

Supervised Machine Learning Algorithms untuk Klasifikasi Penyakit Jantung

Denny Fajar Riadi*, Wahyu Aji Eko Prabowo

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

Email: ¹111202214084@mhs.dinus.ac.id, ²prabowo@dsn.dinus.ac.id

Email Penulis Korespondensi: 111202214084@mhs.dinus.ac.id

Submitted 07-01-2026; Accepted 24-02-2026; Published 28-02-2026

Abstrak

Penyakit jantung merupakan salah satu penyebab utama kematian di dunia sehingga diperlukan metode prediksi yang akurat untuk mendukung deteksi dini dan pengambilan keputusan medis. Penelitian ini bertujuan menganalisis dan membandingkan kinerja tiga algoritma machine learning terawasi, yaitu K-Nearest Neighbors (KNN), Support Vector Machine (SVM), dan Random Forest (RF) dalam klasifikasi penyakit jantung menggunakan dataset Cleveland Heart Disease yang terdiri dari 303 data pasien dengan 13 fitur klinis. Tahapan penelitian meliputi pra-pemrosesan data, pembagian dataset menjadi data latih (80%) dan data uji (20%), pelatihan model, serta optimasi hyperparameter menggunakan GridSearchCV dengan 5-fold cross-validation. Setelah proses optimasi, dilakukan tahap prediksi menggunakan data uji dan evaluasi performa model untuk menilai kemampuan generalisasi. Evaluasi performa dilakukan menggunakan accuracy, precision, recall, F1-score, AUC-ROC, dan confusion matrix. Hasil eksperimen menunjukkan bahwa KNN dan Random Forest memperoleh nilai accuracy tertinggi sebesar 90,16%. Model KNN memiliki nilai recall sebesar 1,0000 yang menunjukkan sensitivitas sempurna terhadap kelas positif, sedangkan Random Forest memberikan performa yang lebih seimbang antara precision dan recall dengan nilai AUC tertinggi sebesar 0,9481. Berdasarkan hasil tersebut, KNN dinilai paling sesuai untuk kebutuhan skrining medis karena mampu mendeteksi seluruh pasien dengan penyakit jantung tanpa menghasilkan false negative. Penelitian ini diharapkan dapat menjadi referensi dalam penerapan machine learning berbasis data klinis sebagai alat bantu deteksi dini penyakit jantung.

Kata Kunci: Klasifikasi Penyakit Jantung; Pembelajaran Mesin; K-Nearest Neighbors; Support Vector Machine; Random Forest

Abstract

Heart disease is one of the leading causes of death worldwide, requiring accurate predictive methods to support early detection and clinical decision making. This study aims to analyze and compare the performance of three supervised machine learning algorithms, namely K-Nearest Neighbors (KNN), Support Vector Machine (SVM), and Random Forest (RF), in classifying heart disease using the Cleveland Heart Disease dataset consisting of 303 patient records with 13 clinical features. The research stages include data preprocessing, splitting the dataset into 80% training data and 20% testing data, model training, and hyperparameter optimization using GridSearchCV with 5-fold cross-validation. After optimization, prediction was performed using test data followed by performance evaluation to assess generalization ability. Model performance was evaluated using accuracy, precision, recall, F1-score, AUC-ROC, and confusion matrix. The results show that KNN and Random Forest achieved the highest accuracy of 90.16%. The KNN model obtained a recall value of 1.0000, indicating perfect sensitivity in detecting positive cases, while Random Forest demonstrated a more balanced performance between precision and recall with the highest AUC value of 0.9481. Based on these findings, KNN is considered the most suitable model for medical screening purposes, as it successfully detected all positive heart disease patients without producing false negatives. This study is expected to serve as a reference for implementing clinical databased machine learning as a decision support tool for early heart disease detection.

Keywords: Heart Disease Classification; Machine Learning; K-Nearest Neighbors; Support Vector Machine; Random Forest

1. PENDAHULUAN

Penyakit jantung merupakan salah satu masalah kesehatan yang berdampak besar karena dapat berkembang dari faktor risiko yang umum (misalnya tekanan darah, kadar gula, kolesterol, kebiasaan hidup) menjadi kondisi berat seperti gagal jantung atau komplikasi lain. Dampak ini mendorong kebutuhan akan pendekatan prediksi yang lebih cepat dan konsisten sebagai dukungan keputusan klinis maupun edukasi pencegahan [1]. Data dari Organisasi Kesehatan Dunia (WHO) menunjukkan bahwa penyakit jantung menyumbang sepertiga dari semua kematian di seluruh dunia. Sistem perawatan kesehatan di seluruh dunia menghadapi tantangan yang signifikan karena meningkatnya angka kematian dan penyakit gagal jantung. Sekitar 17,9 juta orang meninggal setiap tahun di seluruh dunia karena penyakit gagal jantung, yang lebih sering terjadi di Asia [2]. Sejumlah penelitian di Indonesia menunjukkan bahwa prediksi penyakit jantung berbasis data klinis dapat dilakukan menggunakan berbagai algoritma *supervised learning* dengan hasil yang cukup menjanjikan, baik untuk kasus penyakit jantung secara umum maupun kasus gagal jantung. Hal ini terlihat dari studi komparasi algoritma *machine learning* dan ensemble dalam prediksi penyakit jantung [1], penelitian perbandingan metode untuk mendeteksi penyakit jantung [3].

Secara konvensional, penilaian risiko penyakit jantung dilakukan melalui interpretasi klinis dan pemeriksaan penunjang, tetapi pada praktiknya sering menghadapi keterbatasan sumber daya, variasi antar pengamat, serta kebutuhan interpretasi data yang banyak. Karena itu, penerapan *machine learning* menjadi alternatif yang relevan, khususnya pada data tabular klinis, karena mampu mempelajari pola dari data historis pasien untuk menghasilkan prediksi kelas secara terukur [4]. Banyak penelitian lokal menguji pendekatan ini dengan membandingkan beberapa algoritma populer seperti KNN, SVM, *Random Forest*, *Logistic Regression*, *Naive Bayes*, hingga *Gradient Boosting*. Contohnya, studi komparasi

KNN dan *Random Forest* [5], penelitian penggunaan *Random Forest* untuk prediksi penyakit jantung [6], serta evaluasi *Random Forest* vs *Gradient Boosting* untuk klasifikasi penyakit jantung [7].

Di antara algoritma yang sering digunakan, KNN dipilih karena sederhana dan efektif untuk klasifikasi berbasis kedekatan, sementara SVM dikenal kuat pada pemisahan kelas dengan margin optimal, dan *Random Forest* sering unggul pada data non-linear karena bersifat ensemble dan relatif stabil [8]. Sejumlah studi Indonesia membandingkan SVM dengan algoritma lain seperti *Logistic Regression* maupun *Decision Tree* [9], membahas penerapan SVM untuk prediksi penyakit jantung [10], serta menguji *Random Forest* baik pada prediksi penyakit jantung maupun gagal jantung [6], [11], [12]. Selain itu, terdapat juga penelitian yang menggabungkan atau membuat model hibrida berbasis SVM dan *Random Forest* untuk meningkatkan kinerja klasifikasi gagal jantung [13].

Namun demikian, hasil penelitian prediksi penyakit jantung pada penelitian-penelitian terlebih dahulu sering dipengaruhi oleh perbedaan rancangan eksperimen, seperti langkah pra-pemrosesan, pemilihan fitur, strategi pembagian data, serta penalaan *hyperparameter*. Beberapa penelitian menekankan pentingnya proses komparasi yang konsisten agar hasil antar model lebih adil, misalnya pada studi komparatif berbagai model *data mining* untuk deteksi penyakit jantung dengan seleksi fitur [14] dan penelitian yang menggunakan seleksi fitur seperti *Recursive Feature Elimination with Cross-validation* (RFECV) untuk deteksi gagal jantung [15],[16]. Selain itu, penentuan parameter model juga menjadi faktor penting karena konfigurasi yang berbeda dapat menghasilkan performa yang berbeda, sebagaimana diindikasikan oleh kajian implementasi *hyperparameter tuning* pada jurnal ML lokal [17].

