

Analisis Pola Pembelian Menggunakan Algoritma *FP-Growth* pada Dataset Online Retail

Andi Prastomo¹, Riezca Talita Trista², Yuli Haryanto³

Teknik Informatika, FTIK, Universitas Indraprasta PGRI, Jakarta, Indonesia^{1,2,3}

*Email Korespondensi: andi_prastomo@yahoo.co.id

Sejarah Artikel:

Diterima 02-08-2026
Disetujui 07-08-2026
Diterbitkan 09-08-2026

ABSTRACT

Heart disease is one of the leading causes of death worldwide, highlighting the need for accurate predictive methods to support early detection. This study aims to analyze and compare the performance of Decision Tree, Naïve Bayes, K-Nearest Neighbor (KNN), and Random Forest algorithms for heart disease prediction using the UCI Heart Disease Dataset. A quantitative experimental approach was employed using 920 patient records with 16 attributes. Data preprocessing included removing irrelevant attributes, handling missing values using median and mode imputation, applying One-Hot Encoding to categorical variables, standardizing numerical variables using StandardScaler, and transforming the target variable into binary classes. The dataset was divided into 80% training data and 20% testing data. The four classification models were evaluated using Accuracy, Precision, Recall, and F1-Score. The results showed that Naïve Bayes achieved the best performance, with an Accuracy of 86.96%, Precision of 87.50%, Recall of 89.22%, and F1-Score of 88.35%, outperforming the other algorithms. These findings indicate that Naïve Bayes is effective for heart disease classification and has potential as a predictive model to support early detection.

Keywords: Machine Learning, Classification, Heart Disease, Naïve Bayes, UCI Heart Disease Dataset.

ABSTRAK

Penyakit jantung merupakan salah satu penyebab utama kematian di dunia sehingga diperlukan metode prediksi yang akurat untuk mendukung deteksi dini. Penelitian ini bertujuan menganalisis dan membandingkan performa algoritma *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor* (KNN), dan *Random Forest* dalam memprediksi penyakit jantung menggunakan UCI Heart Disease Dataset. Penelitian menggunakan pendekatan kuantitatif dengan metode eksperimen terhadap 920 data pasien yang terdiri atas 16 atribut. Tahap *preprocessing* meliputi penghapusan atribut yang tidak relevan, penanganan *missing value* menggunakan median dan modus, transformasi atribut kategorikal menggunakan *One-Hot Encoding*, standarisasi atribut numerik menggunakan *StandardScaler*, serta transformasi target menjadi klasifikasi biner. Dataset dibagi menjadi 80% data pelatihan dan 20% data pengujian. Evaluasi model menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Hasil penelitian menunjukkan bahwa *Naïve Bayes* memperoleh performa terbaik dengan *Accuracy* 86,96%, *Precision* 87,50%, *Recall* 89,22%, dan *F1-Score* 88,35%, serta mengungguli ketiga algoritma lainnya. Temuan ini menunjukkan bahwa *Naïve Bayes* efektif digunakan untuk klasifikasi penyakit jantung dan berpotensi mendukung proses deteksi dini.

Katakunci: Machine Learning, Klasifikasi, Penyakit Jantung, *Naïve Bayes*, UCI Heart Disease Dataset.

PENDAHULUAN

Perkembangan Artificial Intelligence (AI) dan *Machine Learning* (ML) telah membawa perubahan signifikan dalam sektor kesehatan, terutama dalam mendukung pengambilan keputusan berbasis data yang lebih cepat, akurat, dan efisien. Berbagai algoritma *machine learning* mampu menganalisis data klinis untuk mengidentifikasi pola penyakit, memprediksi risiko kesehatan, serta mendukung diagnosis secara lebih objektif dibandingkan pendekatan konvensional (Nandy et al., 2024; Sharma et al., 2025). Dalam beberapa tahun terakhir, penerapan algoritma seperti *Random Forest*, *Support Vector Machine* (SVM), *Naïve Bayes*, dan *K-Nearest Neighbor* (KNN) telah menunjukkan kinerja yang baik dalam berbagai tugas klasifikasi penyakit berdasarkan data medis terstruktur (Al-Nafjan et al., 2025). Pemanfaatan teknologi ini tidak hanya meningkatkan akurasi prediksi, tetapi juga membantu tenaga medis dalam melakukan deteksi dini sehingga penanganan penyakit dapat dilakukan secara lebih tepat waktu (Kasbe, 2023; Yasmeen et al., 2024). Meskipun demikian, keberhasilan penerapan *machine learning* masih dipengaruhi oleh kualitas *dataset*, teknik *preprocessing*, pemilihan algoritma, serta interpretabilitas model yang digunakan. Oleh karena itu, pengembangan model prediksi berbasis *machine learning* menjadi salah satu pendekatan yang terus dikembangkan untuk meningkatkan kualitas layanan kesehatan dan mendukung sistem diagnosis berbasis data.

