

Analisis Perbandingan 3 Model Machine Learning Pada Deteksi Penyakit Jantung

Adinda Pangestu¹, Diah Widi Astuti², Aulia Safira Putri³, Aghni Syahrul Muarof⁴, Tegar Wijaya Kusuma⁵,
Fiqri Novan Dwi Andika⁶

¹Program Studi Informatika, Fakultas Ilmu Komputer, Universitas Amikom Purwokerto, Purwokerto, Indonesia

²Program Studi Informatika, Fakultas Saintek, UIN Prof. K.H. Saifuddin Zuhri Purwokerto, Purwokerto, Indonesia

surel: ¹adindapangestu17@gmail.com, ²widiidiah288@gmail.com, ³auliasafirap79@gmail.com, ⁴syahrularof@gmail.com, ⁵tegarwijayakusuma2001@gmail.com, ⁶fikrinofan11@gmail.com

Info Artikel

Sejarah artikel:

Diterima 26 Juni 2025

Revisi 30 juli 2025

Diterima 5 agustus 2025

Kata kunci:

Penyakit jantung
Machine learning
Random Forest
Logistic Regression
KNN
Klasifikasi

ABSTRAK

Penyakit jantung merupakan penyebab utama kematian di Indonesia dan dunia, sehingga deteksi dini menjadi sangat penting. Penelitian ini membandingkan performa tiga algoritma machine learning—Logistic Regression, Random Forest, dan K-Nearest Neighbors (KNN) dalam memprediksi penyakit jantung menggunakan data dari Kaggle. Proses penelitian mencakup pengumpulan data, preprocessing, pembagian data, pemodelan, dan evaluasi menggunakan metrik akurasi, presisi, recall, dan F1-score. Hasil evaluasi menunjukkan bahwa model KNN memiliki performa terbaik dengan akurasi 85,88% dan keseimbangan yang tinggi antara presisi, recall, dan F1-score (masing-masing 87,23%). Model Random Forest menyusul dengan akurasi 82,35% dan AUC tertinggi sebesar 0,88, sedangkan Logistic Regression menunjukkan performa terendah dengan akurasi 77,65%. Berdasarkan hasil tersebut, model KNN direkomendasikan sebagai metode yang paling efektif untuk deteksi dini penyakit jantung, dengan Random Forest sebagai alternatif yang stabil dan andal.

Penulis yang sesuai:

Adinda Pangestu

Program Studi Informatika Fakultas Ilmu Komputer Universitas Amikom Purwokerto

Email: adindapangestu17@gmail.com

1. PENDAHULUAN

Jantung merupakan organ yang sangat penting dalam tubuh manusia karena berperan sebagai inti dari sistem peredaran darah. Organ ini bekerja tanpa henti untuk memompa darah ke seluruh bagian tubuh melalui jaringan pembuluh darah yang kompleks. Kelainan detak jantung dapat terjadi ketika kecepatannya kurang dari 60 bpm yang dikenal sebagai bradikardia. Selain itu, kelainan detak jantung juga dapat terjadi ketika kecepatannya melebihi 100 bpm yang dikenal sebagai takikardia[1]. Dengan tugas utamanya mengalirkan darah yang kaya akan oksigen dan nutrisi, jantung juga membantu mengangkut limbah hasil metabolisme untuk dibuang dari tubuh. Fungsi vital ini membuat jantung tak tergantikan dalam menjaga kelangsungan hidup serta stabilitas berbagai sistem organ lainnya[2].



Namun, di balik peranannya yang krusial tersebut, jantung termasuk salah satu organ yang paling rentan mengalami gangguan atau kerusakan.

Menurut WHO, penyakit jantung mencakup berbagai gangguan pada jantung dan pembuluh darah yang mengganggu fungsi normalnya, mulai dari ringan hingga berat, dan bisa berakibat fatal jika tidak ditangani[3]. Jenisnya antara lain penyakit arteri koroner, gagal jantung, aritmia, serta kelainan katup jantung. Menjaga kesehatan jantung sangat penting karena dampaknya besar terhadap kualitas hidup[4]. Kondisi ini sangat memprihatinkan dan menunjukkan rendahnya kesadaran masyarakat terhadap pentingnya deteksi dini dan pencegahan penyakit jantung. Banyak orang tidak menyadari gejalanya atau menganggapnya sepele, sehingga penyakit sering terdeteksi saat sudah parah[5]. Pemeriksaan jantung belum menjadi kebiasaan, padahal penyakit ini bisa menyerang siapa saja, termasuk yang tampak sehat. Karena itu, penting untuk rutin memeriksakan diri dan menjalani pola hidup sehat guna mencegah risiko fatal[6]. Machine Learning (ML) merupakan salah satu penerapan dari Kecerdasan Buatan (AI) yang berfokus pada pengembangan sistem yang mampu belajar secara mandiri tanpa perlu diprogram berulang kali. ML memerlukan data (biasanya disebut sebagai data pelatihan) sebagai bagian dari proses pembelajaran sebelum menghasilkan output. Menggunakan algoritma Machine Learning untuk mendiagnosis penyakit kanker diharapkan dapat menghasilkan prediksi yang akurat dan memungkinkan peningkatan dalam deteksi dini serta perawatan yang lebih efektif bagi pasien[7].