Sebagian besar penelitian terdahulu masih memiliki keterbatasan tertentu. Banyak studi hanya membandingkan dua algoritma atau menggunakan metrik evaluasi yang terbatas, seperti akurasi saja, tanpa mempertimbangkan *precision*, *recall*, dan *F1-score* secara komprehensif [1], [3], [18], [19]. Padahal, dalam konteks medis, kesalahan klasifikasi, khususnya *False Negative* (pasien sakit yang diprediksi sehat), dapat menimbulkan risiko yang sangat besar [11], [13], [20], [21]. Agusyl dan Firmansyah [6] melaporkan bahwa *Random Forest* mampu menghasilkan akurasi tinggi dan performa yang stabil pada data klinis penyakit jantung karena kemampuannya menangani hubungan non-linear antar fitur [6]. Ridwan dkk. [7] juga menegaskan bahwa *Random Forest* cenderung lebih konsisten dibandingkan algoritma lain pada data kesehatan yang kompleks [7]. Di sisi lain, Fredilio dkk. [12] menunjukkan bahwa KNN memiliki sensitivitas yang tinggi dalam mendeteksi kasus gagal jantung, meskipun performanya sangat bergantung pada pemilihan parameter dan pra-pemrosesan data [12]. Sementara itu, Hidayat dkk. [10] menyatakan bahwa SVM mampu memisahkan kelas penyakit jantung dengan baik, namun memerlukan penyesuaian parameter kernel dan regularisasi yang tepat agar terhindar dari *overfitting* [10]. Selain itu, analisis terhadap kemampuan generalisasi model dan potensi *overfitting* sering kali belum dibahas secara mendalam. Beberapa penelitian juga menggunakan konfigurasi default algoritma tanpa melakukan optimasi parameter yang memadai, sehingga performa model yang diperoleh belum tentu optimal [17], [22]. Hal ini menunjukkan adanya celah penelitian (*research gap*) yang masih perlu dijawab.

Sebagian besar penelitian terdahulu masih berfokus pada metrik evaluasi berbasis satu nilai ambang (*threshold*) seperti akurasi, sehingga belum sepenuhnya mencerminkan kemampuan model dalam membedakan kelas secara menyeluruh. Pendekatan tersebut berpotensi menghasilkan penilaian yang bias, khususnya pada konteks medis yang sensitif terhadap kesalahan klasifikasi. Oleh karena itu, terdapat celah penelitian (*research gap*) dalam pemanfaatan evaluasi berbasis *Receiver Operating Characteristic* (ROC) dan *Area Under the Curve* (AUC) untuk menilai kemampuan diskriminatif model secara lebih stabil dan independen terhadap satu nilai *threshold* tertentu. Penelitian ini diharapkan dapat mengisi celah tersebut dengan menyajikan evaluasi performa yang lebih komprehensif sebagai dasar pemilihan model yang tepat untuk klasifikasi penyakit jantung.

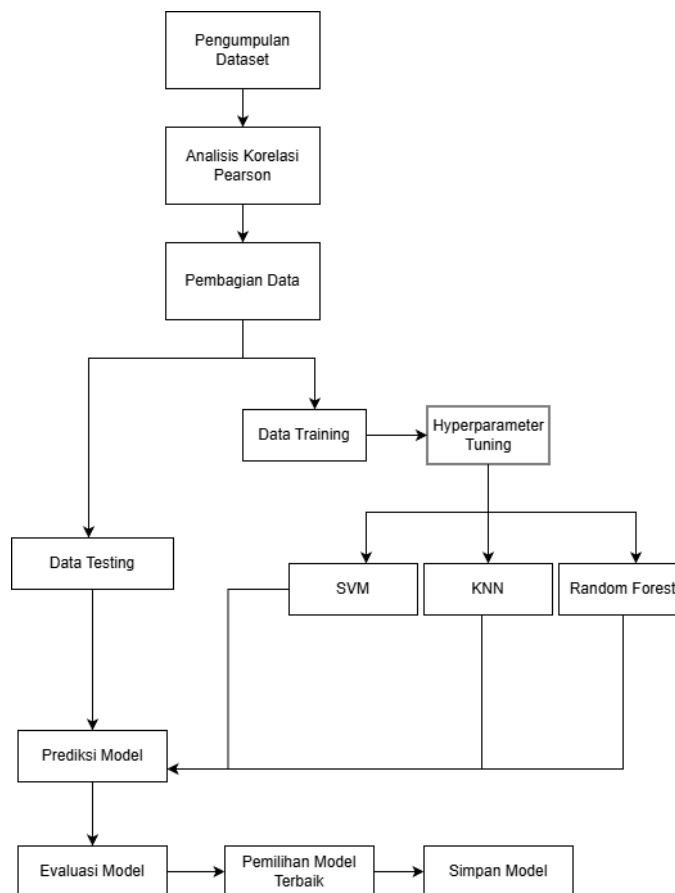
Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk melakukan analisis dan perbandingan performa tiga algoritma supervised *machine learning*, yaitu *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF), dalam klasifikasi penyakit jantung menggunakan *dataset* medis. Evaluasi performa tidak hanya didasarkan pada akurasi, tetapi juga menggunakan metrik yang lebih komprehensif, meliputi *precision*, *recall*, dan *F1-score*, untuk memperoleh gambaran kinerja model yang lebih seimbang. Selain itu, penelitian ini menerapkan skema pelatihan dan pengujian yang sistematis guna memastikan kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Gambar 1 menggambarkan alur tahapan penelitian yang dimulai dari proses pengumpulan *dataset* Heart Disease Cleveland yang diperoleh dari Kaggle dan bersumber dari UCI Machine Learning Repository. *Dataset* ini terdiri dari 303 sampel pasien dengan 13 fitur klinis dan 1 target yang merepresentasikan kondisi penyakit jantung. *Dataset* Cleveland banyak digunakan dalam penelitian prediksi penyakit jantung karena fitur klinisnya relevan dan telah menjadi standar pembanding pada berbagai studi klasifikasi penyakit jantung [1], [6], [23]. Setelah *dataset* dikumpulkan, dilakukan tahap preprocessing data yang meliputi pembersihan data dan normalisasi fitur agar sesuai dengan algoritma *machine learning* yang digunakan [14], [19]. Selanjutnya, *dataset* dibagi menjadi data training (80%) dan data testing (20%) untuk melatih

model dan menguji kemampuan generalisasi. Data *testing* disimpan secara terpisah dan tidak pernah digunakan selama proses pelatihan maupun *hyperparameter tuning*, sehingga benar-benar terisolasi dari model. Langkah ini dilakukan untuk mencegah terjadinya *data leakage* dan memastikan bahwa evaluasi performa model merepresentasikan kemampuan generalisasi terhadap data yang belum pernah dilihat sebelumnya. Pada tahap pelatihan model, digunakan tiga algoritma *supervised learning*, yaitu *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF). Algoritma tersebut dipilih karena kemampuannya dalam menangani data kesehatan dan tugas klasifikasi, serta telah banyak digunakan dalam penelitian prediksi penyakit jantung [1], [10], [12], [24]. Setelah model dasar dibangun, dilakukan *hyperparameter tuning* menggunakan *Grid Search* dan *Cross-validation* untuk mengoptimalkan kinerja model [7], [17], [22]. Model yang telah dioptimalkan kemudian dievaluasi menggunakan metrik akurasi, *precision*, *recall*, *F1-score*, dan *confusion matrix* guna menilai performa dan kemampuan generalisasi model. Model dengan performa terbaik selanjutnya dipilih sebagai model akhir yang diharapkan dapat digunakan sebagai alat bantu skrining dini penyakit jantung [14], [16], [20].



Gambar 1. Tahapan Penelitian

2.2 Data Penelitian

Penelitian ini menggunakan *dataset* Heart Disease Cleveland yang berasal dari UCI Machine Learning Repository dan diperoleh melalui platform Kaggle [25]. *Dataset* ini bertujuan untuk memprediksi ada atau tidaknya penyakit jantung berdasarkan sejumlah fitur klinis pasien. *Dataset* Cleveland merupakan salah satu *dataset* standar yang paling banyak digunakan dalam penelitian prediksi penyakit jantung karena bersifat terbuka dan mudah direplikasi oleh peneliti lain [3], [6], [23]. *Dataset* ini terdiri dari 303 data pasien dengan 13 fitur yang digunakan sebagai variabel prediktor dan 1 target yang merepresentasikan kondisi penyakit jantung [25]. Variabel target umumnya dikodekan secara biner, di mana nilai 0 menunjukkan pasien tidak terindikasi penyakit jantung dan nilai 1 menunjukkan pasien terindikasi penyakit jantung. Fitur-fitur dalam *dataset* ini mencakup data klinis seperti usia, jenis kelamin, tekanan darah, kadar kolesterol, denyut jantung maksimum, dan hasil pemeriksaan medis lainnya yang relevan untuk tugas klasifikasi [1], [12], [24].