Penyakit jantung merupakan salah satu penyebab utama kematian di dunia sehingga deteksi dini menjadi langkah penting untuk menurunkan angka mortalitas dan meningkatkan keberhasilan penanganan pasien. Perkembangan *machine learning* telah membuka peluang dalam pengembangan sistem prediksi penyakit jantung yang mampu menganalisis data klinis secara cepat dan akurat. Berbagai penelitian menunjukkan bahwa penerapan algoritma klasifikasi, seperti *Random Forest*, *Support Vector Machine* (SVM), dan model hibrida, mampu meningkatkan akurasi prediksi penyakit jantung dibandingkan pendekatan konvensional (Almutairi et al., 2025; Ingole et al., 2024; Paracha et al., 2025). Oleh karena itu, pemanfaatan *machine learning* dalam prediksi penyakit jantung menjadi salah satu pendekatan yang potensial untuk mendukung deteksi dini dan membantu tenaga medis dalam pengambilan keputusan klinis.

Pemanfaatan algoritma *machine learning* dalam prediksi penyakit jantung telah menunjukkan peningkatan akurasi dibandingkan metode diagnostik konvensional. Berbagai algoritma klasifikasi, seperti *Random Forest*, *Logistic Regression*, *K-Nearest Neighbor* (KNN), dan *Naïve Bayes*, telah diterapkan untuk menganalisis data klinis, termasuk usia, tekanan darah, kadar kolesterol, serta riwayat kesehatan pasien (Chikhale, 2023; Fahim et al., 2024; S. Yadav et al., 2024). Keberhasilan model prediksi tidak hanya dipengaruhi oleh pemilihan algoritma, tetapi juga oleh kualitas *dataset*, teknik *preprocessing*, dan metode evaluasi yang digunakan (Ayankoya et al., 2025). Oleh karena itu, penggunaan *dataset* publik seperti *UCI Heart Disease Dataset* menjadi penting karena menyediakan data yang terstandarisasi dan banyak digunakan sebagai acuan dalam penelitian prediksi penyakit jantung, sehingga memungkinkan perbandingan performa berbagai algoritma klasifikasi secara objektif.

Berbagai penelitian telah menerapkan algoritma klasifikasi seperti *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor* (KNN), *Random Forest*, dan *Support Vector Machine* (SVM) untuk memprediksi penyakit jantung dengan tingkat akurasi yang bervariasi. Sebagian besar penelitian melaporkan bahwa *Random Forest* dan *Support Vector Machine* menghasilkan performa yang tinggi, sedangkan KNN, *Decision Tree*, dan *Naïve Bayes* juga menunjukkan hasil yang kompetitif pada kondisi tertentu (Biddinika et al., 2024a; Prasetyo et al., 2023; Simanjuntak et al., 2025). Perbedaan hasil tersebut menunjukkan bahwa performa algoritma sangat dipengaruhi oleh karakteristik *dataset*, teknik *preprocessing*, metode validasi, serta konfigurasi parameter yang digunakan (Haider et al., 2025; Suwondo et al., 2025).

Meskipun berbagai penelitian telah membandingkan algoritma klasifikasi untuk prediksi penyakit jantung, hasil yang diperoleh masih menunjukkan performa yang beragam. Perbedaan tersebut dipengaruhi oleh variasi karakteristik *dataset*, tahapan *preprocessing*, metode validasi, teknik seleksi fitur, serta konfigurasi parameter yang digunakan pada masing-masing penelitian. Selain itu, sebagian besar penelitian menggunakan kombinasi algoritma, teknik optimasi, atau *dataset* yang berbeda sehingga hasil yang diperoleh sulit dibandingkan secara langsung. Kondisi tersebut menyebabkan belum terdapat kesepakatan mengenai algoritma klasifikasi yang paling sesuai untuk prediksi penyakit jantung. Oleh karena itu, diperlukan penelitian yang mengevaluasi beberapa algoritma menggunakan *dataset*, tahapan *preprocessing*, dan metode evaluasi yang seragam agar perbandingan performa dapat dilakukan secara lebih objektif.

Berdasarkan kondisi tersebut, penelitian ini menawarkan analisis komparatif terhadap empat algoritma klasifikasi, yaitu *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor (KNN)*, dan *Random Forest*, menggunakan *UCI Heart Disease Dataset* sebagai sumber data yang sama. Seluruh algoritma diuji melalui tahapan *preprocessing* yang identik, meliputi penghapusan atribut yang tidak relevan, penanganan *missing value*, transformasi data kategorikal, standardisasi atribut numerik, serta transformasi variabel target menjadi klasifikasi biner. Selain itu, seluruh model dievaluasi menggunakan metrik yang sama, yaitu *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Pendekatan ini diharapkan mampu menghasilkan perbandingan performa yang lebih objektif sehingga dapat memberikan informasi yang lebih akurat mengenai algoritma yang paling sesuai untuk prediksi penyakit jantung pada *UCI Heart Disease Dataset*.