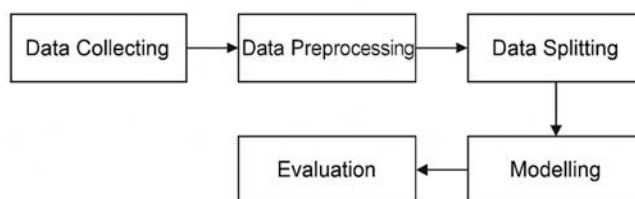
Pada penelitian dengan judul “Implementasi Data Mining Untuk Memprediksi Penyakit Jantung Menggunakan Metode K-Nearest Neighbor dan Logistic Regression” oleh [8] menunjukkan hasil akurasi bahwa metode logistic regression lebih unggul dibanding dengan metode K-Nearest Neighbor (KNN). Pada riset sebelumnya, tingkat akurasi yang dihasilkan masih tergolong rendah dan belum optimal. Oleh karena itu, pada penelitian ini akan dilakukan perbandingan akurasi dari tiga model yang berbeda guna mengetahui model mana yang memiliki performa terbaik dalam mengklasifikasikan data. Kondisi ini sangat memprihatinkan dan menunjukkan rendahnya kesadaran masyarakat terhadap pentingnya deteksi dini dan pencegahan penyakit jantung. Banyak orang tidak menyadari gejalanya atau menganggapnya sepele, sehingga penyakit sering terdeteksi saat sudah parah. Pemeriksaan jantung belum menjadi kebiasaan, padahal penyakit ini bisa menyerang siapa saja, termasuk yang tampak sehat[9]. Karena itu, penting untuk rutin memeriksakan diri dan menjalani pola hidup sehat guna mencegah risiko fatal.[6]. Klasifikasi adalah jenis analisis data yang dapat membantu orang memprediksi label kelas sampel harus diklasifikasikan. Berbagai macam teknik klasifikasi telah diusulkan dalam bidang-bidang seperti pembelajaran mesin, sistem pakar dan statistik[10].

Salah satu solusi potensial untuk mengatasi permasalahan ini adalah dengan memanfaatkan machine learning, yang memiliki kemampuan untuk menganalisis gejala-gejala yang berkaitan dengan penyakit jantung. Data historis pasien dapat digunakan untuk melatih algoritma, sehingga sistem mampu mengidentifikasi pola-pola tersembunyi yang mungkin tidak terdeteksi melalui metode konvensional. Dalam penerapannya, terdapat berbagai metode pemodelan machine learning yang dapat digunakan. Namun, dalam penelitian ini, penulis melakukan perbandingan antara tiga metode, yaitu Logistic Regression, Random Forest, dan K-Nearest Neighbors (KNN). Ketiga metode ini digunakan untuk menyelesaikan masalah klasifikasi biner, di mana model memprediksi probabilitas suatu kejadian target berdasarkan variabel-variabel independen, dengan dua kemungkinan hasil akhir, yaitu 0 atau 1 [11]. Tujuan dari penelitian ini adalah mengembangkan dan membandingkan akurasi beberapa model deteksi penyakit jantung berdasarkan data yang tersedia. Dengan mempertimbangkan keterbatasan sumber data dan cakupan variabel, penelitian ini difokuskan untuk menghasilkan model yang cukup akurat dan dapat menjadi referensi awal bagi tenaga medis dalam mendeteksi risiko penyakit jantung secara dini.

2. METODE

Penelitian ini menunjukkan tahapan dari penelitian ini yang dimana terbagi menjadi 5 tahap yaitu Data Collecting proses pengambilan data, kemudian Data Preprocessing yaitu proses pembersihan data dari missing value, duplicate value dan data outlier, kemudian Data Splitting. Tahap selanjutnya yaitu Modelling menggunakan algoritma

Logistic Regression, Random Forest, dan KNN. Dan terakhir yaitu Evaluasi terhadap model yang digunakan. Jenis data yang digunakan dalam penelitian ini bersifat rasio yang merupakan data yang berbentuk angka sebenarnya[12].