Tabel 1 menunjukkan ringkasan *dataset* yang digunakan dalam penelitian ini. *Dataset* Cleveland dipilih karena telah banyak digunakan dalam penelitian sebelumnya sebagai acuan evaluasi kinerja algoritma *machine learning* pada kasus penyakit jantung. Beberapa penelitian menunjukkan bahwa *dataset* ini efektif digunakan untuk membandingkan performa berbagai algoritma klasifikasi seperti KNN, SVM, dan *Random Forest*, sehingga sesuai dengan tujuan penelitian ini [1], [18], [23].

Tabel 1. Sampel *Dataset*

Age	Sex	CP	Trestbps	Chol	Fbs	Restecg	Thalach	Exang	Oldpeak	Slope	Ca	Thal	Target
63	1	0	145	233	1	2	150	0	2.3	2	0	2	0
67	1	3	160	286	0	2	108	1	1.5	1	3	1	1
67	1	3	120	229	0	2	129	1	2.6	1	2	3	1
37	1	2	130	250	0	0	187	0	3.5	2	0	1	0
41	0	1	130	204	0	2	172	0	1.4	0	0	1	0

Tabel 1 menyajikan representasi sebagian data dari *dataset* Heart Disease Cleveland yang digunakan dalam penelitian ini [25]. Tabel tersebut menampilkan sejumlah fitur klinis utama yang relevan dalam klasifikasi penyakit jantung, antara lain usia (Age), jenis kelamin (Sex), tipe nyeri dada (Chest Pain/CP), tekanan darah istirahat (Trestbps), kadar kolesterol serum (Chol), gula darah puasa (Fbs), hasil elektrokardiografi istirahat (Restecg), denyut jantung maksimum (Thalach), induksi angina akibat olahraga (Exang), depresi segmen ST (Oldpeak), kemiringan segmen ST (Slope), jumlah pembuluh darah utama (Ca), dan hasil thallium stress test (Thal). Secara keseluruhan, *dataset* Cleveland terdiri dari 13 fitur yang digunakan sebagai variabel input dan 1 target yang menunjukkan kondisi penyakit jantung.

2.3 Data Penelitian

Tabel 2 menyajikan proporsi pembagian *dataset*. Pada tahap ini, *dataset* dibagi menjadi dua bagian utama, yaitu data latih dan data uji. Data latih digunakan untuk melatih model agar mampu mempelajari pola dari data, sedangkan data uji digunakan untuk mengevaluasi kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya (*unseen*).

Tabel 2. Pembagian *Dataset*

Jenis Data	Jumlah Data	Persentase (%)
Data Latih	242	80
Data Uji	61	20
Total	303	100

Sebagaimana ditunjukkan pada Tabel 2, sebanyak 80% data dialokasikan sebagai data latih, dan 20% sisanya sebagai data uji untuk evaluasi akhir model. Proses pembagian ini dilakukan secara acak, namun menggunakan parameter *random_state* tertentu agar hasilnya dapat direproduksi. Teknik pembagian semacam ini umum digunakan dalam pembelajaran mesin untuk memastikan bahwa model dapat mengenali pola dari data pelatihan dan mampu melakukan generalisasi terhadap data baru.

2.4 Hyperparameter Tuning

Hyperparameter tuning merupakan proses pencarian kombinasi parameter optimal yang bertujuan untuk meningkatkan kinerja model *machine learning*. Berbeda dengan parameter yang dipelajari langsung dari data selama proses pelatihan, *hyperparameter* ditentukan sebelum pelatihan dimulai dan sangat memengaruhi kemampuan model dalam melakukan generalisasi terhadap data baru. Oleh karena itu, pemilihan *hyperparameter* yang tepat menjadi langkah penting dalam membangun model klasifikasi penyakit jantung yang akurat dan stabil. Beberapa penelitian menunjukkan bahwa penggunaan parameter default sering kali menghasilkan performa yang kurang optimal, sehingga diperlukan proses penyesuaian parameter secara sistematis [7], [17].

Dalam penelitian ini, proses *hyperparameter tuning* dilakukan menggunakan teknik *Grid Search Cross-validation* (GridSearchCV) yang tersedia pada pustaka *scikit-learn*. *GridSearchCV* bekerja dengan cara menguji seluruh kombinasi nilai *hyperparameter* yang telah ditentukan sebelumnya, kemudian mengevaluasi setiap kombinasi menggunakan metode *k-fold cross-validation*. Dengan pendekatan ini, kinerja model tidak hanya bergantung pada satu subset data tertentu, tetapi diuji secara berulang pada beberapa bagian data latih, sehingga hasil yang diperoleh lebih stabil dan representatif. Pendekatan *Grid Search* dengan validasi silang telah banyak digunakan dalam penelitian prediksi penyakit jantung karena mampu meningkatkan performa model dan mengurangi risiko *overfitting* [17], [22].

Pada algoritma *K-Nearest Neighbors* (KNN), *hyperparameter* yang disesuaikan meliputi jumlah tetangga terdekat (number of neighbors) serta jenis metrik jarak yang digunakan dalam proses klasifikasi. Pemilihan jumlah tetangga yang tidak tepat dapat menyebabkan model menjadi terlalu sensitif terhadap *noise* (jika terlalu kecil) atau kehilangan pola lokal yang penting (jika terlalu besar). Oleh karena itu, penyesuaian nilai *k* melalui *GridSearchCV* diperlukan untuk menemukan keseimbangan terbaik antara bias dan variansi model [12], [23].

Untuk algoritma *Support Vector Machine* (SVM), *hyperparameter tuning* difokuskan pada pemilihan fungsi kernel serta parameter regularisasi. Parameter ini berperan penting dalam menentukan bentuk hyperplane pemisah antar kelas. Nilai regularisasi yang terlalu besar dapat menyebabkan model menjadi terlalu kompleks dan rentan terhadap *overfitting*, sedangkan nilai yang terlalu kecil dapat mengakibatkan *underfitting*. Beberapa penelitian menunjukkan bahwa performa SVM pada klasifikasi penyakit jantung sangat dipengaruhi oleh pemilihan parameter kernel dan regularisasi yang tepat, sehingga *tuning* parameter menjadi tahap yang krusial [10], [24].

Sementara itu, pada algoritma *Random Forest*, proses *hyperparameter tuning* mencakup pengaturan jumlah pohon (number of trees), kedalaman maksimum pohon (maximum depth), serta jumlah fitur yang dipertimbangkan pada setiap pemisahan. Parameter-parameter ini memengaruhi kompleksitas model dan kestabilan prediksi. *Random Forest* dengan jumlah pohon yang terlalu sedikit dapat menghasilkan model yang kurang stabil, sedangkan jumlah pohon yang terlalu banyak dapat meningkatkan waktu komputasi tanpa peningkatan performa yang signifikan. Oleh karena itu, penyesuaian parameter *Random Forest* melalui *GridSearchCV* diperlukan untuk memperoleh kombinasi parameter yang efisien dan optimal [6], [7], [18].

Kombinasi *hyperparameter* terbaik yang diperoleh dari proses validasi silang kemudian digunakan untuk melatih ulang masing-masing model menggunakan seluruh data latih. Model hasil pelatihan ulang ini selanjutnya digunakan pada tahap evaluasi menggunakan data uji yang telah disimpan terpisah. Pendekatan *hyperparameter tuning berbasis Grid Search* dan *Cross-validation* ini diharapkan mampu menghasilkan model klasifikasi penyakit jantung yang memiliki kinerja optimal, stabil, dan mampu melakukan generalisasi dengan baik pada data yang belum pernah dilihat sebelumnya [17], [22].

