+FN}$: *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 menganalisis performa lima algoritma pembelajaran mesin non-linier, yang terdiri dari: K-nearest neighbors (KNN), random forest, decision tree, support vector machine (SVM) menggunakan radial basis function (RBF), serta SVM dengan Polynomial. Semua algoritma tersebut diterapkan untuk mengklasifikasikan penyakit diabetes berdasarkan data dari Pima Indians Diabetes. Sebelum menjelaskan hasil pengujian, bagian ini akan menjelaskan langkah-langkah penerapan metode dengan cara yang teratur, sehingga dapat dipahami dengan baik bagaimana algoritma diterapkan.

3.1 Tahapan Penerapan Metode

Tahapan penerapan metode klasifikasi terdiri dari beberapa langkah penting :

a. Pengumpulan dan Pra-pemrosesan Data

Data diambil dari dataset Diabetes Pima Indians yang mencakup 768 datapoin. Proses pra-pemrosesan dilakukan untuk meningkatkan mutu data sebelum model dilatih. Ini mencakup penghapusan data yang sama dan penanganan nilai yang hilang. Selanjutnya, normalisasi dilakukan dengan menggunakan *StandardScaler* agar semua fitur numerik memiliki skala yang seragam (rata-rata = 0, deviasi standar = 1).

b. Pembagian Data

Dataset dibagi menjadi dua segmen, di mana 80% digunakan sebagai data pelatihan dan 20% untuk data pengujian. Pemisahan ini dilakukan secara acak untuk memastikan objektivitas dan mencegah overfitting.

c. Pelatihan Model

Setiap model machine learning yang akan diuji, dilatih dengan menggunakan data pelatihan. Proses ini dilakukan setelah data dinormalisasi, dan evaluasi awal terjadi pada data pelatihan dengan menerapkan metrik evaluasi klasifikasi.

d. Optimasi *Hyperparameter Tuning*

Model yang telah dilatih kemudian dioptimalkan melalui GridSearchCV dengan validasi silang (5-fold *cross-validation*) untuk menemukan parameter yang optimal. Tujuannya agar model tidak hanya tepat tetapi juga konsisten dan dapat diandalkan.

e. Prediksi dan Evaluasi Kinerja

Hasil dari optimisasi hyperparameter tuning kemudian digunakan untuk prediksi klasifikasi penyakit diabetes menggunakan lima model pembelajaran mesin. Hasil prediksi dievaluasi dengan menggunakan metrik: akurasi, presisi, recall, F1-score, dan AUC-ROC, untuk menilai kemampuan model dalam generalisasi terhadap data yang baru.

3.2 Hasil Data Pelatihan

Untuk mengevaluasi seberapa baik model dapat mengenali pola dalam data pelatihan, telah dilakukan uji coba terhadap lima model machine learning yang sudah dilatih. Tabel 3 menunjukkan evaluasi performa dari metrik pengukuran.

Tabel 3. Hasil pengujian data pelatihan (*train*)

Model	Accuracy	Precision	F1-Score	AUC-ROC
KNN	0.837	0.806	0.844	0.931
Random Forest	1.000	1.000	1.000	1.000
Decision Tree	1.000	1.000	1.000	1.000
SVM (RBF)	0.846	0.836	0.848	0.922
SVM (Polynomial)	0.795	0.846	0.777	0.886

Hasil pada Tabel 3 mengungkapkan bahwa model random forest dan decision tree berhasil meraih nilai sangat baik di semua metrik. Ini menunjukkan bahwa kedua model tersebut dapat mendeteksi semua pola dalam data pelatihan tanpa melakukan kesalahan dalam klasifikasi. Namun, hasil sangat baik seperti ini bisa mengindikasikan terjadinya *overfitting*, yaitu keadaan di mana model terlalu terikat pada data pelatihan dan berisiko kurang efektif dalam mengenali pola baru pada data pengujian. Di sisi lain, KNN mencatat hasil yang cukup baik namun lebih realistis, dengan tingkat akurasi 83,7% dan F1-score 84,4%, yang menunjukkan keseimbangan yang memadai antara presisi dan recall. Hasil ini mengindikasikan bahwa KNN mampu menghindari *overfitting* dan memiliki kemampuan generalisasi yang efektif. SVM dengan kernel RBF menunjukkan hasil yang sebanding dengan KNN, sementara SVM Polynomial memiliki tingkat akurasi dan F1-score yang lebih rendah, meskipun presisinya tetap tinggi. Ini menunjukkan bahwa SVM Polynomial cenderung menghasilkan prediksi positif yang akurat, namun kurang efektif dalam menangkap seluruh kasus positif yang sebenarnya.