Gambar 1. Tahap Metode Penelitian

2.1. Data Collecting

Data yang dimanfaatkan dalam penelitian ini sama seperti yang dilakukan oleh [4] yang berasal dari informasi pasien penderita penyakit jantung yang diunduh secara daring melalui situs web Kaggle dan link nya <https://www.kaggle.com/code/jabirmuktabir/klasifikasi-penyakit-jantung/input>. Dataset tersebut mencakup berbagai fitur atau atribut sebagai berikut:

1. age
2. sex
3. chest pain type (4 values)
4. resting blood pressure
5. serum cholestoral in mg/dl
6. fasting blood sugar > 120 mg/dl
7. resting electrocardiographic results (values 0,1,2)
8. maximum heart rate achieved
9. exercise induced angina
10. oldpeak = ST depression induced by exercise relative to rest
11. the slope of the peak exercise ST segment
12. number of major vessels (0-3) colored by flourosopy
13. thal: 0 = normal; 1 = fixed defect; 2 = reversable defect

2.2. Data Preprocessing

Proses mengubah data mentah menjadi format yang mudah dipahami. Proses ini dilakukan karena data mentah sering kali dalam bentuk format yang tidak beraturan[13]. Tujuannya agar bisa dijadikan sumber informasi melalui sekumpulan data yang bisa diteruskan untuk diolah datanya[8]. Tahap pra-pemrosesan data pada penelitian ini dilakukan melalui tiga tahapan utama, yaitu:

2.2.1. Pengecekan Missing

Pengecekan missing value, yaitu nilai-nilai kosong yang mungkin terdapat dalam dataset. Nilai kosong dapat mengganggu proses pelatihan karena model tidak dapat memproses data yang tidak lengkap. Untuk memastikan kelengkapan data, digunakan fungsi `.isnull().sum()` yang menjumlahkan nilai kosong di tiap kolom. Hasil pemeriksaan menunjukkan bahwa seluruh kolom dalam dataset telah terisi secara lengkap, sehingga tidak diperlukan penanganan tambahan terkait missing value.

2.2.2. Penghapusan Data Duplikat

Penghapusan *data duplikat* terjadi ketika terdapat dua atau lebih baris yang identik dalam dataset. Keberadaan data yang sama secara berulang dapat menimbulkan bias dalam pelatihan model dan mengurangi keberagaman data. Untuk mengatasi hal ini, digunakan metode `.drop_duplicates()` pada dataset.

2.2.3. Penghapusan Outlier

Proses ini menghilangkan nilai-nilai ekstrem yang berbeda jauh dari data mayoritas. Outlier dapat mengganggu distribusi data dan mempengaruhi performa model, terutama pada algoritma yang sensitif terhadap nilai numerik seperti regresi atau KNN. Penghapusan outlier dilakukan dengan metode IQR (Interquartile Range).

2.2.4. Standarisasi Fitur

Normalisasi fitur menggunakan metode StandardScaler dari pustaka Scikit-learn. Normalisasi ini penting untuk menyamakan skala antar fitur, karena beberapa algoritma seperti K-Nearest Neighbors dan Logistic Regression sangat dipengaruhi oleh perbedaan skala nilai. StandardScaler akan mengubah nilai setiap fitur agar memiliki rata-rata (mean) sebesar 0 dan standar deviasi sebesar 1. Proses ini membantu model dalam memberikan bobot yang adil pada setiap fitur, sehingga prediksi menjadi lebih akurat dan tidak bias terhadap fitur dengan skala lebih besar

2.3. Data Splitting

Pada tahap *data splitting*, dataset dalam penelitian ini dibagi menjadi dua bagian, yaitu data pelatihan (*training set*) dan data pengujian (*testing set*). Pembagian ini dilakukan untuk membangun dan mengevaluasi model secara adil. Adapun proporsi data yang digunakan adalah 70% untuk pelatihan dan 30% untuk pengujian.

