

KLASIFIKASI PENYAKIT JANTUNG BERDASARKAN FAKTOR KLINIS MENGGUNAKAN ALGORITMA XG BOOST, RANDOM FOREST, SUPPORT VECTOR MACHINE (SVM) DAN K-NEAREST NEIGHBOR (KNN)

Siti Azizah Aini , Diah Fitriani , Mohammad Alif Awwalin,
Ahmad Abil Rizky Ramadhan, Ifnu Wisma Dwi Prastya
Teknik Informatika, Universitas Nahdlatul Ulama Sunan Giri
Jl Ahmad Yani No.10 Jambean, Sukoharjo, Bojonegoro, Jawa Timur
sitiazizahaini@gmail.com

ABSTRAK

Saat ini penyakit jantung menjadi tantangan utama dalam bidang kesehatan masyarakat, diperlukan sarana deteksi dini yang efektif karena penyakit jantung adalah salah satu penyebab kematian paling umum di dunia. Metode mesin pembelajaran dikembangkan untuk membantu proses diagnosis karena data klinis yang kompleks dan metode konvensional yang terbatas. Menggunakan algoritma *Random Forest*, *Support Vector Machine (SVM)*, *XGBoost*, dan *K-Nearest Neighbors (KNN)*, penelitian ini bertujuan untuk mengklasifikasikan penyakit jantung berdasarkan faktor klinis. Dengan jumlah 1.888 dataset pasien dari *Kaggle*. Proses *preprocessing* data, pembagian data latih, dan pembagian data uji adalah semua hasil dari penelitian ini, dengan rasio 80:20, pembuatan model dan penilaian menggunakan akurasi, presisi, recall, dan skor F1. Hasil pengujian menunjukkan bahwa *Random Forest* adalah yang terbaik dengan akurasi sebesar 97,88%, diikuti oleh *XGBoost* sebesar 94,71%, *SVM* sebesar 92,59%, dan *KNN* sebesar 90,21%. Berdasarkan hasil ini, *Random Forest* dianggap sebagai sistem yang paling ideal dan dapat dipercaya untuk mengklasifikasikan penyakit jantung berdasarkan faktor klinis, dan dapat digunakan sebagai sistem pendukung keputusan dalam bidang kesehatan.

Kata kunci : *penyakit jantung, machine learning, Random Forest, SVM, XGBoost, KNN*

1. PENDAHULUAN

Hingga saat ini penyakit jantung tetap menjadi tantangan utama dalam bidang kesehatan masyarakat, sekaligus menjadi salah satu faktor penyebab kematian paling tinggi di seluruh dunia. Berdasarkan laporan terbaru dari *World Health Organization (WHO)*, penyakit kardiovaskular menjadi penanggung jawab atas lebih dari 32% dari semua kematian di seluruh dunia setiap tahun. Hal ini menunjukkan bahwa penyakit ini merupakan ancaman besar bagi kualitas hidup masyarakat di berbagai negara [1] [2]. Permasalahan ini bukan hanya memengaruhi terhadap kesehatan individu, tetapi juga menimbulkan beban ekonomi bagi sistem pelayanan kesehatan, termasuk biaya pengobatan yang meningkat, jangka panjang perawatan, dan penurunan produktivitas masyarakat.

Di Indonesia, peningkatan kasus tidak hanya terjadi pada kelompok usia lanjut, tetapi juga mulai banyak ditemukan pada kelompok usia produktif. Fenomena ini mengindikasikan adanya perubahan pola hidup masyarakat modern, contohnya mengkonsumsi makanan yang tidak sehat, dan berkurangnya aktivitas, serta meningkatnya paparan faktor lain seperti stres, obesitas, dan kebiasaan merokok[3] [4]. Jika tidak ditangani secara tepat, kondisi ini berpotensi meningkatkan angka kesakitan dan kematian dini akibat penyakit jantung.

Secara medis, penyakit jantung dapat dikenali melalui kombinasi gejala klinis, hasil pemeriksaan fisik, indikator laboratorium, serta rekam riwayat kesehatan pasien. Faktor risiko seperti darah tinggi, kolesterol, diabetes, serta riwayat keluarga memiliki peran penting dalam meningkatkan kemungkinan

seseorang mengalami gangguan jantung[5] [6]. Namun demikian, metode diagnosis tradisional yang mengandalkan analisis manual masih memiliki sejumlah keterbatasan, antara lain subjektivitas penilaian dokter, kebutuhan waktu yang relatif lama, serta kesulitan dalam mengolah data klinis dalam jumlah besar dan kompleks secara bersamaan. Keterbatasan ini dapat berpotensi memengaruhi akurasi diagnosis, khususnya pada tahap deteksi dini penyakit jantung.