3.3 Hasil Data Pengujian

Setelah melaksanakan pelatihan pada lima model pembelajaran mesin, tahap berikutnya adalah menilai kinerja setiap model menggunakan data pengujian, yaitu data yang tidak digunakan (*unseen*) saat proses pelatihan. Tujuan dari evaluasi ini adalah untuk memahami seberapa baik model dapat generalisasi saat klasifikasi data baru yang belum pernah dikenali sebelumnya. Hasil dari pengujian pada data uji ini ditampilkan pada Tabel 4 di bawah ini.

Tabel 4. Hasil pengujian data pengujian (*testing*)

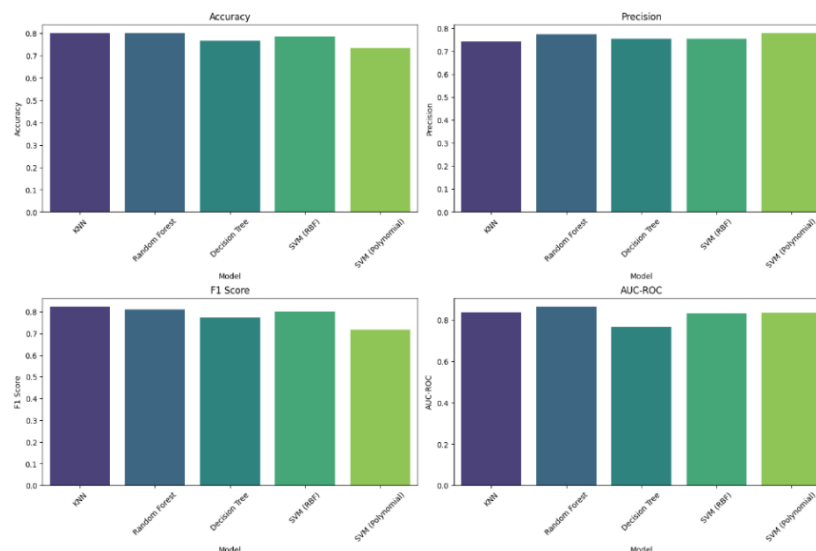
Model	Accuracy	Precision	F1-Score	AUC-ROC
KNN	0.800	0.744	0.823	0.835
Random Forest	0.800	0.775	0.811	0.863
Decision Tree	0.765	0.755	0.773	0.765
SVM (RBF)	0.785	0.754	0.800	0.831
SVM (Polynomial)	0.735	0.779	0.717	0.834

Hasil pada Tabel 4 menunjukkan bahwa random forest adalah model yang memberikan performa paling konsisten dan unggul dalam hal AUC-ROC sebesar 0.863, yang menunjukkan kemampuannya membedakan antara kelas positif dan negatif dengan akurasi tinggi. Temuan ini sejalan dengan penelitian Afuan dkk [21], yang menyatakan bahwa random forest merupakan model non-linier dengan performa yang sangat kompetitif. Sementara itu, model K-nearest neighbors (KNN) menunjukkan performa yang cukup kompetitif dengan akurasi mencapai 80%, *precision* sebesar 74,40%, F1-score 82,30%, dan AUC-ROC sebesar 0,835. Algoritma ini menunjukkan kemampuannya dalam mendeteksi pasien diabetes, terlihat dari nilai F1-score yang tertinggi jika dibandingkan dengan model lainnya. KNN beroperasi berdasarkan prinsip kedekatan jarak antar data, sehingga kinerjanya sangat dipengaruhi oleh distribusi data dan teknik pra-pemrosesan seperti normalisasi. Meskipun F1-score yang tinggi mencerminkan keseimbangan yang baik antara *precision* dan *recall*, nilai *precision* yang relatif rendah menunjukkan bahwa KNN cenderung menghasilkan lebih banyak prediksi positif palsu. Hal ini menjadi perhatian penting, mengingat dalam konteks medis, hasil positif palsu dapat mengakibatkan diagnosis