2.4. Modelling

Penelitian ini membangun model prediksi penyakit jantung menggunakan tiga algoritma machine learning, yaitu Logistic Regression, Random Forest, dan K-Nearest Neighbors (KNN), dengan data yang diambil dari situs Kaggle. *Logistic Regression* memodelkan probabilitas seseorang terkena penyakit jantung berdasarkan 13 variabel independen dengan fungsi logistik sigmoid[14]. *Random Forest* membangun beberapa pohon keputusan dari subset data untuk meningkatkan akurasi dan mencegah overfitting[15]. Sementara itu, *KNN* mengklasifikasikan data baru berdasarkan mayoritas dari k tetangga terdekat. Ketiga metode ini digunakan untuk mengevaluasi dan membandingkan performa prediksi secara menyeluruh dengan akurasi. Akurasi mengukur seberapa tepat model mengklasifikasikan data ke kelas yang benar. Presisi menunjukkan proporsi prediksi positif yang benar, sedangkan recall mengukur seberapa banyak data positif yang berhasil dikenali model. F1-Score adalah rata-rata harmonis dari presisi dan recall. Keempat metrik ini memiliki rumus perhitungan masing-masing, yaitu rumus (1) hingga (4)[14].

2.5. Evaluation Model

Setelah model dibangun, dilakukan proses evaluasi lanjutan menggunakan learning curve, confusion matrix, dan ROC curve. Learning curve membantu memvisualisasikan perubahan performa model berdasarkan jumlah data pelatihan, serta dapat digunakan untuk mendeteksi overfitting atau underfitting. Confusion matrix menyajikan detail evaluasi melalui jumlah true positive (TP), true negative (TN), false positive (FP), dan false negative (FN). Sedangkan ROC curve menggambarkan kemampuan model dalam membedakan antara true positive rate (TPR) dan false positive rate (FPR), di mana semakin mendekati sudut kiri atas grafik menunjukkan kinerja klasifikasi yang semakin baik.

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

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

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

$$F1 - \text{Score} = 2 \times \frac{\text{Recall} \times \text{Presisi}}{\text{Recall} + \text{Presisi}}$$

Keterangan:

TP = True Positive

TN = True Negative

FP = False Positive

FN = False Negative

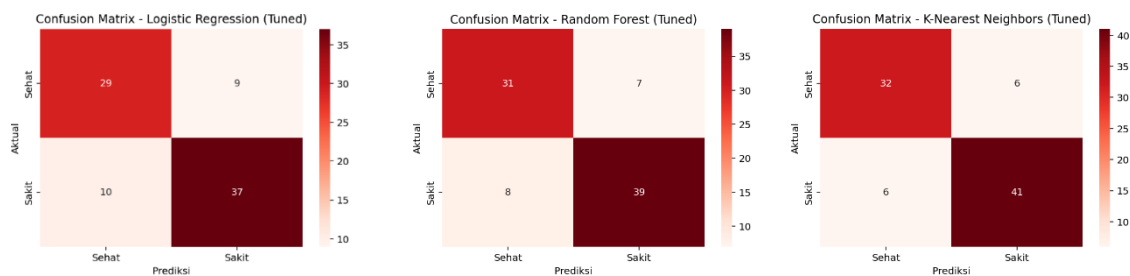
Gambar 2. Rumus Evaluasi

3. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan tiga metode klasifikasi, yaitu Logistic Regression, K-Nearest Neighbors (K-NN), dan Random Forest, untuk memprediksi data yang telah disiapkan. Evaluasi kinerja dari masing-masing model dilakukan dengan menggunakan tiga pendekatan visual dan statistik utama, yaitu Confusion Matrix, ROC Curve, dan Learning Curve. Pada tahap data preprocessing telah dilakukan terdapat hasil yaitu pada pengecekan missing value Semua kolom memiliki nilai lengkap (tidak ada data yang hilang), penghapusan data duplikat: terdapat 723 data duplikat yang dihapus, menyisakan 302 data, penghapusan outlier: menggunakan metode iqr (interquartile range) pada fitur sensitif seperti trestbps, chol, thalach, dan oldpeak, sehingga data bersih menjadi 283 entri. Dan ada tahap data splitting yang telah dilakukan terdapat hasil training set (70%): 198 data dan testing set (30%): 85 data

3.1. Confusion Matrix

Confusion matrix digunakan untuk menggambarkan kinerja model klasifikasi dengan menunjukkan jumlah prediksi yang benar dan salah untuk masing-masing kelas.



Gambar 3. Confusion Matrix

3.1.1. Logistic Regression (Tuned)

Pada logistic regression menghasilkan Akurasi: 77,65%, Presisi: 80,43%, Recall: 78,72%, F1-Score 79,57% dan TP: 37, TN: 29, FP: 9, FN: 10. Analisis model regresi logistik menunjukkan hasil terendah di antara ketiga algoritma yang ada. Berdasarkan confusion matrix

- True Positive (TP) = 37: Pasien yang sebenarnya mengalami penyakit dapat terdeteksi dengan baik.
- False Negative (FN) = 10: Pasien yang seharusnya dicatat sebagai sakit malah teridentifikasi sebagai sehat, yang bisa menyebabkan penundaan dalam perawatan medis.
- False Positive (FP) = 9: Pasien yang sehat keliru ditempatkan dalam kategori sakit, yang mungkin mengarah pada tindakan medis yang tidak perlu.
- True Negative (TN) = 29: Pasien yang sehat berhasil dicatat dengan benar.