Berdasarkan *research gap* dan kebaruan penelitian yang telah diuraikan, penelitian ini bertujuan mengevaluasi dan membandingkan performa algoritma *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor (KNN)*, dan *Random Forest* dalam memprediksi penyakit jantung menggunakan *UCI Heart Disease Dataset*. Evaluasi dilakukan berdasarkan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score* untuk mengidentifikasi algoritma yang memiliki performa klasifikasi terbaik. Hasil penelitian ini diharapkan dapat menjadi referensi dalam pemilihan algoritma klasifikasi yang sesuai untuk pengembangan sistem pendukung keputusan berbasis *machine learning* dalam deteksi dini penyakit jantung serta menjadi acuan bagi penelitian selanjutnya yang menggunakan *dataset* dan metode klasifikasi sejenis. Berdasarkan *research gap* yang telah diuraikan, penelitian ini bertujuan mengevaluasi dan membandingkan performa algoritma *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor (KNN)*, dan *Random Forest* pada *UCI Heart Disease Dataset* menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score* untuk memperoleh algoritma klasifikasi yang paling sesuai dalam prediksi penyakit jantung.

METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen (*experimental research*) untuk membandingkan performa beberapa algoritma klasifikasi dalam memprediksi penyakit jantung menggunakan *UCI Heart Disease Dataset*. Dataset diperoleh dari *UCI Machine Learning Repository* dan terdiri atas 920 data pasien dengan 16 atribut, yang mencakup variabel numerik dan kategorikal. Variabel target pada penelitian ini adalah *num*, yang menunjukkan status penyakit jantung pasien. Sebelum proses pemodelan, variabel target ditransformasikan menjadi klasifikasi biner, yaitu 0 untuk pasien yang tidak memiliki penyakit jantung dan 1 untuk pasien yang memiliki penyakit jantung.

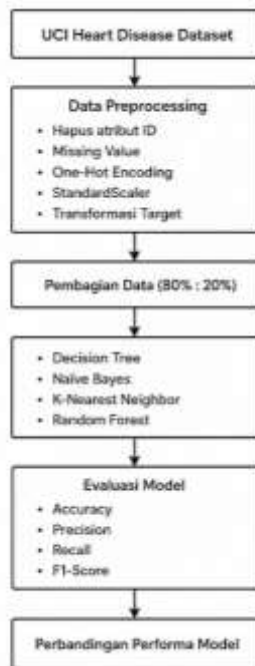
Tahap awal penelitian dilakukan melalui *preprocessing* data untuk meningkatkan kualitas dataset. Proses ini meliputi penghapusan atribut *id* karena tidak memiliki pengaruh terhadap proses klasifikasi, penanganan

missing value menggunakan median untuk atribut numerik dan modus untuk atribut kategorikal, transformasi atribut kategorikal menggunakan One-Hot Encoding, serta standarisasi atribut numerik menggunakan StandardScaler. Seluruh tahapan preprocessing dilakukan menggunakan pustaka Scikit-learn sehingga proses pengolahan data dapat direproduksi secara konsisten.

Setelah preprocessing selesai, dataset dibagi menjadi 80% data latih dan 20% data uji menggunakan metode Stratified Train-Test Split untuk mempertahankan proporsi kelas pada kedua kelompok data. Selanjutnya dilakukan proses pelatihan model menggunakan empat algoritma klasifikasi, yaitu Decision Tree, Naïve Bayes, K-Nearest Neighbor (KNN), dan Random Forest. Keempat algoritma dipilih karena merupakan metode klasifikasi yang umum digunakan pada penelitian prediksi penyakit serta memiliki karakteristik pembelajaran yang berbeda sehingga sesuai untuk analisis komparatif.

Evaluasi performa model dilakukan menggunakan Accuracy, Precision, Recall, dan F1-Score yang dihitung berdasarkan Confusion Matrix. Accuracy digunakan untuk mengukur tingkat ketepatan prediksi secara keseluruhan, Precision menunjukkan ketepatan model dalam mengidentifikasi pasien yang diprediksi memiliki penyakit jantung, Recall mengukur kemampuan model mendeteksi seluruh pasien yang benar-benar memiliki penyakit jantung, sedangkan F1-Score digunakan untuk mengevaluasi keseimbangan antara Precision dan Recall. Seluruh proses analisis data dilakukan menggunakan bahasa pemrograman Python pada platform Google Colaboratory dengan memanfaatkan pustaka Pandas, NumPy, Scikit-learn, Matplotlib, dan Seaborn.