dan penanganan yang tidak perlu. Selain itu, KNN juga memiliki kelemahan dalam hal efisiensi, terutama saat diterapkan pada dataset yang besar, karena proses pencarian tetangga terdekat dapat memakan waktu yang cukup lama. Pada model yang lainnya yaitu decision tree, menunjukkan performa yang relatif lebih rendah dibandingkan dengan model-model lainnya, dengan nilai akurasi mencapai 76,5%, *precision* 75,47%, F1-score 77,29%, dan AUC-ROC sebesar 0,764. Meskipun demikian, model ini memiliki sifat yang intuitif dan mudah diinterpretasikan, sehingga sering dimanfaatkan dalam aplikasi medis untuk memahami jalur keputusan. Namun, hasil penelitian ini mengindikasikan bahwa model decision tree rentan terhadap *overfitting*, terutama ketika pohon berkembang terlalu dalam tanpa adanya mekanisme pruning atau batasan kedalaman. *Overfitting* ini dapat menyebabkan model terlalu menyesuaikan diri dengan data pelatihan dan kehilangan kemampuan untuk menggeneralisasi pada data baru. Walaupun demikian, dengan melakukan tuning lebih lanjut seperti penerapan teknik *pruning* atau pendekatan *ensemble* (misalnya random forest), performa model decision tree dapat ditingkatkan secara signifikan.

Pada model support vector machine (SVM) dengan kernel radial basis function (RBF) menunjukkan kinerja yang baik dalam klasifikasi diabetes, dengan tingkat akurasi mencapai 78,5%, *precision* 75,44%, F1-score 80%, serta AUC-ROC sebesar 0,831. Kernel RBF efektif dalam menangani data non-linear dengan cara memetakan data input ke dalam ruang berdimensi lebih tinggi, yang mempermudah proses pemisahan antar kelas. Hasil ini menunjukkan bahwa hubungan antara fitur dalam data diabetes tidak sepenuhnya linier, sehingga penggunaan kernel non-linear seperti RBF dapat meningkatkan kinerja klasifikasi. Meskipun akurasinya tidak setinggi model random forest, model SVM dengan kernel RBF menawarkan keuntungan dalam stabilitas klasifikasi dan cukup efektif untuk menangani data medis yang kompleks. Namun, kelemahan dari model ini terletak pada sensitivitasnya terhadap pemilihan parameter seperti C dan gamma, yang jika tidak dioptimalkan dengan baik, dapat mengurangi kinerja model. Sementara pada model support vector machine (SVM) dengan kernel Polynomial menunjukkan performa yang paling rendah di antara semua model yang diuji, dengan akurasi mencapai 73,5%, *precision* 77,90%, F1-score 71,65%, dan AUC-ROC sebesar 0,834. Meskipun *precision* dari model ini cukup tinggi, yang menunjukkan bahwa prediksi positifnya cukup akurat, nilai F1-score yang rendah mengindikasikan adanya ketidakseimbangan dalam performa *recall*. SVM Polynomial memiliki sensitivitas yang lebih tinggi terhadap derajat polinomial yang digunakan serta skala data, sehingga jika parameter tidak dituning dengan hati-hati, model ini berisiko mengalami *overfitting* atau *underfitting*. Selain itu, model ini memerlukan waktu komputasi yang lebih lama dibandingkan dengan kernel RBF, yang menjadi tantangan tambahan saat menghadapi dataset yang lebih besar dan kompleks. Hasil ini menegaskan bahwa walaupun SVM Polynomial dapat bermanfaat dalam situasi tertentu, dalam konteks ini, performanya kurang bersaing dibandingkan dengan metode lain seperti Random Forest dan SVM RBF.