Perkembangan teknologi informasi dan ketersediaan data medis dalam skala besar mendorong pemanfaatan teknik berdasarkan data, khususnya *machine learning*, dalam bidang kesehatan. *Machine learning* mampu mengidentifikasi pola tersembunyi dari data historis pasien dan menghasilkan model yang dapat membantu tenaga medis untuk sebuah keputusan klinis[7]. Pendekatan ini dinilai efektif dalam meningkatkan akurasi diagnosis, mempercepat proses analisis, serta mengurangi potensi kesalahan interpretasi terhadap data klinis yang kompleks.

Dalam konteks klasifikasi penyakit jantung, beberapa algoritma *machine learning* telah banyak digunakan, termasuk *Random Forest*, *Support Vector Machine (SVM)*, *Xgboost*, dan *K-Nearest Neighbors*. *Random Forest* adalah teknik kelompok pembelajaran Algoritma ini mengintegrasikan sejumlah pohon keputusan untuk mencapai hasil ramalan yang lebih tepat dan stabil. Dikenal efektif dalam mengelola data yang tidak linier, meminimalkan bahaya *overfitting*, serta merepresentasikan hubungan antara variabel klinis secara optimal [8] [9]. Sebaliknya, *Support Vector Machine (SVM)* terbukti andal dalam

menemukan *hyperplane* optimal untuk membatasi kategori data dengan jarak maksimal. Lebih dari itu, SVM mampu mengolah dataset yang sangat luas dan tidak mengikuti pola linier [10] [11]. Sementara itu, *XG Boost*, atau *Extreme Gradient Boosting*, merupakan metode pembelajaran mesin ensemble yang menerapkan pendekatan boosting melalui penggabungan berbagai *decision tree* [12].

K-Nearest Neighbors (KNN) yaitu model sederhana yang digunakan untuk klasifikasi dan regresi berdasarkan kedekatan jarak antar data [4].

Meskipun empat algoritma tersebut telah digunakan dalam sejumlah penelitian, terkait bagaimana kinerja *Random Forest*, *SVM*, *XG Boost* dan *KNN* dibandingkan secara langsung dalam konteks prediksi penyakit jantung berbasis faktor klinis terkini. Beberapa studi sebelumnya hanya berfokus pada satu algoritma, menggunakan dataset terbatas, atau tidak melibatkan evaluasi komprehensif seperti *accuracy*, *presisi*, *recall*, dan *F1-score*. Selain itu, masih diperlukan analisis mengenai algoritma mana yang lebih optimal untuk diterapkan sebagai sistem pendukung keputusan dalam skenario klinis modern.

Menurut uraian tersebut, penelitian ini sangat penting untuk memberikan kontribusi ilmiah dengan melakukan evaluasi komparatif yang lebih mendalam tentang bagaimana *Random Forest* dan *SVM* bekerja dalam klasifikasi penyakit jantung. Hasil penelitian diharapkan memberikan langkah menentukan keputusan di bidang kesehatan dan menjadi rujukan untuk pengembangan sistem deteksi dini berbasis data. Dengan demikian, riset ini bertujuan untuk menganalisis, membandingkan kinerja berbagai model dalam memprediksi penyakit jantung berdasarkan faktor klinis. Algoritma-algoritma ini termasuk *Random Forest*, *SVM*, *Xgboost*, dan *KNN*.

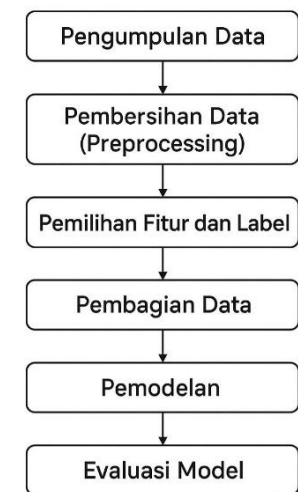
2. TINJAUAN PUSTAKA

Riset mengenai klasifikasi gagal jantung juga dilakukan oleh Laili Nur Farida dan Saiful Bahri [13]. Pada penelitian tersebut, algoritma *SVM* model dengan tingkat hasil 86,92%,. Selanjutnya riset mengenai klasifikasi penyakit jantung dengan memanfaatkan *Machine Learning* pada riset yang dilakukan oleh Deo Haganta Depar dkk [14]. riset menunjukkan bahwa setiap algoritma menghasilkan tingkat akurasi yang berbeda *Random Forest* dengan akurasi 75%. Kemudian riset Dewi Nasien dkk [15]. mengklasifikasikan Penyakit Jantung dengan *KNN* Menggunakan Fitur *PCA* model mencapai akurasi 81.82%.