2.5 Pelatihan Model

Pada tahap ini, konfigurasi *hyperparameter* optimal yang diperoleh dari proses *tuning* sebelumnya digunakan untuk melatih seluruh model klasifikasi penyakit jantung. Proses pelatihan dilakukan menggunakan data training yang telah melalui tahapan pra-pemrosesan dan pembagian data sebelumnya. Tujuan utama tahap pelatihan model adalah membangun model yang mampu mempelajari pola hubungan antara fitur klinis dan kelas target secara optimal, serta memiliki kemampuan generalisasi yang baik terhadap data yang belum pernah dilihat sebelumnya [1], [18].

Model *K-Nearest Neighbors* (KNN) dilatih dengan menggunakan konfigurasi jumlah tetangga terbaik yang diperoleh dari proses *GridSearchCV*. KNN merupakan algoritma berbasis instance yang melakukan klasifikasi berdasarkan kedekatan jarak antar data dalam ruang fitur. Oleh karena itu, proses pelatihan KNN sangat bergantung pada distribusi data dan skala fitur, sehingga normalisasi data menjadi langkah yang penting sebelum pelatihan dilakukan [12], [23]. Beberapa penelitian menunjukkan bahwa KNN mampu memberikan performa yang kompetitif dalam klasifikasi penyakit jantung ketika nilai parameter k dipilih secara optimal dan data telah diproses dengan baik [5], [12]. Namun, KNN juga memiliki keterbatasan pada kompleksitas komputasi ketika jumlah data meningkat, sehingga pemilihan parameter yang tepat menjadi krusial.

Model *Support Vector Machine* (SVM) dilatih menggunakan konfigurasi kernel dan parameter regularisasi terbaik hasil *hyperparameter tuning*. SVM bekerja dengan mencari hyperplane optimal yang memaksimalkan margin pemisah antar kelas, sehingga efektif untuk data kesehatan dengan batas kelas yang kompleks. Proses pelatihan SVM dilakukan menggunakan data training dengan validasi silang untuk menghindari overfitting dan memastikan stabilitas model [10], [24]. Penelitian sebelumnya menunjukkan bahwa SVM mampu memberikan hasil klasifikasi yang baik pada data penyakit jantung, khususnya ketika dikombinasikan dengan pra-pemrosesan yang tepat dan pemilihan parameter kernel yang sesuai [10], [26]. Namun demikian, SVM sensitif terhadap pemilihan parameter, sehingga pelatihan model perlu dilakukan secara hati-hati agar tidak menghasilkan model yang underfitting maupun overfitting.

Selanjutnya, model *Random Forest* (RF) dilatih dengan menggunakan kombinasi parameter terbaik yang mencakup jumlah pohon (number of trees), kedalaman maksimum pohon, serta parameter pemisahan lainnya. *Random Forest* merupakan metode ensemble yang menggabungkan banyak pohon keputusan untuk meningkatkan stabilitas prediksi dan mengurangi variansi model. Proses pelatihan *Random Forest* relatif robust terhadap *noise* dan data non-linear, sehingga banyak digunakan dalam penelitian prediksi penyakit jantung dan gagal jantung [6], [7], [23]. Beberapa studi menunjukkan bahwa *Random Forest* mampu mempertahankan kinerja yang stabil meskipun terdapat variasi distribusi data, sehingga cocok digunakan untuk klasifikasi berbasis data klinis [18], [19].

Seluruh model dilatih menggunakan data training yang sama agar perbandingan kinerja antar algoritma dapat dilakukan secara adil. Teknik validasi silang yang diterapkan pada tahap *tuning* memastikan bahwa proses pelatihan tidak hanya bergantung pada satu subset data tertentu. Model hasil pelatihan selanjutnya digunakan pada tahap evaluasi menggunakan data testing untuk mengukur performa aktual dan kemampuan generalisasi. Pendekatan pelatihan model yang sistematis ini diharapkan dapat menghasilkan model klasifikasi penyakit jantung yang akurat, stabil, dan layak digunakan sebagai dasar pemilihan model terbaik pada tahap selanjutnya [1], [16], [20].

2.6 Evaluasi Model

Untuk menilai kinerja model secara objektif dan menyeluruh, terutama dengan mempertimbangkan adanya ketidakseimbangan kelas dalam *dataset*, penelitian ini menggunakan *Confusion matrix* sebagai dasar perhitungan. Matriks ini memetakan hasil prediksi menjadi empat kategori: *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Berdasarkan elemen-elemen tersebut, kinerja model diukur menggunakan metrik evaluasi utama berikut:

- a. *Accuracy* (Akurasi)

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

b. *Precision* (Presisi)

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

c. *Recall* (Sensitivitas) *Recall* atau sensitivitas

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

d. *F1-score*

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

Selain metrik evaluasi berbasis satu nilai ambang (*threshold*) seperti akurasi, *precision*, *recall*, dan *F1-score*, penelitian ini juga menggunakan *Receiver Operating Characteristic* (ROC) dan *Area Under Curve* (AUC) untuk mengevaluasi kinerja model klasifikasi. Penggunaan *AUC-ROC* penting karena memberikan evaluasi berbasis probabilitas yang lebih komprehensif dan stabil dibandingkan hanya menggunakan akurasi, khususnya pada klasifikasi medis.

e. *ROC* dan *AUC*

$$\text{Sumbu Y: True Positive Rate (TPR)} = TPR \frac{TP}{TP + FN} \quad (5)$$

$$\text{Sumbu X: False Positive Rate (FPR)} = FPR \frac{FP}{FP + TN} \quad (6)$$

$$\text{Nilai AUC} = AUC \int_0^1 TPR(FPR) d(FPR) \quad (7)$$

Keterangan:

TP (*True Positive*): Jumlah data positif (Sakit) yang diprediksi benar sebagai positif.

TN (*True Negative*): Jumlah data negatif (Normal) yang diprediksi benar sebagai negatif.

FP (*False Positive*): Jumlah data negatif yang salah diprediksi sebagai positif.

FN (*False Negative*): Jumlah data positif yang salah diprediksi sebagai negatif.

2.7 Pemilihan Model Terbaik

Pemilihan model terbaik dilakukan berdasarkan hasil evaluasi kinerja algoritma *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF) yang telah dilatih menggunakan 80% data training dan diuji pada 20% data testing yang bersifat unseen. Seluruh model dioptimalkan menggunakan *GridSearchCV* dengan 5-fold *cross-validation* sebelum dilakukan pengujian akhir. Evaluasi performa didasarkan pada metrik akurasi, *precision*, *recall*, *F1-score*, dan *confusion matrix*, serta analisis perbandingan kinerja antara data training dan data testing untuk mengidentifikasi potensi overfitting. Model dengan performa paling stabil dan seimbang pada seluruh metrik evaluasi dipilih sebagai model terbaik dan dijadikan sebagai model akhir dalam penelitian ini

3. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil dan analisis performa tiga model klasifikasi yang digunakan dalam penelitian ini, yaitu *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF). Sebelum evaluasi performa dilakukan, seluruh model dilatih menggunakan data latih sebesar 80% dari total *dataset* Heart Disease Cleveland, sedangkan 20% sisanya digunakan sebagai data uji. Pembagian data ini bertujuan untuk mengevaluasi kemampuan generalisasi model pada data yang benar-benar belum pernah dilihat selama proses pelatihan.

Proses optimasi model dilakukan sebelum evaluasi akhir. Algoritma KNN, SVM, dan *Random Forest* dioptimasi menggunakan teknik *Grid Search Cross-validation* (GridSearchCV) dengan 5-fold *cross-validation* pada data latih. Pada KNN, penyesuaian parameter difokuskan pada jumlah tetangga dan metrik jarak. Pada SVM, optimasi dilakukan pada pemilihan kernel dan parameter regularisasi, sedangkan pada *Random Forest* penyesuaian mencakup jumlah pohon dan

kedalaman maksimum pohon keputusan. Pendekatan ini bertujuan untuk memperoleh konfigurasi parameter yang paling optimal dan stabil sebelum model digunakan untuk pengujian pada data uji.

Evaluasi performa model dilakukan menggunakan data uji yang sepenuhnya terpisah (*unseen data*) dari proses pelatihan dan *tuning*. Metrik evaluasi yang digunakan meliputi *accuracy*, *precision*, *recall*, *F1-score*, *Receiver Operating Characteristic* (ROC) dan *Area Under the Curve* (AUC) serta *confusion matrix* untuk menilai kemampuan klasifikasi masing-masing model. Penggunaan berbagai metrik ini bertujuan untuk memberikan gambaran yang komprehensif mengenai kinerja model, khususnya dalam konteks medis di mana kesalahan klasifikasi terutama *false negative* dapat berdampak signifikan. Selain itu, perbandingan performa antara data latih dan data uji juga dianalisis untuk mengidentifikasi potensi overfitting pada masing-masing model.