Visualisasi kinerja keempat metrik pengukuran dapat dilihat pada Gambar 2 yang menunjukkan perbandingan nilai akurasi, presisi, F1-score, dan AUC-ROC dari kelima jenis model pembelajaran mesin. Dari grafik tersebut, terlihat bahwa model random forest menunjukkan keunggulan di hampir semua metrik evaluasi, dengan nilai AUC-ROC tertinggi yang menunjukkan kemampuannya dalam membedakan antara kelas positif dan negatif dengan lebih konsisten dibandingkan model lainnya. KNN juga menunjukkan performa yang baik, terutama dalam konteks F1-score, yang menunjukkan keseimbangan yang baik antara presisi dan recall. Sementara itu, pada model decision tree cenderung menunjukkan penurunan kinerja saat diuji dengan data baru, yang menandakan potensi *overfitting* meskipun sebelumnya menunjukkan hasil yang sangat baik pada data pelatihan. Model SVM dengan kernel RBF tetap menunjukkan kinerja yang stabil dan tangguh dalam semua metrik, mengindikasikan efektivitas kernel non-linear dalam menangkap hubungan kompleks di antara fitur dalam dataset diabetes. Sementara itu, SVM dengan kernel Polinomial menunjukkan performa terendah dalam hal F1-score dan akurasi, meskipun memiliki nilai presisi yang cukup tinggi, yang menunjukkan bahwa model ini cenderung memberikan banyak prediksi positif yang benar, tetapi kurang efektif dalam mendeteksi seluruh kasus positif.



Gambar 2. perbandingan performa lima metrik pengukuran

4. KESIMPULAN

Penelitian ini menganalisis perbandingan antara lima algoritma machine learning non-linier untuk mengklasifikasikan penyakit diabetes dengan memanfaatkan dataset Pima Indians Diabetes. Kelima model pembelajaran mesin tersebut adalah K-nearest neighbors (KNN), random forest, decision tree, serta support vector machine (SVM) dengan kernel RBF dan Polynomial. Berdasarkan evaluasi menggunakan ukuran seperti akurasi, presisi, F1-score, dan AUC-ROC, model random forest menunjukkan hasil terbaik dengan tingkat akurasi yang tinggi dan kuat dalam kestabilan prediksi, diikuti oleh KNN dan SVM RBF. Visualisasi kinerja mendukung temuan tersebut, menunjukkan bahwa model random forest selalu menunjukkan keunggulan di semua ukuran evaluasi. Kajian ini menyimpulkan bahwa random forest adalah algoritma paling efisien untuk mengembangkan sistem deteksi dini diabetes yang akurat serta dapat diandalkan dalam mendukung keputusan medis yang berbasis data.