3. METODE PENELITIAN

Penelitian kuantitatif merupakan Metode penelitian berfokus pada pengumpulan informasi berbentuk angka yang kemudian dievaluasi melalui prosedur statistik. Pendekatan semacam ini diterapkan untuk mengkaji gejala secara terstruktur dengan memanfaatkan data yang bisa dihitung menggunakan

alat statistik, matematika, atau komputasi. Penelitian ini menerapkan metode kuantitatif guna membandingkan keefektifan algoritma *Random Forest* dan *SVM*, *XGBoost*, *KNN* dalam proses klasifikasi penyakit jantung.



Gambar 1. Tahap Penelitian

3.1. Pengumpulan Data

Riset ini dengan data rekam medis pasien terkait penyakit jantung yang diperoleh dari *kaggle* [16]. Dataset berisi sejumlah atribut klinis yang umumnya digunakan dalam proses diagnosis penyakit jantung. Dataset ini terdiri dari 1888 record dan 14 atribut variabel.

3.2. Pembersihan Data (Preprocessing)

Sebelum dilakukan proses klasifikasi, dataset melalui tahapan *preprocessing* untuk memastikan kualitas data sehingga model dapat dilatih secara optimal. *Preprocessing* yaitu tahap perapian data yang mengandung kesalahan. Tujuan *preprocessing* yaitu untuk memperbaiki data tersebut [17].

3.3. Pemilihan Fitur dan Label

Dalam penelitian ini, fitur dan label dipilih berdasarkan ringkasan dengan klasifikasi penyakit jantung (usia), (jenis kelamin), (trestbps), (chol), (fbs), (thalach), (exang), (oldpeak), (slope), (ca), (restecg) adalah variabel klinis yang sering digunakan untuk menentukan diagnosis penyakit jantung. Fitur-fitur ini dipilih karena terkait langsung dengan kondisi kardiovaskular pasien dan tingkat risiko penyakit jantung [18].

Sementara itu, label (target) pada penelitian ini adalah variabel yang menunjukkan kondisi penyakit jantung pasien, yang dikodekan dalam dua kelas, yaitu kelas 0 yang merepresentasikan pasien tanpa penyakit jantung dan kategori 1 yang merepresentasikan pasien dengan penyakit jantung. Pemilihan label ini bertujuan untuk membangun model klasifikasi biner yang mampu membedakan secara jelas untuk pasien yang terkena dan non terkena penyakit jantung. Dengan pemilihan fitur serta label tepat, model diharapkan

mampu menghasilkan prediksi yang akurat dan dapat diandalkan[19].

3.4. Pembagian Data

Tahap berikutnya adalah membagi dataset menjadi data uji dan data latih, kemudian dievaluasi menggunakan data yang belum pernah dilihat sebelumnya sehingga hasil evaluasi lebih objektif dan menggambarkan kemampuan generalisasi model. Pada penelitian ini digunakan skema pembagian 80% training set dan 20% testing set [2] [16].

3.5. Random Forest

Random Forest merupakan *ensemble* yang berfungsi dengan mengintegrasikan berbagai pohon keputusan guna memprediksi yang lebih tebat dan lebih benar terhadap perubahan. Dalam proses klasifikasi, *Random Forest* membuat beberapa keputusan pohon secara bersamaan, setiap pohon memberikan hasil prediksi, dan kelas yang paling banyak melalui pilihan voting akan menjadi hasil akhir dari model tersebut [20] [21].