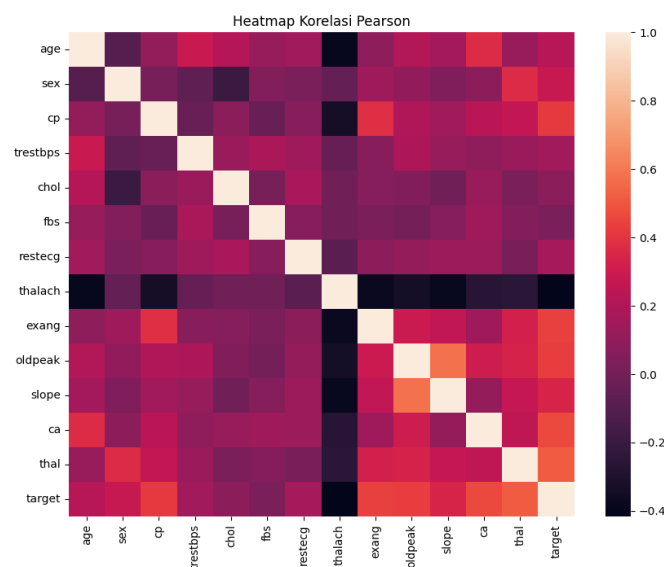
3.1 Tahapan Implementasi Algoritma dalam Klasifikasi Penyakit Jantung

Pada tahap ini dijelaskan proses implementasi algoritma dalam menyelesaikan permasalahan klasifikasi penyakit jantung menggunakan dataset *Heart Disease Cleveland*. Setelah dataset melalui tahap preprocessing, termasuk normalisasi serta pembagian data menjadi data latih (80%) dan data uji (20%), masing-masing algoritma diterapkan pada data latih untuk membangun model klasifikasi. Proses dimulai dengan fitting model menggunakan data latih sehingga algoritma dapat mempelajari pola hubungan antara fitur klinis dan variabel target. Model yang telah terbentuk kemudian digunakan untuk melakukan prediksi terhadap data uji sebagai tahap pengujian kinerja.

Implementasi setiap algoritma dilakukan sesuai dengan mekanisme pembelajarannya masing-masing. Pada algoritma *K-Nearest Neighbors* (KNN), model menyimpan seluruh data latih dan melakukan klasifikasi dengan menghitung jarak antara data uji dan data latih menggunakan metrik jarak tertentu, kemudian menentukan kelas berdasarkan mayoritas dari *k* tetangga terdekat hasil optimasi *GridSearchCV*. Pada *Support Vector Machine* (SVM), model membentuk *hyperplane* optimal yang memisahkan kelas dalam ruang fitur multidimensi berdasarkan parameter kernel dan regularisasi yang telah dioptimasi. Sementara itu, *Random Forest* membangun sejumlah pohon keputusan menggunakan teknik bootstrap sampling dan pemilihan subset fitur secara acak, dengan hasil akhir ditentukan melalui voting mayoritas. Prediksi yang dihasilkan oleh masing-masing model selanjutnya dievaluasi menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, *AUC-ROC*, serta *confusion matrix* untuk menilai kinerja klasifikasi.

3.2 Hasil Preprocessing Data dan Analisis Korelasi

Pada Gambar 2 menyajikan hasil analisis korelasi. Tahap preprocessing data dilakukan untuk memastikan kualitas dan kesiapan *dataset* sebelum digunakan dalam proses pelatihan dan evaluasi model klasifikasi.



Gambar 2. Analisis Korelasi

Dataset Heart Disease Cleveland terlebih dahulu melalui proses pemeriksaan data dan normalisasi fitur numerik untuk menyamakan skala antar variabel. Proses normalisasi ini penting karena beberapa algoritma yang digunakan dalam penelitian ini, seperti *K-Nearest Neighbors* (KNN) dan *Support Vector Machine* (SVM), sangat sensitif terhadap perbedaan skala fitur, sehingga dapat memengaruhi hasil klasifikasi apabila tidak ditangani dengan baik. Sebagai bagian dari tahap preprocessing, dilakukan analisis korelasi menggunakan metode Pearson Correlation untuk mengidentifikasi hubungan linier antar fitur serta keterkaitannya dengan variabel target penyakit jantung. Hasil analisis korelasi tersebut divisualisasikan dalam bentuk heatmap korelasi, yang memberikan gambaran visual mengenai kekuatan dan arah hubungan antar variabel, baik antar fitur maupun antara fitur dengan variabel target, sebagaimana ditunjukkan pada Gambar 2

Berdasarkan hasil visualisasi pada Gambar 2, terlihat bahwa beberapa fitur memiliki tingkat korelasi yang relatif lebih tinggi terhadap variabel target. Fitur *thalach* (denyut jantung maksimum) menunjukkan korelasi negatif yang cukup kuat terhadap target, yang mengindikasikan bahwa semakin tinggi nilai denyut jantung maksimum, kecenderungan pasien terindikasi penyakit jantung cenderung menurun. Sebaliknya, fitur *oldpeak*, *ca*, dan *thal* menunjukkan korelasi positif yang lebih kuat, yang menandakan bahwa peningkatan nilai pada fitur-fitur tersebut berasosiasi dengan meningkatnya risiko penyakit jantung. Selain hubungan antara fitur dan variabel target, analisis korelasi juga menunjukkan adanya hubungan antar fitur tertentu, seperti korelasi antara *oldpeak* dan *slope*, serta hubungan negatif antara *exang* dan *thalach*. Temuan ini mengindikasikan bahwa beberapa fitur klinis saling berkaitan dalam merepresentasikan kondisi fisiologis pasien. Namun demikian, sebagian besar fitur tidak menunjukkan korelasi yang sangat tinggi satu sama lain, sehingga risiko multikolinearitas yang dapat mengganggu proses pelatihan model relatif rendah.

3.3 Hasil Evaluasi Model Klasifikasi

Tabel 3 menyajikan Evaluasi performa model dilakukan untuk menilai kemampuan masing-masing algoritma dalam mengklasifikasikan penyakit jantung berdasarkan *dataset* Heart Disease Cleveland. Tiga algoritma yang dievaluasi dalam penelitian ini adalah *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF). Evaluasi dilakukan pada data latih (*training*) dan data uji (*testing*) menggunakan beberapa metrik klasifikasi, yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan *Area Under Curve* (AUC). Penggunaan metrik yang beragam bertujuan untuk memberikan gambaran menyeluruh mengenai kinerja model serta mendeteksi potensi *overfitting*.

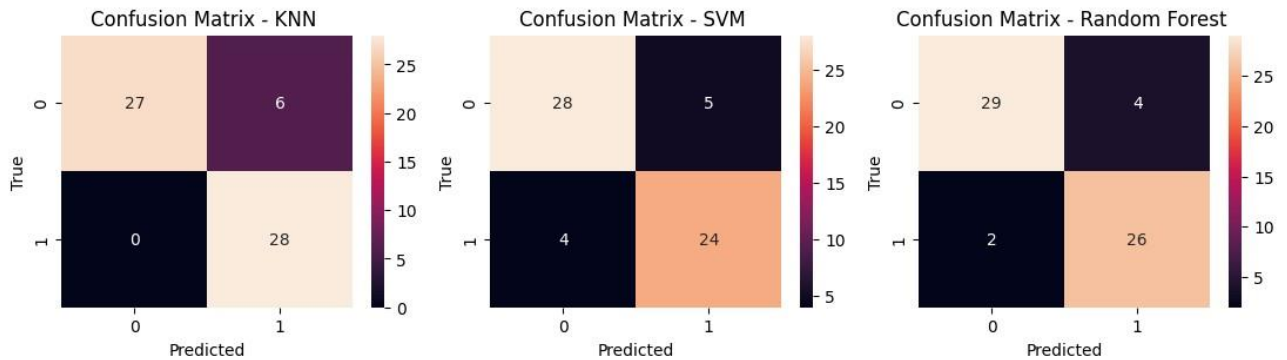
Tabel 3. Hasil Evaluasi Performansi Model Pada Data Uji