Tahapan penelitian secara keseluruhan ditunjukkan pada Gambar 1.



Gambar 1. Alur Penelitian

HASIL DAN PEMBAHASAN

Deskripsi Dataset

Penelitian ini menggunakan *UCI Heart Disease Dataset* yang diperoleh dari *UCI Machine Learning Repository*. Sebelum dilakukan proses *preprocessing* dan pemodelan, terlebih dahulu diidentifikasi karakteristik dataset yang digunakan. Ringkasan karakteristik dataset disajikan pada Tabel 1.

Tabel 1. Karakteristik UCI Heart Disease Dataset

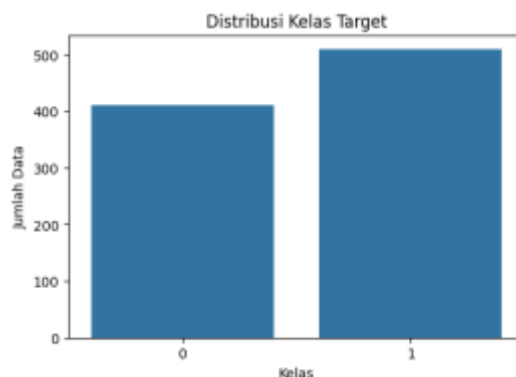
Karakteristik	Keterangan
Nama Dataset	UCI Heart Disease Dataset
Sumber Dataset	UCI Machine Learning Repository
Jumlah Data	920 data
Jumlah Variabel	15 variabel
Variabel Numerik	7 variabel
Variabel Kategorikal	8 variabel
Variabel Target	num
Jenis Klasifikasi	Biner (0 = Tidak Sakit, 1 = Sakit)

Sumber: UCI Machine Learning Repository, Python (2026).

Berdasarkan Tabel 1, dataset terdiri atas 920 data pasien dengan 15 variabel, yang meliputi 7 variabel numerik dan 8 variabel kategorikal. Variabel target yang digunakan adalah num, yang pada penelitian ini ditransformasikan menjadi klasifikasi biner, yaitu 0 untuk pasien yang tidak memiliki penyakit jantung dan 1 untuk pasien yang memiliki penyakit jantung. Transformasi tersebut dilakukan agar sesuai dengan tujuan penelitian, yaitu membandingkan performa beberapa algoritma klasifikasi biner dalam memprediksi penyakit jantung.

Hasil Preprocessing Data

Selain analisis statistik deskriptif, dilakukan analisis distribusi kelas target untuk mengetahui proporsi data pasien yang termasuk ke dalam kategori memiliki penyakit jantung dan tidak memiliki penyakit jantung. Distribusi kelas target disajikan pada Gambar 1.



Gambar 1. Distribusi Kelas Target UCI Heart Disease Dataset

Berdasarkan Gambar 1, dataset terdiri atas 508 pasien yang memiliki penyakit jantung dan 412 pasien yang tidak memiliki penyakit jantung. Distribusi kelas yang relatif seimbang menunjukkan bahwa dataset layak digunakan untuk proses pelatihan dan pengujian model klasifikasi tanpa memerlukan penanganan khusus terhadap ketidakseimbangan data.

Hasil Evaluasi Algoritma dan Pembahasan

Sebelum proses pemodelan dilakukan, dataset terlebih dahulu dianalisis untuk mengidentifikasi kualitas data, khususnya keberadaan nilai yang hilang (*missing value*). Identifikasi *missing value* penting dilakukan karena keberadaan data yang tidak lengkap dapat memengaruhi proses pelatihan model dan menurunkan kinerja algoritma klasifikasi. Hasil identifikasi jumlah *missing value* pada setiap variabel disajikan pada Tabel 2.

Tabel 2. Jumlah Missing Value Sebelum Preprocessing

Variabel	Missing Value
trestbps	59
chol	30
fbs	90
restecg	2
thalach	55
exang	55
oldpeak	62
slope	309
ca	611
thal	486

Sumber: Hasil Olah Data Python (2026).

Berdasarkan Tabel 2, beberapa variabel masih mengandung *missing value* dengan jumlah yang bervariasi. Variabel *ca* memiliki jumlah *missing value* tertinggi, yaitu sebanyak 611 data, diikuti oleh variabel *thal* sebanyak 486 data dan *slope* sebanyak 309 data. Selain itu, variabel *trestbps*, *chol*, *fbs*, *thalach*, *exang*, *oldpeak*, dan *restecg* juga memiliki sejumlah data yang tidak lengkap. Untuk mempertahankan seluruh 920 data tanpa mengurangi jumlah observasi, penelitian ini menerapkan teknik imputasi pada tahap preprocessing. Nilai yang hilang pada atribut numerik diisi menggunakan nilai median, sedangkan atribut kategorikal diisi menggunakan nilai modus sebelum dilakukan proses transformasi data dan pemodelan.