Model ini memiliki akurasi pelatihan tinggi (84,85%) namun menurun saat pengujian (77,65%) hingga menunjukkan keterbatasan generalisasi. Cocok untuk edukasi atau sistem yang butuh interpretasi, tapi kurang ideal untuk diagnosis yang butuh akurasi tinggi.

3.1.2. K-Nearest Neighbors (Disesuaikan)

Algoritma Nearest Neighbor (K-NN) merupakan algoritma klasifikasi berdasarkan kedekatan jarak suatu data dengan data yang lain. Pada algoritma K-NN, data berdimensi q , jarak dari data tersebut ke data yang lain dapat dihitung. Nilai jarak inilah yang digunakan sebagai nilai kedekatan/kemiripan antara data uji dengan data latih. Nilai K pada K-NN berarti K-data terdekat dari data uji, hasilnya Akurasi: 85,88%, Presisi: 87,23%, Recall: 87,23%, F1-Score: 87,23% dan TP: 41, TN: 32, FP: 6, FN: 6. Analisis:

- Model KNN menunjukkan hasil terbaik di antara semua model yang ada.
- Tingkat recall yang paling tinggi mencapai 87,23%, menunjukkan kemampuan model dalam mengidentifikasi pasien yang benar-benar menderita penyakit sangat baik.
- Jumlah False Negative (FN) hanya sebanyak 6, yang berarti kemungkinan untuk gagal mendeteksi pasien yang sakit sangat kecil. Ini merupakan poin penting dalam dunia medis karena kesalahan FN bisa berisiko besar.
- Jumlah False Positive (FP) juga terendah (6) jika dibandingkan dengan model-model lain.
- Confusion matrix menunjukkan klasifikasi yang paling merata serta tingkat kesalahan yang terendah.

KNN memperoleh akurasi pelatihan 100% (yang mengindikasikan overfitting), tetapi hasil pengujian tetap unggul dengan 85,88%, menjadikan model ini tetap yang terbaik. Namun, penting untuk diingat bahwa KNN mungkin menghadapi tantangan ketika berhadapan dengan dataset besar karena biaya komputasi yang tinggi dan ketergantungan pada pemilihan nilai parameter k .

3.1.3. Random Forest (Tuned)

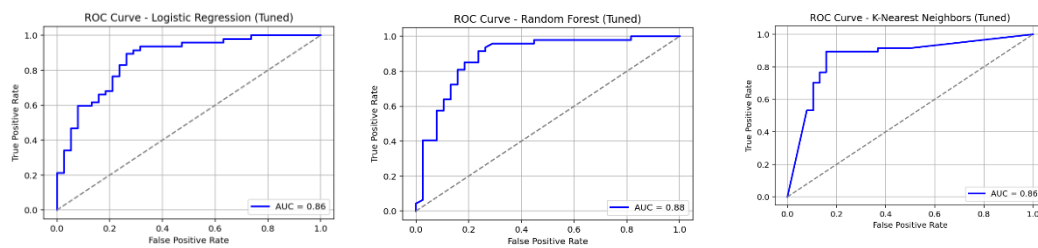
Pada Random Forest menghasilkan Akurasi: 82,35%, Presisi: 84,78%, Recall: 82,97%, F1-Score: 83,87%, TP: 39, TN: 31, FP: 7, FN: 8, Analisis:

- Random Forest menunjukkan kinerja yang seimbang dan kuat dalam mendeteksi pasien sakit maupun sehat. Pada confusion matrix:
- False Positive dan False Negative relatif rendah, yang menunjukkan risiko kesalahan diagnosis lebih kecil dibandingkan Logistic Regression.
- True Negative = 31 berarti kemampuan model dalam mendeteksi pasien sehat lebih baik dibandingkan Logistic Regression.
- Overfitting minimal: meskipun akurasi training mencapai 100%, akurasi testing masih terjaga tinggi (82,35%), mencerminkan kemampuan generalisasi yang baik.

Random Forest sangat efektif menangani hubungan non-linear dan outlier dalam data. Model ini dapat diandalkan untuk diterapkan dalam sistem pendukung keputusan klinis (clinical decision support system), terutama bila dibutuhkan kestabilan dan interpretasi fitur yang lebih kuat.