Model	Accuracy	Precision	Recall	F1-score	AUC-ROC
KNN	0,9016	0,8235	1,0000	0,9032	0,9242
SVM	0,8525	0,8276	0,8571	0,8421	0,9416
Random Forest	0,9016	0,8667	0,9794	0,8966	0,9481

Berdasarkan hasil evaluasi pada Tabel 3, ketiga algoritma menunjukkan kemampuan klasifikasi yang baik, namun dengan karakteristik performa yang berbeda. Model KNN memperoleh akurasi sebesar 0,9016 pada data uji dengan nilai *recall* mencapai 1,0000. *Recall* menunjukkan kemampuan yang sangat baik dalam klasifikasi tiap kelas. Namun demikian, nilai *precision* sebesar 0,8235 menunjukkan masih terdapat beberapa *false positive*, walaupun kondisi ini masih dapat ditoleransi dalam konteks skrining medis. Sementara itu, model SVM menunjukkan performa yang relatif stabil antara data latih dan data uji dengan akurasi sebesar 0,8525 pada data uji serta nilai *AUC* sebesar 0,9416. Nilai ini menunjukkan bahwa SVM memiliki kemampuan pemisahan kelas yang baik, meskipun nilai *recall* yang lebih rendah dibandingkan KNN menunjukkan bahwa masih terdapat sebagian kecil kasus yang tidak terdeteksi. Di sisi lain, *Random Forest* juga menunjukkan performa yang baik dan seimbang. Model ini menghasilkan akurasi sebesar 0,9016 pada data uji, serta memiliki keseimbangan yang cukup baik antara *precision* 0,8667 dan *recall* 0,9268, dengan nilai *F1-score* sebesar 0,8966. Nilai *AUC* sebesar 0,9481 menunjukkan kemampuan diskriminasi kelas yang sangat baik. Hal ini menunjukkan bahwa *Random Forest* tetap merupakan model yang kuat, meskipun tidak melampaui KNN dalam hal sensitivitas terhadap kelas positif. Secara keseluruhan, hasil evaluasi menunjukkan bahwa ketiga algoritma mampu melakukan klasifikasi penyakit jantung dengan performa yang baik. Namun, KNN menunjukkan keunggulan paling signifikan melalui nilai *recall* yang sempurna dan akurasi yang tinggi, sehingga dinilai lebih sesuai untuk konteks medis yang menekankan pentingnya meminimalkan *false negative*. Hasil evaluasi ini menjadi dasar untuk analisis lebih lanjut melalui visualisasi *confusion matrix* pada subbab berikutnya serta penentuan model terbaik pada tahap akhir penelitian.

3.3.1 Visualisasi Confusion matrix

Gambar 3 menampilkan *confusion matrix* dari tiga algoritma klasifikasi yang digunakan, yaitu *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest*, yang dievaluasi menggunakan data uji. *Confusion matrix* digunakan untuk menganalisis secara rinci jenis kesalahan klasifikasi yang terjadi, khususnya *false positive* dan *false negative*, yang memiliki implikasi penting dalam konteks diagnosis penyakit jantung. Berdasarkan *confusion matrix* KNN, model ini berhasil mendeteksi seluruh pasien yang benar-benar menderita penyakit jantung, yang ditunjukkan oleh tidak adanya *false negative* (FN = 0). Namun, masih terdapat 6 *false positive*, yaitu pasien sehat yang diprediksi sebagai sakit. Hasil ini menunjukkan bahwa KNN memiliki sensitivitas yang sangat tinggi, tetapi cenderung menghasilkan alarm palsu yang lebih banyak, sehingga menurunkan nilai *precision*. Pada *confusion matrix* SVM, terlihat adanya 4 *false negative* dan 5 *false positive*. Hal ini menunjukkan bahwa SVM masih melewatkan sebagian kecil kasus penyakit jantung, sehingga sensitivitasnya lebih rendah dibandingkan KNN dan *Random Forest*. Meskipun demikian, distribusi kesalahan yang relatif seimbang antara kelas positif dan negatif menunjukkan bahwa SVM memiliki kemampuan pemisahan kelas yang cukup stabil.

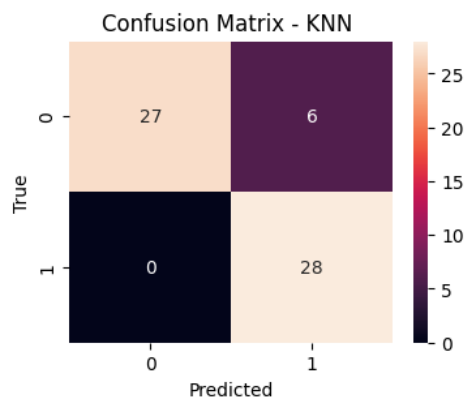


Gambar 3. Confusion matrix untuk Ketiga Model:KNN,SVM dan RF

Sementara itu, *Random Forest* menunjukkan hasil *confusion matrix* yang paling seimbang di antara ketiga model. Model ini hanya menghasilkan 2 *false negative* dan 4 *false positive*, dengan jumlah prediksi benar yang tinggi pada kedua kelas. Jumlah *false negative* yang relatif rendah menunjukkan bahwa *Random Forest* mampu mendeteksi sebagian besar pasien yang menderita penyakit jantung, sekaligus tetap menjaga kesalahan klasifikasi pada pasien sehat. Secara keseluruhan, analisis *confusion matrix* pada Gambar 3. memperkuat hasil evaluasi metrik sebelumnya. KNN unggul dalam sensitivitas, SVM menunjukkan performa moderat, sedangkan *Random Forest* memberikan keseimbangan terbaik antara deteksi pasien sakit dan pengendalian kesalahan klasifikasi. Oleh karena itu, *Random Forest* dinilai paling sesuai untuk digunakan pada tugas klasifikasi penyakit jantung dalam penelitian ini.

3.4 Pemilihan Model Terbaik

Gambar 4 menunjukkan distribusi prediksi benar (*True Positive* dan *True Negative*) serta prediksi salah (*False Positive* dan *False Negative*). Algoritma *K-Nearest Neighbors* (KNN) terbukti menjadi model yang paling efisien dan unggul dalam penelitian ini, dengan nilai akurasi sebesar 0,9016 pada data uji dan nilai *recall* mencapai 1,0000. Analisis *confusion matrix* dilakukan untuk memvalidasi kualitas prediksi model secara lebih rinci.



Gambar 4. Confusion matrix Model KNN

Hasil analisis *confusion matrix* menunjukkan bahwa sebagian besar data uji terklasifikasi dengan benar, dan yang paling penting adalah tidak ditemukannya *false negative* pada model *K-Nearest Neighbors* (KNN). Kondisi ini menunjukkan bahwa KNN memiliki kemampuan yang sangat baik dalam mendeteksi seluruh pasien yang benar-benar menderita penyakit jantung. Dalam konteks skrining medis, hal ini sangat krusial karena kesalahan *false negative* dapat menyebabkan keterlambatan diagnosis dan penanganan pasien. Nilai *recall* sebesar 1,0000 mengonfirmasi bahwa seluruh anggota kelas positif berhasil teridentifikasi dengan baik. Sementara itu, nilai *precision* sebesar 0,8235 dan *F1-score* sebesar 0,9032 menunjukkan bahwa model masih memiliki keseimbangan kinerja yang baik meskipun terdapat beberapa *false positive* yang masih terjadi, namun kondisi ini masih dapat ditoleransi dalam sistem skrining medis.

Dua faktor utama berkontribusi terhadap keunggulan kinerja algoritma *K-Nearest Neighbors* (KNN) dalam penelitian ini. Pertama, karakteristik KNN sebagai algoritma berbasis kedekatan memungkinkan model memanfaatkan pola distribusi data secara langsung, sehingga mampu mengklasifikasikan pasien berdasarkan kemiripan karakteristik klinis dengan kasus-kasus sebelumnya. Pendekatan ini membuat KNN efektif dalam mengenali pola kelas positif sehingga mampu mendeteksi seluruh pasien yang benar-benar menderita penyakit jantung. Kedua, proses *hyperparameter tuning* menggunakan *GridSearchCV* berperan penting dalam menentukan nilai parameter *k* yang paling optimal, sehingga model tidak terlalu sensitif terhadap *noise* namun tetap mampu mempertahankan kemampuan generalisasi yang baik.