Setelah melalui tahap *preprocessing*, *dataset* digunakan untuk membangun model klasifikasi menggunakan empat algoritma, yaitu Decision Tree, Naïve Bayes, K-Nearest Neighbor (KNN), dan Random Forest. Evaluasi dilakukan untuk mengukur kemampuan masing-masing algoritma dalam memprediksi penyakit jantung menggunakan metrik Accuracy, Precision, Recall, dan F1-Score. Hasil evaluasi setiap algoritma disajikan pada Tabel 3.

Tabel 3. Hasil Evaluasi Algoritma Klasifikasi

Algoritma	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Decision Tree	76,63	77,06	82,35	79,62
Naïve Bayes	86,96	87,50	89,22	88,35
K-Nearest Neighbor	83,70	82,73	89,22	85,85
Random Forest	83,70	84,62	86,27	85,44

Sumber: Hasil Olah Data Python (2026).

Berdasarkan Tabel 3, seluruh algoritma mampu melakukan klasifikasi terhadap UCI Heart Disease Dataset dengan tingkat performa yang berbeda. Algoritma Naïve Bayes memperoleh hasil terbaik dengan nilai Accuracy sebesar 86,96%, Precision 87,50%, Recall 89,22%, dan F1-Score 88,35%. Hasil tersebut menunjukkan bahwa pendekatan probabilistik yang digunakan Naïve Bayes mampu mengenali pola hubungan antaratribut pada dataset secara lebih efektif dibandingkan algoritma lainnya. Selain memiliki akurasi tertinggi, algoritma ini juga menghasilkan keseimbangan yang baik antara Precision dan Recall, sehingga mampu meminimalkan kesalahan klasifikasi pada pasien yang memiliki maupun tidak memiliki penyakit jantung.

Algoritma K-Nearest Neighbor (KNN) dan Random Forest menunjukkan performa yang relatif seimbang dengan nilai Accuracy masing-masing sebesar 83,70%. Meskipun demikian, KNN memiliki nilai Recall yang lebih tinggi (89,22%) dibandingkan Random Forest (86,27%), yang menunjukkan kemampuan lebih baik dalam mendeteksi pasien yang benar-benar memiliki penyakit jantung. Sebaliknya, Random Forest menghasilkan nilai Precision yang lebih tinggi (84,62%), sehingga lebih baik dalam mengurangi kesalahan prediksi positif. Perbedaan tersebut menunjukkan bahwa setiap algoritma memiliki karakteristik yang berbeda sesuai dengan mekanisme pembelajaran yang digunakan.

Sementara itu, Decision Tree memperoleh Accuracy sebesar 76,63%, merupakan nilai terendah dibandingkan algoritma lainnya. Hasil tersebut menunjukkan bahwa penggunaan satu pohon keputusan memiliki keterbatasan dalam mempelajari pola data yang kompleks sehingga kemampuan generalisasinya lebih rendah. Walaupun demikian, Decision Tree tetap memiliki keunggulan dari sisi interpretasi karena aturan keputusan yang dihasilkan lebih mudah dipahami dibandingkan algoritma lainnya.

Secara keseluruhan, hasil penelitian menunjukkan bahwa Naïve Bayes merupakan algoritma yang paling optimal untuk memprediksi penyakit jantung pada UCI Heart Disease Dataset. Temuan ini mengindikasikan bahwa karakteristik data yang digunakan lebih sesuai dimodelkan menggunakan pendekatan probabilistik dibandingkan metode berbasis pohon keputusan maupun metode berbasis jarak. Oleh karena itu, Naïve Bayes dapat dipertimbangkan sebagai alternatif model klasifikasi yang efektif dalam pengembangan sistem pendukung keputusan untuk prediksi penyakit jantung.

Pembahasan

Hasil penelitian menunjukkan bahwa algoritma *Naïve Bayes* memberikan performa terbaik dalam memprediksi penyakit jantung menggunakan *UCI Heart Disease Dataset* dengan nilai *Accuracy* sebesar 86,96%, *Precision* 87,50%, *Recall* 89,22%, dan *F1-Score* 88,35%. Hasil tersebut menunjukkan bahwa pendekatan probabilistik yang diterapkan oleh *Naïve Bayes* mampu mengidentifikasi hubungan antaratribut klinis secara efektif sehingga menghasilkan klasifikasi yang lebih akurat dibandingkan algoritma *Decision Tree*, *K-Nearest Neighbor (KNN)*, dan *Random Forest*. Meskipun algoritma ini menggunakan asumsi

independensi antaratribut, asumsi tersebut tidak mengurangi kemampuan model dalam mengklasifikasikan data pasien pada *UCI Heart Disease Dataset*. Hal ini mengindikasikan bahwa sebagian besar atribut klinis, seperti usia, tekanan darah, kadar kolesterol, dan detak jantung maksimum, memiliki distribusi yang cukup representatif sehingga dapat dimodelkan dengan baik menggunakan pendekatan probabilistik.