3.6. Support Vector Machine (SVM)

Merupakan model yang beroperasi menemukan garis atau *hyperplane* optimal untuk membedakan data ke dalam berbagai kategori dengan jarak terbesar. SVM menggunakan titik-titik penting yang disebut *support vector* dan dapat menangani data linear maupun non-linear melalui penggunaan kernel [22].

$$F(x) = \text{sign}(w \cdot x + b) \quad (1)$$

- Mengklasifikasikan data dengan melihat posisi data terhadap garis / *hyperplane*
- w dan b membentuk garis pemisah terbaik
- sign menentukan kelas (+1 atau -1)
- Data terdekat ke garis disebut *support vector*

3.7. XGBOOST

XGBoost (Extreme Gradient Boosting) adalah model berbasis *ensemble* yang menggunakan teknik *boosting*, yaitu menggabungkan banyak *decision tree* sederhana untuk membentuk model yang kuat dan akurat. *XGBoost* merupakan pengembangan dari algoritma *Gradient Boosting* yang dioptimalkan agar lebih cepat, efisien, dan akurat[12].

3.8. KNN (K Nearest Neighbors)

KNN adalah model sederhana yang digunakan untuk klasifikasi dan regresi berdasarkan kedekatan jarak antar data. KNN termasuk metode *lazy learning*, karena tidak membangun model saat pelatihan, melainkan menyimpan seluruh data latih dan melakukan perhitungan saat proses prediksi[23].

$$d(x, y) = \sum_{i=1}^n (x_i - y_i)^2 \quad (1)$$

- $d(x,y)$: jarak antara data uji dan data latih
- x : data uji yang akan diprediksi
- y : data latih
- x_i : nilai fitur ke- i pada data uji
- y_i : nilai fitur ke- i pada data latih
- n : jumlah fitur yang digunakan

3.9. Evaluasi Model

Evaluasi model adalah tahapan mengukur dan menilai seberapa baik suatu model dalam klasifikasi data. Evaluasi model dilakukan dengan menggunakan beberapa metrik seperti akurasi, presisi, recall, dan $F1_Score$ digunakan untuk mendapatkan gambaran lebih lengkap tentang performa model.

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan data

Sumber data dari *Kaggle* dipergunakan untuk mengklasifikasikan penyakit jantung berdasarkan faktor-faktor klinis. Data ini diperoleh pada tanggal 19 November 2025.

Tabel 1. Raw Data

no	age	sex	cp	trestbps	chol	fbs	restecg	thalachh	exang	oldpeak	slope	ca	thal	Trgt
0	63	1	3	145	233	1	0	150	0	2.3	0	0	1	1
1	37	1	2	130	250	0	1	187	0	3.5	0	0	2	1
2	41	0	1	130	204	0	0	172	0	1.4	2	0	2	1
3	56	1	1	120	236	0	1	178	0	0.8	2	0	2	1
4	57	0	0	120	354	0	1	163	1	0.6	2	0	2	1

Berdasarkan tabel 1, menyajikan informasi klinis tentang pasien yang berhubungan dengan kesehatan jantung, termasuk usia, gender, tekanan darah, kadar kolesterol, gula darah, hasil pemeriksaan EKG, frekuensi detak jantung tertinggi, reaksi saat berolahraga, serta parameter medis tambahan lainnya. Setiap baris merepresentasikan satu pasien, sementara kolom target menunjukkan hasil diagnosis, yaitu ada atau tidaknya penyakit jantung. Data ini digunakan untuk menganalisis hubungan faktor-faktor klinis dengan risiko penyakit jantung dan sebagai dasar pemodelan prediksi.

4.2. Pembersihan Data (Preprocessing)

Tahapan *preprocessing* yang dilakukan meliputi pengecekan *missing value*, encoding data kategorikal, dan normalisasi (*scaling*) data numerik.