Kombinasi karakteristik algoritma dan *tuning* parameter ini menjadikan KNN memiliki kinerja yang stabil tanpa indikasi overfitting, yang terlihat dari konsistensi performanya pada data latih dan data uji. Secara keseluruhan, temuan ini menunjukkan bahwa *K-Nearest Neighbors* (KNN) memiliki kemampuan klasifikasi yang andal pada *dataset* penyakit jantung Cleveland. Dengan nilai akurasi sebesar 0,9016, nilai *recall* sebesar 1,0000, serta kinerja yang konsisten pada berbagai metrik evaluasi, KNN dinilai sangat layak untuk digunakan sebagai model utama dalam sistem pendukung keputusan klinis untuk skrining dini penyakit jantung. Kemampuan KNN dalam mendeteksi seluruh pasien pada kelas positif tanpa menghasilkan *false negative* menjadi keunggulan yang sangat penting dalam konteks medis, sehingga model ini berpotensi memberikan dukungan yang signifikan terhadap proses diagnosis dan pengambilan keputusan klinis.

3.5 Pembahasan

Tabel 4 menyajikan validasi kontribusi dan posisi penelitian ini dalam ranah akademik, hasil terbaik yang diperoleh dibandingkan dengan penelitian-penelitian terdahulu yang menggunakan metode serupa pada *dataset* penyakit jantung. Penelitian ini berfokus pada analisis dan perbandingan kinerja algoritma *machine learning* klasik dalam melakukan klasifikasi penyakit jantung menggunakan *dataset* Cleveland Heart Disease. Berdasarkan hasil eksperimen, algoritma *K-Nearest Neighbors* (KNN) dan *Random Forest* (RF) menunjukkan kinerja yang lebih unggul dibandingkan *Support Vector Machine* (SVM). Kedua model tersebut menghasilkan nilai akurasi tertinggi sebesar 90,16%, dengan karakteristik performa yang berbeda.

Tabel 4. Komparansi Hasil Dengan Penelitian Terdahulu

Peneliti (Tahun)	Metode/Arsitektur	Akurasi	Keterangan
Agusyul & Firmansyah (2023)	<i>Random Forest</i>	88,00%	RF stabil pada data klinis, sensitif terhadap fitur dominan
Fredilio (2023)	KNN	86,40%	<i>Recall</i> tinggi, sensitif terhadap pemilihan parameter <i>k</i>
Hidayat (2024)	SVM	84,70%	Performa bergantung pada kernel dan regularisasi
Metode Usulan Yang Dilakukan Pada Penelitian Ini (2026)	KNN + <i>GridSearchCV</i>	90,16%	Akurasi tinggi dengan sensitivitas sempurna (<i>recall</i> = 1,0000), seluruh pasien positif terdeteksi

Analisis pada Tabel 4 menunjukkan bahwa metode yang digunakan dalam penelitian ini melampaui capaian akurasi beberapa penelitian terdahulu secara signifikan. Peningkatan performa paling jelas terlihat pada algoritma *Random Forest* dan KNN, yang memperoleh akurasi di atas 90% tanpa menggunakan teknik *oversampling* atau penambahan data sintesis. Jika dibandingkan dengan penelitian Agusyul & Firmansyah [6] yang menggunakan *Random Forest* dengan akurasi 88,00% serta penelitian Fredilio [12] yang menerapkan KNN dengan akurasi 86,40%, hasil penelitian ini menunjukkan peningkatan akurasi hingga 90,16%. Performa tersebut dicapai melalui kombinasi pra-pemrosesan yang tepat, pembagian data yang ketat, serta optimasi *hyperparameter* menggunakan *GridSearchCV* mampu meningkatkan kinerja model secara efektif.

Keunggulan performa KNN dalam penelitian ini terutama terlihat dari nilai *recall* yang mencapai 1,0000, yang menunjukkan bahwa seluruh pasien dengan penyakit jantung pada data uji berhasil terdeteksi tanpa menghasilkan *false negative*. Kondisi ini sangat penting dalam konteks medis karena kesalahan *false negative* dapat menyebabkan keterlambatan diagnosis dan penanganan pasien. Meskipun demikian, KNN masih menghasilkan sejumlah *false positive*, namun hal ini masih dapat ditoleransi pada sistem skrining medis karena pasien yang terdeteksi positif masih dapat menjalani pemeriksaan lanjutan. Di sisi lain, *Random Forest* tetap menunjukkan performa yang kuat dan seimbang dengan akurasi tinggi sebesar 90,16%, nilai *AUC* yang besar, serta keseimbangan antara *precision* dan *recall* yang baik, meskipun sensitivitasnya masih berada di bawah KNN. Dengan demikian, *Random Forest* memberikan performa yang stabil dan komprehensif, namun KNN tetap menjadi model dengan sensitivitas terbaik dalam mendeteksi pasien dengan penyakit jantung.

Secara keseluruhan, hasil penelitian ini membuktikan bahwa penerapan algoritma *machine learning* klasik yang dikombinasikan dengan strategi optimasi *hyperparameter* yang sistematis mampu menghasilkan model klasifikasi penyakit jantung yang akurat dan andal. Penelitian ini juga menegaskan bahwa tanpa menggunakan arsitektur *deep learning* yang kompleks, algoritma seperti KNN dan *Random Forest* masih mampu memberikan performa yang kompetitif sebagai alat bantu pengambilan keputusan medis. Berdasarkan nilai akurasi yang tinggi, sensitivitas yang sempurna, serta kemampuan mendeteksi seluruh pasien pada kelas positif, KNN dinilai sebagai model yang paling sesuai untuk kebutuhan skrining dini penyakit jantung dalam penelitian ini.

4. KESIMPULAN

Berdasarkan rangkaian eksperimen dan analisis perbandingan yang telah dilakukan terhadap algoritma *K-Nearest Neighbors* (KNN), *Support Vector Machine* (SVM), dan *Random Forest* (RF) dalam klasifikasi penyakit jantung menggunakan *dataset* Cleveland Heart Disease, dapat disimpulkan bahwa pemilihan algoritma memiliki pengaruh yang signifikan terhadap kinerja model. Hasil pengujian menunjukkan bahwa KNN merupakan model dengan performa terbaik dalam penelitian ini, khususnya dalam konteks skrining medis. Model KNN mampu mencapai akurasi sebesar 90,16% pada data uji dan menghasilkan nilai *recall* sebesar 1,0000, yang berarti seluruh pasien dengan penyakit jantung pada data uji berhasil terdeteksi tanpa menghasilkan *false negative*. Kondisi ini sangat penting dalam praktik klinis karena kesalahan *false negative* dapat menyebabkan keterlambatan diagnosis dan penanganan pasien. Meskipun demikian, *Random Forest* tetap menunjukkan performa yang kuat dan seimbang dengan akurasi tinggi dan nilai *AUC* yang besar, namun sensitivitasnya masih berada di bawah KNN. Keunggulan KNN dalam penelitian ini tidak terlepas dari proses optimasi *hyperparameter* menggunakan *GridSearchCV* dengan *5-fold cross-validation* yang berperan penting dalam meningkatkan stabilitas dan kemampuan generalisasi model. Temuan ini juga menegaskan bahwa evaluasi model tidak hanya perlu didasarkan pada akurasi, tetapi juga pada metrik lain seperti *precision*, *recall*, *F1-score*, dan *AUC* untuk memperoleh gambaran kinerja yang lebih komprehensif. Meskipun hasil yang diperoleh cukup memuaskan, penelitian ini masih memiliki keterbatasan karena hanya menggunakan satu *dataset* dengan jumlah fitur klinis yang terbatas. Oleh karena itu, penelitian selanjutnya disarankan untuk menggunakan *dataset* yang lebih beragam serta melibatkan populasi pasien yang lebih luas guna menguji kemampuan generalisasi model. Selain itu, pengembangan ke depan dapat diarahkan pada implementasi sistem pendukung keputusan berbasis *machine learning* yang dilengkapi dengan konsep *Explainable AI*, sehingga hasil prediksi dapat lebih mudah dipahami dan dipercaya oleh tenaga medis dalam praktik klinis.