Keunggulan *Naïve Bayes* pada penelitian ini juga diduga dipengaruhi oleh proses *preprocessing* yang diterapkan sebelum pemodelan. Tahapan penghapusan atribut yang tidak relevan, penanganan *missing value*, *one-hot encoding*, standardisasi atribut numerik menggunakan *StandardScaler*, serta transformasi variabel target menjadi klasifikasi biner menghasilkan dataset yang lebih konsisten untuk proses pelatihan model. Kondisi tersebut memungkinkan *Naïve Bayes* membangun distribusi probabilitas setiap atribut secara lebih stabil sehingga menghasilkan keseimbangan antara nilai *Precision* dan *Recall*. Nilai *Recall* sebesar 89,22% menunjukkan bahwa model memiliki kemampuan yang baik dalam mengidentifikasi pasien yang benar-benar menderita penyakit jantung, sedangkan nilai *Precision* sebesar 87,50% menunjukkan bahwa sebagian besar prediksi positif yang dihasilkan model merupakan prediksi yang benar. Keseimbangan kedua metrik tersebut tercermin pada nilai *F1-Score* sebesar 88,35%, yang menunjukkan bahwa model memiliki performa klasifikasi yang baik secara keseluruhan.

Temuan penelitian ini berbeda dengan beberapa penelitian terdahulu yang melaporkan bahwa *Random Forest* merupakan algoritma terbaik dalam prediksi penyakit jantung. C. Yadav dan Chandrakar (2025) melaporkan bahwa *Random Forest* mampu mencapai tingkat akurasi sebesar 90,16% karena memanfaatkan mekanisme *ensemble learning* yang menggabungkan banyak pohon keputusan sehingga lebih efektif dalam memodelkan hubungan nonlinier antarvariabel. Temuan serupa juga disampaikan oleh Samyal dan Singh (2025) serta Sawan et al. (2023) yang menyatakan bahwa metode *ensemble* menghasilkan performa klasifikasi yang lebih stabil pada data medis yang memiliki kompleksitas tinggi. Perbedaan hasil tersebut menunjukkan bahwa keunggulan suatu algoritma tidak bersifat universal, melainkan sangat bergantung pada karakteristik dataset, jumlah data, kualitas atribut, serta strategi *preprocessing* dan *hyperparameter tuning* yang diterapkan.

Selain itu, Barry et al. (2023) menunjukkan bahwa algoritma *K-Nearest Neighbor (KNN)* mampu mencapai akurasi sebesar 91% setelah menerapkan teknik *Synthetic Minority Over-sampling Technique (SMOTE)* untuk mengatasi ketidakseimbangan kelas. Kondisi tersebut berbeda dengan penelitian ini yang tidak menerapkan teknik *oversampling* maupun *undersampling* karena distribusi kelas pada *UCI Heart Disease Dataset* masih relatif seimbang. Perbedaan strategi penanganan data tersebut menjadi salah satu faktor yang memengaruhi variasi performa antaralgoritma. Dengan demikian, hasil penelitian ini menguatkan bahwa pemilihan teknik *preprocessing* memiliki pengaruh yang signifikan terhadap performa model klasifikasi.

Di sisi lain, hasil penelitian ini sejalan dengan penelitian Durgadevi (2016) yang menyatakan bahwa *Naïve Bayes* tetap mampu menghasilkan performa klasifikasi yang kompetitif meskipun menggunakan asumsi independensi antaratribut. Selain memiliki tingkat akurasi yang baik, *Naïve Bayes* juga memiliki kompleksitas komputasi yang relatif rendah dibandingkan algoritma berbasis *ensemble*. Keunggulan tersebut menjadikan *Naïve Bayes* lebih efisien dalam proses pelatihan maupun prediksi, terutama ketika diterapkan pada dataset berukuran kecil hingga sedang. Oleh karena itu, algoritma ini memiliki potensi untuk diimplementasikan pada sistem pendukung keputusan yang membutuhkan proses prediksi secara cepat dengan sumber daya komputasi yang terbatas.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa tidak terdapat satu algoritma yang selalu memberikan performa terbaik pada seluruh kasus prediksi penyakit jantung. Kinerja algoritma dipengaruhi