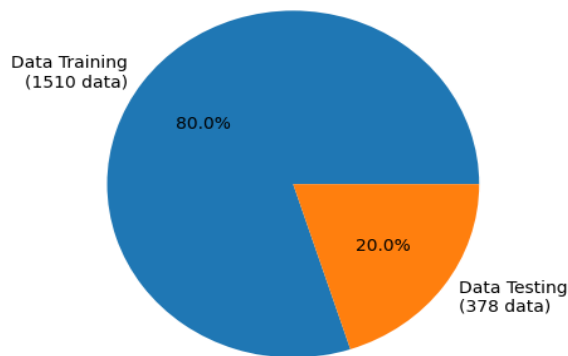
	age	sex	cp	trestbps	chol	fbs	restecg	thalachh	exang	oldpeak	slope	ca	thal	
0	-0.155167	0.667065	1.681797	-0.489197	0.712971	-0.412968	-0.927995	-2.351282	1.423395	0.801202	0.936587	1.224206	3.452187	
1	0.287330	0.667065	0.778055	-0.204473	-0.349398	-0.412968	2.238722	0.031328	-0.702546	-0.575152	0.936587	0.248717	3.452187	
2	1.173233	0.667065	1.681797	0.364974	0.712971	2.421495	-0.927995	1.071013	-0.702546	0.285069	-0.672339	0.248717	-0.521591	
3	-1.372032	0.667065	1.681797	0.934421	-0.048445	-0.412968	-0.927995	1.244293	-0.702546	-0.231064	0.936587	1.224206	-0.521591	
4	1.836068	-1.499104	0.778055	-1.229478	0.372890	2.421495	-0.927995	-0.835075	-0.702546	-0.919241	0.936587	0.248717	-0.521591	

Gambar 2. Preprocessing

Berdasarkan gambar 2, dataset fitur medis pasien yang sudah dibersihkan, digunakan sebagai input model *machine learning* untuk menganalisis atau memprediksi kondisi kesehatan (khususnya penyakit jantung).

4.3. Pembagian data

Pada tahap ini pembagian dataset menjadi data latih sebanyak 1510 dan data uji sebanyak 378 data.



Gambar 3. Split Data

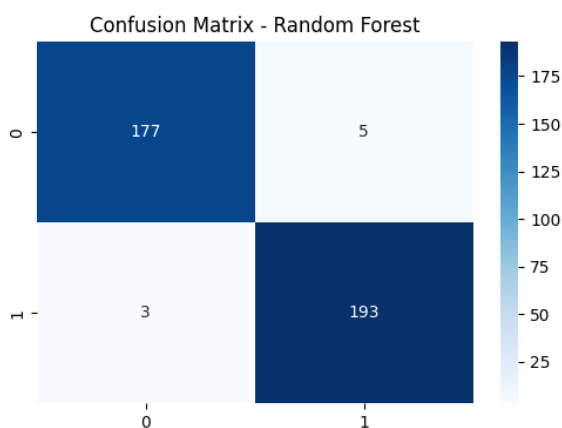
Berdasarkan gambar 3, data penelitian dibagi menjadi data latih sebesar 80% dan data uji sebesar 20%. Pembagian ini bertujuan agar model dapat dilatih menggunakan sebagian besar data dan diuji menggunakan data yang belum pernah dilihat sebelumnya untuk mengukur performa model secara objektif.

4.4. Evaluasi Model

Evaluasi menggunakan empat algoritma yaitu *XGBoost*, *Random Forest*, *Support Vector Machine* dan *K-Nearest Neighbors*.

4.5. Random Forest

Hasil Laporan Klasifikasi *Random Forest* menunjukkan kinerja model yang luar biasa dalam mengklasifikasikan data. Untuk kelas 0, model mencapai nilai akurasi sebesar 97.88%, dengan nilai presisi 0,98, recall 0,97, dan f1-score 0,98, dengan jumlah data 182.

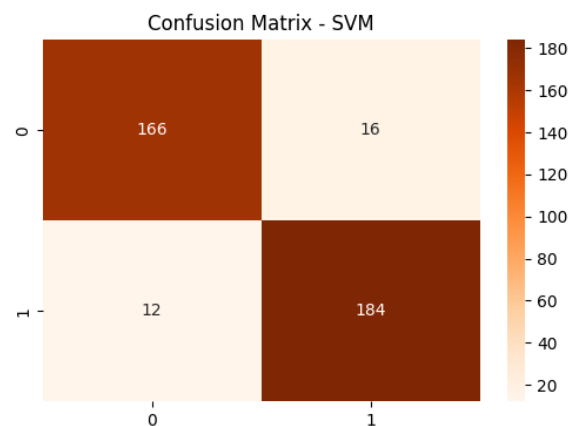


Gambar 3. Confusion Matrix RF

Gambar 3 menunjukkan *Random Forest*, yang menunjukkan kemampuan model untuk mengklasifikasikan data dengan sangat baik. 177 data dari kelas 0 diprediksi dengan benar, 5 data dari kelas 1 diprediksi dengan benar, dan hanya 3 data dari kelas 0 diprediksi dengan salah.