3.2. Roc Curve

ROC curve (Receiver Operating Characteristic) digunakan untuk mengevaluasi trade-off antara True Positive Rate (Sensitivity) dan False Positive Rate (1-Specificity).



Gambar 4. ROC Curve

3.2.1 Logistic Regression

Logistic Regression adalah model linier yang banyak digunakan untuk klasifikasi biner karena kesederhanaannya dan kemampuannya untuk memberikan interpretasi yang jelas terhadap koefisien. Akurasi pada data uji sebesar 77,65%, menunjukkan performa moderat, Nilai recall 78,72% menandakan bahwa model masih cukup banyak gagal mendeteksi pasien sakit (FN = 10). Presisi = 80,43%, artinya 19,57% dari pasien yang diprediksi sakit sebenarnya sehat. Dari ROC curve, AUC sebesar 0,86 menunjukkan model cukup baik dalam membedakan antara kelas sakit dan sehat, meskipun tidak optimal. Model ini mudah dijelaskan dan cocok untuk baseline, namun performanya kurang ideal untuk aplikasi medis yang membutuhkan deteksi yang lebih akurat.

3.2.2 K-Nearest Neighbors

KNN adalah model non-parametrik yang mengklasifikasikan data berdasarkan kedekatan jarak dari data tetangga terdekat. Akurasi tertinggi (85,88%), menunjukkan kemampuan generalisasi yang sangat baik meskipun akurasi training 100%. Recall 87,23% dan Presisi 87,23%, menandakan model sangat mampu mengenali pasien yang benar-benar sakit sekaligus menghindari kesalahan pada pasien sehat. FN dan FP hanya masing-masing 6, terendah dari semua model. AUC sebesar 0,86, sama seperti Logistic Regression namun dengan akurasi yang lebih tinggi. KNN adalah model paling efektif dalam mendeteksi penyakit jantung berdasarkan hasil ini, sangat cocok untuk aplikasi prediktif pada data yang bersih dan dengan jumlah fitur yang tidak terlalu besar.

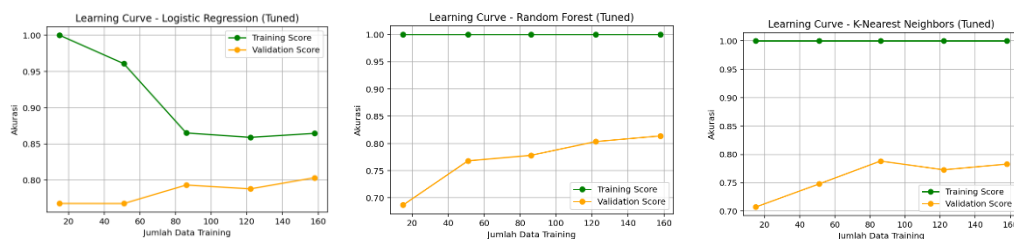
3.2.3. Random Forest

Random Forest merupakan model ensemble yang membangun banyak decision tree dan menggabungkan hasilnya untuk meningkatkan akurasi dan mengurangi overfitting. Akurasi testing sebesar 82,35%, lebih tinggi dari Logistic Regression. Recall 82,97%, dengan hanya 8 pasien sakit tidak terdeteksi. Presisi tinggi (84,78%), sehingga

model jarang memberikan prediksi salah pada pasien sehat. AUC tertinggi di antara ketiga model, yakni 0,88, mengindikasikan kemampuan klasifikasi yang sangat baik. Model ini seimbang dan kuat, ideal digunakan dalam sistem deteksi dini karena mampu meminimalkan kesalahan diagnosis dengan generalisasi yang baik.

3.3 Learning Curve

Learning curve menggambarkan hubungan antara ukuran data pelatihan dan kinerja model terhadap data pelatihan dan validasi.



Gambar 5. Learning Curve

3.3.1 Random Forest

Memiliki kurva pelatihan dan validasi yang hampir sejajar dan tinggi, menandakan model yang stabil, tidak overfit, serta memiliki kemampuan generalisasi yang sangat baik. Pada grafik learning curve, tampak bahwa akurasi training konstan di 1.0, menunjukkan model menghafal data pelatihan (overfitting). akurasi validasi meningkat dari 0.69 ke 0.82, tapi tetap terdapat celah signifikan antara training dan validasi. Hasil Evaluasi Akurasi (training): 100%, Akurasi (testing): 82.35%, Presisi: 0.848, Recall: 0.830, F1-Score: 0.839.