oleh berbagai faktor, seperti karakteristik dataset, kualitas data, tahapan *preprocessing*, metode validasi, serta konfigurasi parameter model (Biddinika et al., 2024b; Goel & Pandey, 2023; Kusuma & Jothi, 2021). Pada penelitian ini, *Naïve Bayes* memberikan performa terbaik dibandingkan *Decision Tree*, *K-Nearest Neighbor (KNN)*, dan *Random Forest*, sehingga dapat direkomendasikan sebagai algoritma klasifikasi untuk prediksi penyakit jantung menggunakan *UCI Heart Disease Dataset*. Temuan ini memberikan kontribusi bagi pengembangan sistem pendukung keputusan berbasis *machine learning*, khususnya dalam pemilihan algoritma klasifikasi yang efektif untuk membantu deteksi dini penyakit jantung. Meskipun demikian, penelitian ini masih terbatas pada penggunaan satu dataset publik dan empat algoritma klasifikasi. Penelitian selanjutnya dapat memperluas kajian dengan menggunakan dataset yang lebih beragam, menerapkan teknik *feature selection* dan *hyperparameter tuning*, serta membandingkan algoritma lain seperti *Extreme Gradient Boosting (XGBoost)*, *Light Gradient Boosting Machine (LightGBM)*, atau *Support Vector Machine (SVM)* untuk memperoleh model prediksi yang lebih optimal.

KESIMPULAN

Penelitian ini bertujuan membandingkan performa empat algoritma klasifikasi, yaitu *Decision Tree*, *Naïve Bayes*, *K-Nearest Neighbor (KNN)*, dan *Random Forest*, dalam memprediksi penyakit jantung menggunakan *UCI Heart Disease Dataset*. Hasil penelitian menunjukkan bahwa seluruh algoritma mampu melakukan klasifikasi terhadap data pasien dengan tingkat performa yang berbeda. Berdasarkan hasil evaluasi menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*, algoritma *Naïve Bayes* memperoleh performa terbaik dengan nilai *Accuracy* sebesar 86,96%, *Precision* 87,50%, *Recall* 89,22%, dan *F1-Score* 88,35%. Hasil tersebut menunjukkan bahwa pendekatan probabilistik yang diterapkan oleh *Naïve Bayes* lebih sesuai dengan karakteristik dataset yang digunakan dibandingkan algoritma *Decision Tree*, *K-Nearest Neighbor*, dan *Random Forest*. Dengan demikian, *Naïve Bayes* dapat direkomendasikan sebagai model klasifikasi untuk prediksi penyakit jantung pada *UCI Heart Disease Dataset*. Meskipun demikian, penelitian ini masih terbatas pada penggunaan satu dataset publik dan empat algoritma klasifikasi. Oleh karena itu, penelitian selanjutnya disarankan menggunakan dataset yang lebih beragam, menerapkan teknik *hyperparameter tuning*, serta membandingkan algoritma lain seperti *Support Vector Machine (SVM)*, *XGBoost*, atau *LightGBM* untuk memperoleh model prediksi yang lebih akurat dan memiliki kemampuan generalisasi yang lebih baik.

DAFTAR PUSTAKA

- Almutairi, M., Dardouri, S., Almutairi, M., & Dardouri, S. (2025). Intelligent Hybrid Modeling for Heart Disease Prediction. *Information*, 16(10), 869. <https://doi.org/10.3390/info16100869>
- Al-Nafjan, A., Aljuhani, A., Alshebel, A., AlHarbi, A., & Alshehri, A. (2025). Artificial Intelligence in Predictive Healthcare: A Systematic Review. *Stomatology*, 14(19), 6752. <https://doi.org/10.3390/jcm14196752>
- Ayankoya, F., Kuyoro, S. O., Nzenwata, U., & Edwin, E. (2025). A Systematic Review of Machine Learning and Deep Learning Approaches for Heart Disease Prediction. *Asian Journal of Research in Computer Science*. <https://doi.org/10.9734/ajrcos/2025/v18i12799>