REFERENCES

- [1] H. B. Novitasari et al., "K-nearest neighbor analysis to predict the accuracy of product delivery using administration of raw material model in the cosmetic industry (PT Cedefindo)," in *Journal of Physics: Conference Series*, Institute of Physics Publishing, Nov. 2019. doi: 10.1088/1742-6596/1367/1/012008.
- [2] A. Rohman and S. Mujiyono, "Komparasi Algoritma Machine Learning dan Ensemble Methods dalam Prediksi Penyakit Jantung dengan Dataset yang Bervariasi," 2022.
- [3] Y. Amelia, "PERBANDINGAN METODE MACHINE LEARNING UNTUK MENDETEKSI PENYAKIT JANTUNG," 2023. [Online]. Available: <http://jom.fti.budiluhur.ac.id/index.php/IDEALIS/index>
- [4] U. Aisyah Pringsewu, A. Jurnaidi Wahidin, A. Eko Setiawan, and P. Bintoro, "Aisyah Journal of Informatics and Electrical Engineering MACHINE LEARNING UNTUK KLASIFIKASI PENYAKIT JANTUNG," [Online]. Available: <http://jti.aisyahuniversity.ac.id/index.php/AJIEE>
- [5] E. Sahelvi, P. Cikita, R. M. Sapitri, R. Rahmaddeni, and L. Efrizoni, "Perbandingan Algoritma K-Nearest Neighbors dan Random Forest untuk Rekomendasi Gaya Hidup Sehat dalam Mencegah Penyakit Jantung," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 5, no. 3, pp. 830–840, Jun. 2025, doi: 10.57152/malcom.v5i3.1972.
- [6] A. Y. Agusyl and F. Firmansyah, "Prediksi Penyakit Jantung Menggunakan Algoritma Random Forest," *Jurnal Minfo Polgan*, vol. 12, no. 2, Nov. 2023, doi: 10.33395/jmp.v12i2.13214.
- [7] R. Ridwan, H. H. Handayani, S. A. P. Lestari, and Y. Cahyana, "Evaluasi Kinerja Algoritma Random Forest Dan Gradient Boosting Untuk Klasifikasi Penyakit Jantung," *Jurnal Komtika (Komputasi dan Informatika)*, vol. 9, no. 1, pp. 112–124, Jun. 2025, doi: 10.31603/komtika.v9i1.13450.
- [8] S. A. Putri, N. Selayanti, M. Kristanaya, M. P. Azzahra, M. G. Navsih, and K. M. Hindrayani, "Penerapan Machine Learning Algoritma Random Forest Untuk Prediksi Penyakit Jantung," *Seminar Nasional Sains Data*, vol. 2024, [Online]. Available: <https://www.kaggle.com/datasets/fedesoriano/heart-failure-prediction>.
- [9] G. Ayu, P. Febriyanti, and A. Baita, "Comparison of Support Vector Machine and Decision Tree Algorithm Performance with Undersampling Approach in Predicting Heart Disease Based on Lifestyle," 2025. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [10] R. Hidayat, Y. S. Sy, T. Sujana, M. Husnah, H. T. Saputra, and F. Okmayura, "Implementasi Machine Learning Untuk Prediksi Penyakit Jantung Menggunakan Algoritma Support Vector Machine," *BIOS: Jurnal Teknologi Informasi dan Rekayasa Komputer*, vol. 5, no. 2, pp. 161–168, Sep. 2024, doi: 10.37148/bios.v5i2.152.
- [11] T. Z. Jasman, E. Hasmin, Sunardi, C. Susanto, and W. Musu, "Perbandingan Logistic Regression, Random Forest, dan Perceptron pada Klasifikasi Pasien Gagal Jantung," *CSRID (Computer Science Research and Its Development Journal)*, vol. 14, no. 3, pp. 271–286, Dec. 2022, doi: 10.22303/csr.14.3.2022.271-286.
- [12] F. Fredilio, J. Rahmad, S. H. Sinurat, D. R. H. Sitompul, D. J. Ziegel, and E. Indra, "Perbandingan Algoritma K-Nearest Neighbors (K-NN) dan Random Forest terhadap Penyakit Gagal Jantung," *Jurnal Teknologi Informatika dan Komputer*, vol. 9, no. 1, pp. 471–486, Mar. 2023, doi: 10.37012/jtik.v9i1.1432.
- [13] M. Sajid Abdillah, H. Mulyo, G. Wahyu, and N. Wibowo, "Heart Failure Classification Using a Hybrid Model Based on SVM and Random Forest," *Journal of Dinda Data Science, Information Technology, and Data Analytics*, vol. 5, no. 2, pp. 208–219, 2025, [Online]. Available: <http://journal.ittelkom-pwt.ac.id/index.php/dinda>
- [14] W. C. Wahyudin, T. Sutikno, R. Umar, and A. Ridwan, "Comparative Performance Analysis of Data Mining Models for Heart Disease Detection with Feature Selection Implementation Perbandingan Kinerja Model Data Mining Dalam Deteksi Penyakit Jantung Dengan Penerapan Feature Selection," 2025.

- [15] A. Setyawan, N. Sulistianingsih, and R. Rismayati, "Perbandingan Algoritma Machine Learning Untuk Deteksi Gagal Jantung Berbasis Seleksi Fitur RFECV dan Penyeimbangan Data SMOTE."
- [16] S. Yuliasari and A. Rahmatulloh, "Performance Analysis and Accuracy of Machine Learning Algorithms for Heart Disease Prediction Evaluasi Kinerja dan Akurasi Algoritma Machine Learning untuk Prediksi Penyakit Jantung," *Jurnal Informatika dan Teknologi Informasi*, vol. 22, no. 3, pp. 98–106, 2025, doi: 10.31515/telematika.v22i3.14022.
- [17] J. Homepage et al., "MALCOM: Indonesian Journal of Machine Learning and Computer Science Implementation of Hyperparameter Tuning for Classification Models in Heart Disease Risk Prediction Penerapan Hyperparameter Tuning pada Model Klasifikasi untuk Prediksi Risiko Penyakit Jantung," vol. 5, pp. 1181–1189, 2025, doi: 10.57152/malcom.v5i4.2138.
- [18] A. S. Prabowo and F. I. Kurniadi, "Analisis Perbandingan Kinerja Algoritma Klasifikasi dalam Mendeteksi Penyakit Jantung."
- [19] J. Dwi Muthohhar and A. Prihanto, "Analisis Perbandingan Algoritma Klasifikasi untuk Penyakit Jantung," *Journal of Informatics and Computer Science*, vol. 04, 2023.
- [20] F. S. Salam Nagalay, D. Rahma Aryanti, M. Liandani, F. Sains dan Teknologi, U. Wira Buana, and F. Ilmu Kesehatan, "PENGEMBANGAN APLIKASI PREDIKSI RISIKO PENYAKIT JANTUNG BERBASIS MODEL KLASIFIKASI RANDOM FOREST MENGGUNAKAN FRAMEWORK STREAMLIT," *Jurnal Kesehatan Wira Buana*, vol. 9, pp. 2541–5387, 2025.
- [21] M. A.-Z. Faradeya and E. R. Subhiyakto, "Klasifikasi Penyakit Gagal Jantung Menggunakan Algoritma Naive Bayes," *Jurnal Algoritma*, vol. 22, no. 1, pp. 115–127, May 2025, doi: 10.33364/algoritma/v.22-1.2178.
- [22] A. R. Ramadhan et al., "Comparison of Machine Learning Models for Heart Disease Classification with Web-Based Implementation," 2024. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [23] A. Samosir, M. Hasibuan, W. E. Justino, and T. Hariyono, "Komparasi Algoritma Random Forest, Naïve Bayes dan K-Nearest Neighbor Dalam klasifikasi Data Penyakit Jantung".
- [24] L. N. Farida and S. Bahri, "Klasifikasi Gagal Jantung menggunakan Metode SVM (Support Vector Machine)," *Komputika : Jurnal Sistem Komputer*, vol. 13, no. 2, pp. 149–156, Oct. 2024, doi: 10.34010/komputika.v13i2.11330.
- [25] Ritwik_B3, "<https://www.kaggle.com/datasets/ritwikb3/heart-disease-cleveland>."
- [26] F. Handayani et al., "JEPIN (Jurnal Edukasi dan Penelitian Informatika) Komparasi Support Vector Machine, Logistic Regression Dan Artificial Neural Network dalam Prediksi Penyakit Jantung".