4.6. Support Vector Machine (SVM)

Hasil laporan klasifikasi *SVM* menunjukkan kinerja sangat baik dengan akurasi sebesar 92.59% dalam mengklasifikasikan data sentimen. Pada kelas negatif (0), model memperoleh nilai precision 0,93, recall 0,91, dan F1-score 0,92, sedangkan pada kelas positif (1) diperoleh precision 0,92, recall 0,94, dan F1-score 0,93, yang menunjukkan kemampuan model dalam mengenali kedua kelas secara efektif. Nilai macro average dan weighted average yang sama-sama berada pada angka 0,93 menandakan bahwa model *SVM* memiliki kinerja yang seimbang dan stabil pada seluruh kelas, sehingga dapat disimpulkan bahwa algoritma *SVM* layak digunakan untuk analisis sentimen dalam penelitian ini.



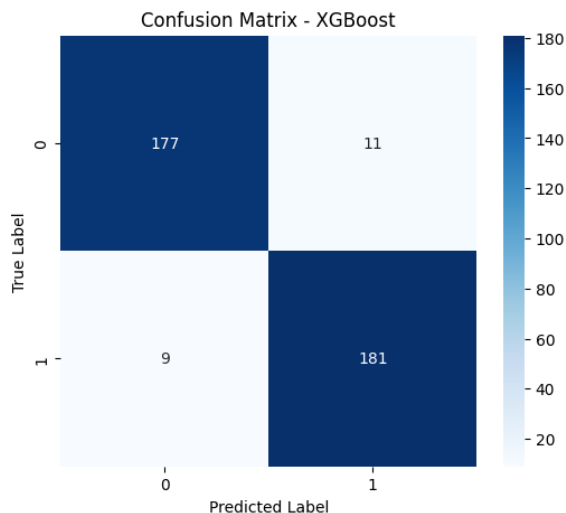
Gambar 4. Confusion Matrix SVM

Berdasarkan gambar 4, hasil dari matriks gangguan, model *Support Vector Machine* (*SVM*) berhasil mengklasifikasikan sentimen data dengan baik. Sebanyak 166 data sentimen negatif dan 184 data sentimen positif berhasil diprediksi secara benar. Meski begitu, masih ada 16 data sentimen negatif yang salah dikategorikan sebagai positif dan 12 data sentimen positif yang salah diberi label negatif. Hasil ini menunjukkan bahwa jumlah kesalahan klasifikasi yang terjadi tidak terlalu besar dibandingkan dengan jumlah prediksi yang benar, sehingga menunjukkan bahwa model *SVM* memiliki kemampuan yang cukup baik dan stabil dalam membedakan antara sentimen positif dan negatif pada data penelitian ini.

4.7. XGBoost

Hasil pengujian, algoritma *XGBoost* menghasilkan nilai akurasi sebesar 94,71%, yang menunjukkan bahwa model mampu mengklasifikasikan data dengan tingkat ketepatan yang sangat baik. Nilai *precision*, *recall*, dan *f1-score*

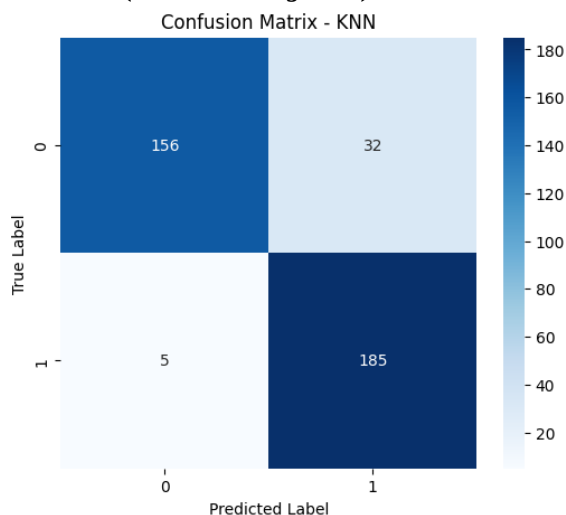
untuk kedua kelas berada pada kisaran 0,94–0,95, yang menandakan keseimbangan yang baik antara kemampuan model dalam memprediksi kelas secara tepat dan mendeteksi data pada masing-masing kelas. Dengan demikian, *XGBoost* dapat dikatakan memiliki performa yang stabil dan andal dalam melakukan klasifikasi pada dataset yang digunakan.