- Barry, K. A., Manzali, Y., Labaybi, O., Rachid, F., & Elfar, M. (2023). Heart Disease Prediction and Classification Using Machine Learning Algorithms. 71–76. <https://doi.org/10.1109/cist56084.2023.10410012>
- Biddinika, M. K., Masitha, A., Herman, H., & Fatimah, V. A. N. (2024a). Machine Learning Techniques for Heart Disease Prediction Using a Multi-Algorithm Approach. *Jurnal Informatika*. <https://doi.org/10.30595/juita.v12i2.24153>
- Biddinika, M. K., Masitha, A., Herman, H., & Fatimah, V. A. N. (2024b). Machine Learning Techniques for Heart Disease Prediction Using a Multi-Algorithm Approach. *Jurnal Informatika*. <https://doi.org/10.30595/juita.v12i2.24153>
- Chikhale, S. R. (2023). Review on Heart Disease Prediction Using Machine Learning Algorithms. *International Journal for Research in Applied Science and Engineering Technology*. <https://doi.org/10.22214/ijraset.2023.57048>
- Durgadevi, A. (2016). Comparative Study of Data Mining Classification Algorithms in Heart Disease Prediction. <http://www.paperpublications.org/download.php?file=Comparative%20Study%20of%20Data%20Mining-501.pdf&act=book>
- Fahim, M., Hassan, M. K., Karpagavalli, S. M., Jansi, S., Jha, R. K., & Mahesh, T. R. (2024). Advancing Cardiovascular Health Prediction: Machine Learning Algorithm Analysis. <https://doi.org/10.1109/csnt60213.2024.10546094>
- Goel, N., & Pandey, A. (2023). Comparative Analysis of Single Classifier Models against Aggregated Fusion Models for Heart Disease Prediction. *International Conference on Database Theory*, 576–580. <https://doi.org/10.1109/ICDT57929.2023.10150611>
- Haider, M., Hussain, M., & Insany, G. P. (2025). Heart Attack Prediction Using Machine Learning Models: A Comparative Study of Naive Bayes, Decision Tree, Random Forest, and K-Nearest Neighbors. 121. <https://doi.org/10.3390/engproc2025107121>
- Ingle, B. S., Ramineni, V., Jayaram, V., Banarse, A. R., Krishnappa, M. S., Pulipeta, N. K., Parlapalli, V., & Pandey, G. (2024). Prediction and Early Detection of Heart Disease: A Hybrid Neural Network and SVM Approach. 282–286. <https://doi.org/10.1109/mcsoc64144.2024.00054>
- Kasbe, T. (2023). Machine Learning in Healthcare (pp. 1–13). IGI Global. <https://doi.org/10.4018/978-1-6684-8974-1.ch001>
- Kusuma, S., & Jothi, K. R. (2021). Cardiovascular Disease Prediction and Comparative Analysis of Varied Classifier Techniques. <https://doi.org/10.1109/GCAT52182.2021.9587734>
- Nandy, A., Choudhary, J. S., Fredrick, J., Zacharia, T. S., Joseph, T., & Kaur, V. (2024). Machine Learning for Healthcare. 59–77. <https://doi.org/10.1201/9781032694870-5>
- Paracha, W. T., Nasir, J. A., Farhan, M., Paracha, M. F. K., & Iqbal, M. J. (2025). Optimizing Heart Disease Prediction with Machine Learning and Feature Selection Techniques. <https://doi.org/10.5281/zenodo.16487579>
- Prasetyo, S. Y., Saputri, H. A., Nabiilah, G. Z., & Wulandari, A. (2023). Model-Based Learning Techniques for Accurate Heart Disease Risk Prediction. 266–271. <https://doi.org/10.1109/icmeralda60125.2023.10458149>
- Samyal, V., & Singh, S. (2025). Performance Comparison of Machine Learning Models for Heart Disease Prediction. *Ibn Al-Haitham*. <https://doi.org/10.71097/ijrsat.v16.i4.9857>

- Sawan, V., Jugnu, K., Roy, D. P., & Pandey, R. (2023). Comparative assessment of machine learning algorithms for heart disease prediction. *International Journal of Engineering Applied Science and Technology*. <https://doi.org/10.33564/ijeast.2023.v08i03.016>
- Sharma, P., Kumra, R., & Bhaggi, E. (2025). Transforming Healthcare. *Advances in Computational Intelligence and Robotics Book Series*, 413–430. <https://doi.org/10.4018/979-8-3373-2597-2.ch014>
- Simanjuntak, M. S., Robet, R., & Hoki, L. (2025). A Comparative Study of Machine Learning and Deep Learning Models for Heart Disease Classification. *Journal of Applied Informatics and Computing*. <https://doi.org/10.30871/jaic.v9i6.11546>
- Suwondo, A., Kusriani, K., Utami, E., & Yuana, K. A. (2025). Survey of the Performance of Classification Algorithms on Heart Disease Datasets: A Systematic Literature Review. 539–546. <https://doi.org/10.1109/ies67184.2025.11161964>
- Yadav, C., & Chandrakar, S. (2025). Machine Learning-Based Heart Disease Prediction: A Comparative Evaluation of Classification Methods. *International Journal of Advanced Research in Science, Communication and Technology*, 240–247. <https://doi.org/10.48175/ijarsct-29634>
- Yadav, S., Yadav, A., Srivastava, S. P., & Rai, P. (2024). Heart Disease Prediction Using Machine Learning. *Indian Scientific Journal Of Research In Engineering And Management*, 08(008), 1–4. <https://doi.org/10.55041/ijrsrem37304>
- Yasmeen, Z., Machi, S., Maguluri, K. K., Mandala, G., & Danda, R. (2024). Transforming Patient Outcomes: Cutting-Edge Applications of AI and ML in Predictive Healthcare. *South Eastern European Journal of Public Health*, 1704–1712. <https://doi.org/10.70135/seejph.vi.220>