Model memiliki performa testing yang baik meskipun overfitting. Nilai F1 tinggi menunjukkan keseimbangan antara presisi dan recall. Confusion matrix menunjukkan kesalahan yang lebih kecil dibanding Logistic Regression: FP: 7 dan FN: 8. Random Forest menunjukkan performa yang kuat, tetapi terlalu kompleks terhadap data pelatihan. Hal ini bisa diatasi dengan melakukan pruning atau pengaturan parameter lebih lanjut (seperti max_depth atau min_samples_leaf).

3.3.2 Logistic Regression

Pada grafik learning curve, tampak bahwa akurasi training menurun dari 1.0 ke sekitar 0.86 seiring bertambahnya jumlah data pelatihan, akurasi validasi meningkat secara perlahan dari sekitar 0.77 ke 0.80. Hal ini menunjukkan bahwa model mengalami sedikit underfitting, tetapi stabilitas model cukup baik, mengingat gap antara kurva training dan validasi relatif kecil. Hasil Evaluasi akurasi (training): 84.85%, akurasi (testing): 77.65%, Presisi 0.804, Recall 0.787, F1-Score: 0.796

Presisi dan recall relatif seimbang, yang penting dalam konteks medis agar model tidak terlalu sering salah mendeteksi kondisi kritis (false negative rendah). Confusion matrix menunjukkan FP 9 dan FN 10, yang berarti masih terdapat kesalahan prediksi pada kedua kelas. Model ini cenderung stabil, cocok untuk situasi di mana interpretabilitas dan generalisasi penting, meskipun akurasi tidak sebaik model lain.

3.3.3 K-Nearest Neighbors

K-NN menunjukkan adanya overfitting, ditandai dengan perbedaan yang signifikan antara skor pelatihan dan validasi, di mana model sangat baik pada data pelatihan namun kurang generalisasi ke data validasi. Pada grafik learning curve, tampak bahwa sama seperti Random Forest, KNN memiliki akurasi training sempurna (1.0). Akurasi validasi meningkat dari 0.71 ke 0.79, namun lebih fluktuatif. Hasil Evaluasi akurasi (training): 100%, akurasi (testing): 85.88% (tertinggi dari semua model), Presisi: 0.872, Recall: 0.872, F1-Score: 0.872

KNN memiliki performa prediksi terbaik di data testing. Confusion matrix menunjukkan FP 6 dan FN 6, kesalahan paling sedikit di antara ketiga model. KNN sangat efektif untuk dataset ini. Namun, overfitting masih menjadi perhatian. Performa tinggi bisa dijaga dengan pendekatan validasi silang dan penyesuaian nilai k yang optimal.

Tabel 1. Tabel Model

Model	Akurasi	AUC Score	TN	FP	FN	TP
Logistic Regression	77.65%	0.7751	29	9	10	37
Random Forest	82.35%	0.8228	31	7	8	39
K-Nearest Neighbors	85.88%	0.8572	32	6	6	41

4. KESIMPULAN

Hasil analisis dari prediksi penyakit jantung menggunakan tiga metode Logistic Regression, Random Forest, dan K-Nearest Neighbors (KNN) dapat disimpulkan bahwa Model K-Nearest Neighbors (KNN) menunjukkan performa terbaik dalam mendeteksi penyakit jantung dengan akurasi tertinggi sebesar 85,88%, serta presisi, recall, dan F1-score yang sama tinggi (87,23%). KNN juga mencatat jumlah True Positive (TP) tertinggi dan False Negative (FN) terendah, menandakan kemampuannya mengenali pasien yang benar-benar sakit dengan kesalahan minimal. Di posisi kedua, Random Forest meraih akurasi 82,35% dengan performa yang stabil, didukung TP dan TN yang tinggi serta kesalahan klasifikasi yang rendah, menunjukkan generalisasi yang baik. Sementara itu, Logistic Regression memiliki akurasi terendah, yakni 77,65%, dengan jumlah FP dan FN yang lebih tinggi, sehingga kinerjanya kurang optimal dalam menangani kompleksitas data penyakit jantung.

Berdasarkan hasil evaluasi, disarankan untuk menggunakan model K-Nearest Neighbors (KNN) dalam sistem deteksi penyakit jantung karena memberikan akurasi dan kinerja terbaik. Namun, untuk penggunaan pada skala besar atau sistem real-time, perlu dipertimbangkan optimasi model atau pemilihan algoritma lain yang lebih efisien secara komputasi, seperti Random Forest. Selain itu, pemilihan parameter seperti nilai k pada KNN dan jumlah pohon pada Random Forest juga perlu disesuaikan agar model dapat bekerja lebih optimal. Penggunaan Logistic Regression sebaiknya dibatasi pada kasus sederhana atau sebagai model baseline, karena kinerjanya kurang memadai untuk kasus penyakit jantung yang kompleks.