Gambar 5. *Confusion Matrix XGBoost*

Berdasarkan gambar 5, *Confusion matrix* menunjukkan bahwa model *XGBoost* mampu melakukan klasifikasi dengan sangat baik. Dari total 188 data kelas 0, sebanyak 177 data berhasil diprediksi dengan benar, sementara 11 data salah diklasifikasikan sebagai kelas 1. Untuk kelas 1, dari 190 data, sebanyak 181 data diprediksi dengan benar dan 9 data salah diprediksi sebagai kelas 0.

4.8. KNN (K-Nearest Neighbors)



Gambar 6. *Confusion Matrix KNN*

Dengan akurasi 90,21% yang dihasilkan oleh model KNN, sebagian besar data dapat diklasifikasikan dengan baik. Secara keseluruhan, nilai *F1-score* 0,90 menunjukkan performa model yang

baik dan cukup seimbang, dengan *recall* kelas 0 yang sangat tinggi (0,97) tetapi *recall* kelas 0 yang lebih rendah (0,83), yang menunjukkan bahwa prediksi kelas 0 cukup tepat tetapi masih ada data yang salah diklasifikasikan.

Berdasarkan Gambar 6, *Confusion Matrix KNN* menunjukkan bahwa model berhasil mengklasifikasikan 156 data kelas 0 dan 185 data kelas 1 dengan benar. Namun, terdapat 32 data kelas 0 yang salah diprediksi sebagai kelas 1 dan 5 data kelas 1 yang salah diprediksi sebagai kelas 0.

5. KESIMPULAN DAN SARAN

Dari hasil riset yang telah diterapkan melewati beberapa proses pembersihan, seperti memeriksa data yang hilang, mengubah kategori data menjadi bentuk yang dapat dipahami oleh komputer, serta menstandarkan nilai fitur, sehingga data siap digunakan untuk membuat model. Hasil menggambarkan jika algoritma *machine learning* mampu memberikan hasil yang baik dalam mengklasifikasikan penyakit jantung berdasarkan data klinis. Kinerja model menunjukkan bahwa algoritma *Random Forest* memiliki performa terbaik dengan akurasi mencapai 97,88%, untuk kedua kelas. Algoritma *XGBoost* dan *SVM* juga menunjukkan performa yang baik dengan akurasi masing - masing sebesar 94,71% dan 92,59%, sedangkan *KNN* mendapat akurasi 90,21%. Dari keempat algoritma tersebut, *Random Forest* lebih baik dalam menangani data klinis yang rumit dan memberikan prediksi yang lebih stabil. Oleh karena itu, *Random Forest* direkomendasikan sebagai model yang paling efektif untuk digunakan dalam sistem bantuan keputusan medis agar dapat membantu mendeteksi dini penyakit jantung. Untuk penelitian berikutnya, disarankan menggunakan dataset yang lebih besar, menambahkan fitur klinis lainnya, serta mengeksplorasi algoritma *deep learning* untuk meningkatkan akurasi dan kemampuan model dalam menerapkan data baru.

DAFTAR PUSTAKA

- [1] M. Anita *et al.*, “KLASIFIKASI FAKTOR RISIKO PENYAKIT JANTUNG MENGGUNAKAN MACHINE,” vol. 16, no. c, pp. 68–78, 2025.
- [2] M. Napiah and S. Heristian, “Perbandingan Algoritma Machine Learning pada Klasifikasi Penyakit Jantung,” vol. 6, no. 1, pp. 46–51, 2024.
- [3] U. A. Pringsewu, “Volume 6 Issue 1 Aisyah Journal of Informatics and Electrical Engineering Aisyah Journal of Informatics and Electrical Engineering,” vol. 6, no. 1, pp. 145–150.
- [4] D. Fabiyanto and Z. P. Putra, “Validasi Efektivitas Logistic Regression untuk Diagnosa Penyakit Jantung melalui Pendekatan Machine Learning,” vol. 16, no. 2, pp. 158–170, 2024.
- [5] P. Algoritma, R. Forest, N. Bayes, and A. S.