REFERENCES

- [1] S. Wirapati, D. Luh, G. Astuti, and M. Kom, "Sistem Pakar Untuk Membantu Diagnosis Diabetes Menggunakan Machine Learning Dengan Algoritma Jaringan Saraf Tiruan," *J. Elektron. Ilmu Komput. Udayana*, vol. 11, no. 4, 2022.
- [2] R. Rizki, R. Athallah, I. Cholissodin, and P. P. Adikara, "Prediksi Potensi Pengidap Penyakit Diabetes berdasarkan Faktor Risiko Menggunakan Algoritme Kernel K-Nearest Neighbor," 2022. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [3] J. J. Pangaribuan, "mendiagnosis penyakit diabetes melitus dengan menggunakan metode extreme learning machine," 2022.
- [4] M. D. Purbolaksono, M. Irvan Tantowi, A. Imam Hidayat, and A. Adiwijaya, "Perbandingan Support Vector Machine dan Modified Balanced Random Forest dalam Deteksi Pasien Penyakit Diabetes," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 5, no. 2, pp. 393–399, Apr. 2021, doi: 10.29207/resti.v5i2.3008.
- [5] N. Maulidah, R. Supriyadi, D. Y. Utami, F. N. Hasan, A. Fauzi, and A. Christian, "Prediksi Penyakit Diabetes Melitus Menggunakan Metode Support Vector Machine dan Naive Bayes," *Indones. J. Softw. Eng.*, vol. 7, no. 1, pp. 63–68, 2021, [Online]. Available: <http://ejournal.bsi.ac.id/ejurnal/index.php/ijse63>
- [6] A. Oktaviana, D. P. Wijaya, A. Pramuntadi, and D. Heksaputra, "Prediksi Penyakit Diabetes Melitus Tipe 2 Menggunakan Algoritma K-Nearest Neighbor (K-NN)," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 3, pp. 812–818, May 2024, doi: 10.57152/malcom.v4i3.1268.
- [7] K. A. Saputro, E. M. Atsir, and H. Hasanah, "perbandingan tingkat akurasi penyakit diabetes menggunakan metode regresi logistik dan random forest," vol. 4, no. 2, pp. 159–166, 2024.
- [8] F. R. Lumbanraja, F. Lufiana, Y. Heningtyas, and K. Muludi, "implementasi support vector machine (svm) untuk klasifikasi penderita diabetes mellitus 1,*)." "
- [9] W. Apriliah *et al.*, "SISTEMASI: Jurnal Sistem Informasi Prediksi Kemungkinan Diabetes pada Tahap Awal Menggunakan Algoritma Klasifikasi Random Forest," 2021. [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [10] P. N. Bayes, R. Logistik, R. Forest, and N. Bayes, "SENASTIKA Universitas Malikussaleh," pp. 1–10.
- [11] I. M. Karo Karo and H. Hendriyana, "Klasifikasi Penderita Diabetes menggunakan Algoritma Machine Learning dan Z-Score," *J. Teknol. Terpadu*, vol. 8, no. 2, pp. 94–99, 2022, doi: 10.54914/jtt.v8i2.564.
- [12] A. W. Mucholladin, F. A. Bachtar, and M. T. Furqon, "Klasifikasi Penyakit Diabetes menggunakan Metode Support Vector Machine," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 5, no. 2, pp. 622–633, 2021, [Online]. Available: <http://j-ptiik.ub.ac.id>
- [13] G. Airlangga, "Comparative Analysis of Machine Learning Models for Predicting Diabetes: Unveiling the Superiority of Advanced Ensemble Methods," *G-Tech J. Teknol. Terap.*, vol. 8, no. 2, pp. 1272–1280, 2024, doi: 10.33379/gtech.v8i2.4246.
- [14] D. A. Albert and H. Leong, "Comparative Performance Analysis of Support Vector Machine and Random Forest on Diabetes Patient Data From Hospitals in the United States," *Proxies J. Inform.*, vol. 7, no. 2, pp. 85–101, 2024, doi: 10.24167/proxies.v7i2.12469.
- [15] H. Alrasyid, A. Homaidi, M. Kom, Z. Fatah, and M. Kom, "Comparison Support Vector Machine and Random Forest Algorithms in Detect Diabetes," vol. 1, no. 1, pp. 447–453, 2024.
- [16] N. A. Maulidiyyah, T. Trimono, A. T. Damaliana, and D. A. Prasetya, "Comparison of Decision Tree and Random Forest Methods in the Classification of Diabetes Mellitus," *JIKO (Jurnal Inform. dan Komputer)*, vol. 7, no. 2, pp. 79–87, 2024, doi: 10.33387/jiko.v7i2.8316.
- [17] B. Alsadi *et al.*, "An ensemble-based machine learning model for predicting type 2 diabetes and its effect on bone health," *BMC Med. Inform. Decis. Mak.*, vol. 24, no. 1, pp. 1–15, 2024, doi: 10.1186/s12911-024-02540-0.
- [18] S. A. Apurva, "compititive study of machine learning algorithms comparative study machine for diabetes risk prediction : a focus on k-nn , svm prediction, and random forest," vol. 19, no. 1, pp. 39–43, 2024.
- [19] Z. Maisat, E. Darmawan, and A. Fauzan, "Implementasi Optimasi Hyperparameter GridSearchCV Pada Sistem Prediksi Serangan Jantung Menggunakan SVM Implementation of GridSearchCV Hyperparameter Optimization in Heart Attack Prediction System Using SVM," *Unipdu*, vol. 13, no. 1, pp. 8–15, 2023.
- [20] A. Hamied Nababan and M. Y. Hutagalung, "Hyperparameter Tuning Pada Model Stance Detection Menggunakan GridSearchCV," *J. Sains dan Teknol.*, vol. 5, no. 1, pp. 205–209, 2023, [Online]. Available: <https://doi.org/10.55338/saintek.v5i1.1505>.
- [21] L. Afuan and R. R. Isnanto, "A comparative study of machine learning algorithms for fall detection in technology-based healthcare system: analyzing SVM, KNN, decision tree, random forest, LSTM, and CNN," *E3S Web Conf.*, vol. 605, 2025, doi: 10.1051/e3sconf/202560503051.