REFERENSI

- [1] H. A. Supriyatna and Z. Yuwaldi Away, "Desain Sistem Internet of Things (IoT) Untuk Pemantauan dan Prediksi Gejala Serangan Jantung," *J. Komputer, Inf. Teknol. dan Elektro*, vol. 4, no. 1, pp. 31–39, 2019, [Online]. Available: <https://jurnal.usk.ac.id/kitektro/article/view/12886>
- [2] M. Lestari, "Penerapan Algoritma Klasifikasi Nearest Neighbor (K-NN) untuk Mendeteksi Penyakit Jantung," *Fakt. Exacta*, vol. 7, no. September 2010, pp. 366–371, 2014.
- [3] L. Ghani and H. Novriani, "Dominant Risk Factors for Coronary Heart Disease in Indonesia," *Bul. Penelit. Kesehat.*, vol. 44, no. 3, pp. 153–164, 2016.
- [4] R. Hidayat, Y. S. Sy, T. Sujana, M. Husnah, and H. T. Saputra, "Implementasi Machine Learning Untuk Prediksi Penyakit Jantung Menggunakan Algoritma Support Vector Machine," vol. 5, no. 2, pp. 161–168, 2024.
- [5] J. Ginting, R. Ginting, and H. Hartono, "Deteksi Dan Prediksi Penyakit Diabetes Melitus Tipe 2 Menggunakan Machine Learning (Scooping Review)," *J. Keperawatan Prior.*, vol. 5, no. 2, pp. 93–105, 2022, doi: 10.34012/jukep.v5i2.2671.
- [6] T. Indriyani, F. F. Rozi, M. Hakimah, N. Fakhur, and R. R. Muhima, "Metode Decision Tree C4.5 untuk Klasifikasi penyakit Jantung," 2024.
- [7] S. Heristian, "Perbandingan Algoritma Machine Learning pada Klasifikasi Penyakit Jantung," *J. Infortech*, vol. 6, no. 1, pp. 46–51, 2024, doi: 10.31294/infortech.v6i1.21888.
- [8] D. Sitanggang, N. Nicholas, V. Wilson, A. R. A. Sinaga, and A. D. Simanjuntak, "Implementasi Data Mining Untuk Memprediksi Penyakit Jantung Menggunakan Metode K-Nearest Neighbor Dan Logistic Regression," *J. Tek. Inf. dan Komput.*, vol. 5, no. 2, p. 493, 2022, doi: 10.37600/tekinkom.v5i2.698.
- [9] A. M. Soesanto, "Penyakit Jantung Katup di Indonesia: masalah yang hampir terlupakan," *Indones. J. Cardiol.*, vol. 33, no. 4, pp. 205–8, 2012, doi: 10.30701/ijc.v33i4.20.
- [10] A. Riski, "Analisis Komparasi Algoritma Klasifikasi Data Mining Untuk Prediksi Penderita Penyakit Jantung," *J. Tek. Inform. Kaputama*, vol. 3, no. 1, pp. 22–28, 2019, [Online]. Available: <https://jurnal.kaputama.ac.id/index.php/JTIK/article/view/141/156>
- [11] J. Ilmiah and T. Informasi, "Prediksi Pasien Terindikasi Penyakit Jantung Menggunakan Metode Logistic Regression," vol. 4, no. 1, pp. 19–23, 2024.
- [12] I. Alhabib, "Komparasi Metode Deep Learning, Naïve Bayes Dan Random Forest Untuk Prediksi Penyakit Jantung," *INFORMATICS Educ. Prof. J. Informatics*, vol. 6, no. 2, p. 176, 2022, doi: 10.51211/itbi.v6i2.1881.



-
- [13] D. L. Mihardja and H. Siswoy, "Prevalensi dan faktor determinan penyakit jantung di Indonesia," *Bul. Penelit. Kesehat.*, vol. 37, no. 3, pp. 142–159, 2019.
- [14] C. Raras, A. Widiawati, L. Nurazizah, and I. R. Yunita, "Implementasi Algoritma Logistic Regression pada Pembuatan Website Sederhana untuk Prediksi Penyakit Jantung," *Infotekmesin*, vol. 15, no. 01, pp. 117–122, 2024, doi: 10.35970/infotekmesin.v15i1.2048.
- [15] N. S. S. Silmi Ath Thahirah Al Azhima, D. Darmawan, N. Fahmi Arief Hakim, I. Kustiawan, M. Al Qibtiya, "Hybrid Machine Learning Model Untuk Memprediksi Penyakit," *J. Teknol. Terpadu*, vol. 8, no. 1, pp. 40–46, 2022.