- Budiman, "Jurnal Sains Informatika Terapan (JSIT)," vol. 4, no. 2, pp. 77–84, 2025.
- [6] A. S. Prabowo and F. I. Kurniadi, "Analisis Perbandingan Kinerja Algoritma Klasifikasi dalam Mendeteksi Penyakit Jantung," 2023.
- [7] C. Algoritma, "KLASIFIKASI PENYAKIT JANTUNG DENGAN MENGGUNAKAN," vol. 7, no. 2, 2022.
- [8] V. N. Mei, D. Setiawan, A. Muhammad, S. Herawati, and F. Dewi, "Penerapan Algoritma Klasifikasi untuk Deteksi Dini Penyakit Jantung Koroner Berdasarkan Gejala Klinis koroner berdasarkan data klinis pasien menggunakan algoritma pembelajaran mesin [7]. komputasional berbasis algoritma pembelajaran mesin untuk klasifikasi penyakit jantung," vol. 5, pp. 18–26, 2025.
- [9] A. Sunyoto and H. Al Fatta, "Klasifikasi Penyakit Jantung Menggunakan Random Forest Clasifier," vol. VII, no. September, pp. 31–40, 2023.
- [10] A. Nurmasani and Y. Pristyanto, "ALGORITME R. P. See, "Klasifikasi Penyakit Jantung Menggunakan Decision Tree dan KNN Menggunakan Ekstraksi Fitur PCA," vol. 4, no. 1, pp. 1–6, 2024.
- [16] P. P. Jantung, F. Handayani, K. S. Kusuma, H. L. Asbudi, R. G. Purnasiwi, and R. Kusuma, "Komparasi Support Vector Machine , Logistic Regression Dan Artificial Neural Network dalam," vol. 7, no. 3, pp. 329–334, 2021.
- [17] M. K. Insan, U. Hayati, and O. Nurdian, "ANALISIS SENTIMEN APLIKASI BRIMO PADA ULASAN PENGGUNA DI," vol. 7, no. 1, pp. 478–483, 2023.
- [18] J. D. Muthohhar and A. Prihanto, "Analisis Perbandingan Algoritma Klasifikasi untuk Penyakit Jantung," vol. 04, pp. 298–304, 2023.
- [19] J. Jtik, J. Teknologi, B. Hirwono, A. Hermawan, and D. Avianto, "Implementasi Metode Naïve Bayes untuk Klasifikasi Penderita Penyakit Jantung," vol. 7, no. 3, 2023.
- STACKING UNTUK," vol. VIII, 2021.
- [11] P. Studi, M. Teknologi, U. Malikussaleh, B. Pulo, and K. Lhokseumawe, "PERBANDINGAN METODE MACHINE LEARNING MENGGUNAKAN METODE SUPPORT VECTOR MACHINE DAN ARTIFICIAL NEURAL," vol. 9, no. 2, 2025.
- [12] P. Kinerja *et al.*, "Jurnal Computer Science and Information Technology (CoSciTech)," vol. 3, no. 3, pp. 453–463, 2022.
- [13] L. N. Farida and S. Bahri, "Komputika : Jurnal Sistem Komputer Klasifikasi Gagal Jantung Menggunakan Metode SVM (Support Vector Machine) Classification of Heart Failure using the SVM (Support Vector Machine) Method," vol. 13, pp. 0–7, 2025, doi: 10.34010/komputika.v13i2.11330.
- [14] D. H. Depari *et al.*, "Perbandingan Model Decision Tree , Naive Bayes dan Random Forest untuk Prediksi Klasifikasi Penyakit Jantung," vol. 4221, pp. 239–248, 2022.
- [15] D. Nasien, R. Syahputra, A. A. Marunduri, and [20] J. Prima, J. Sistem, I. Komputer, P. Penyakit, and G. Jantung, "Prediksi penyakit gagal jantung dengan menggunakan random forest," vol. 5, no. 2, pp. 176–181, 2022.
- [21] L. Regression and R. Forest, "Perbandingan Logistic Regression , Random Forest , dan Perceptron pada Klasifikasi Pasien Gagal Jantung," vol. 14, no. 3, pp. 271–286, 2022.
- [22] K. Algoritma, K. S. Dan, C. Dalam, and P. Penyakit, "I n f o r m a t i k a," vol. 13, no. 2, pp. 31–41, 2021.
- [23] Y. Pratama, A. Prayitno, D. Nazrian, N. Aini, Y. R. R, and E. Rasywir, "BULLETIN OF COMPUTER SCIENCE RESEARCH Klasifikasi Penyakit Gagal Jantung Menggunakan Algoritma K-Nearest Neighbor," vol. 3, no. 1, pp. 52–56, 2022, doi: 10.47065/bulletincsr.v3i1.203.