



Analisis Performa XGBoost dan Gaussian Naive Bayes untuk Klasifikasi Dini Penyakit Hipertensi

Ni Made Ochiana Septhi Pratiwi^{1*}, Adie Wahyudi Oktavia Gama²

^{1,2}Universitas Pendidikan Nasional, Indonesia

E-mail: ochianasp.made@gmail.com¹, adiewahyudi@undiknas.ac.id²

*Penulis korespondensi: ochianasp.made@gmail.com

Abstract. Hypertension is one of the leading causes of premature death globally that often goes undetected due to minimal clinical symptoms, earning it the nickname “silent killer.” The application of artificial intelligence (AI), particularly Machine Learning, is a strategic approach to early detection, but the main challenge lies in balancing diagnostic accuracy with detection sensitivity so that no patients at risk are overlooked. This study aims to analyze and compare the performance of the Extreme Gradient Boosting (XGBoost) algorithm with the Cost-Sensitive strategy compared to Gaussian Naive Bayes (GNB) as a baseline in hypertension risk classification. The dataset used included 1,985 electronic medical records with 9 clinical attributes, which were evaluated using the 10-Fold Cross-Validation method to determine model validity. The test results showed that XGBoost consistently outperformed GNB across all evaluation metrics. XGBoost recorded superior performance with an Accuracy of 92.19% and an AUC of 0.9752, far surpassing GNB, which obtained an Accuracy of 84.13%. The application of Cost-Sensitive Learning in XGBoost proved effective in overcoming performance trade-offs by producing a Recall of 91.26% and a Precision of 93.53%. Furthermore, Feature Importance analysis identified Blood Pressure History, Smoking Status, and Family History as the most dominant risk factors, which is in line with global medical guidelines. Based on these results, it is concluded that XGBoost is a more reliable and accurate method to be applied in early detection systems for hypertension compared to classical probabilistic approaches.

Keywords: Early Detection; Gaussian Naive Bayes; Hypertension; Machine Learning; XGBoost.

Abstrak. Hipertensi adalah salah satu penyebab utama kematian dini di seluruh dunia yang seringkali tidak terdeteksi karena gejala klinisnya minimal, sehingga dijuluki “pembunuh senyap”. Penerapan kecerdasan buatan (AI), khususnya *Machine Learning*, merupakan pendekatan strategis untuk deteksi dini, tetapi tantangan utamanya terletak pada menyeimbangkan akurasi diagnostik dengan sensitivitas deteksi agar tidak ada pasien berisiko yang terlewatkan. Studi ini bertujuan untuk menganalisis dan membandingkan kinerja algoritma *Extreme Gradient Boosting* (XGBoost) dengan strategi *Cost-Sensitive* dibandingkan dengan *Gaussian Naive Bayes* (GNB) sebagai dasar dalam klasifikasi risiko hipertensi. Dataset yang digunakan mencakup 1.985 rekam medis elektronik dengan 9 atribut klinis, yang dievaluasi menggunakan metode *10-Fold Cross-Validation* untuk menentukan validitas model. Hasil pengujian menunjukkan bahwa XGBoost secara konsisten mengungguli GNB di semua metrik evaluasi. XGBoost mencatatkan kinerja superior dengan Akurasi 92,19% dan AUC 0,9752, jauh melampaui GNB yang memperoleh Akurasi 84,13%. Penerapan Pembelajaran Sensitif Biaya (*Cost-Sensitive Learning*) dalam XGBoost terbukti efektif dalam mengatasi trade-off kinerja dengan menghasilkan Recall 91,26% dan Presisi 93,53%. Lebih lanjut, analisis Kepentingan Fitur (*Feature Importance*) mengidentifikasi Riwayat Tekanan Darah, Status Merokok, dan Riwayat Keluarga sebagai faktor risiko paling dominan, yang sejalan dengan pedoman medis global. Berdasarkan hasil ini, disimpulkan bahwa XGBoost adalah metode yang lebih andal dan akurat untuk diterapkan dalam sistem deteksi dini hipertensi dibandingkan dengan pendekatan probabilistik klasik.

Kata kunci: Deteksi Dini; Gaussian Naive Bayes; Hipertensi; Pembelajaran Mesin; XGBoost.

1. PENDAHULUAN

Saat ini Penyakit Tidak Menular (PTM) menjadi penyebab utama kematian dan kecacatan di seluruh dunia, termasuk di Indonesia. Penyakit hipertensi menjadi salah satu kontributor utama dalam kelompok tersebut (Septian et al., 2025). Secara global, hipertensi juga menjadi faktor risiko utama untuk penyakit kardiovaskular dan stroke (Boateng & Ampofo, 2023). Penyakit ini juga menjadi salah satu tantangan besar bagi sistem kesehatan

global dan menjadi penyebab utama kematian dini di seluruh dunia (n.d.). Setiap tahunnya, hipertensi menyebabkan sekitar 8 juta kematian, dengan sekitar 1,5 juta kasus kematian di Asia Tenggara (Yana Afrina, 2024). Menurut *World Health Organization* (WHO), diperkirakan lebih dari 1,28 miliar orang dewasa di seluruh dunia menderita penyakit hipertensi, dengan dua pertiga di antaranya tinggal di negara-negara berpendapatan rendah hingga menengah (2023). Penyakit ini sering disebut sebagai *silent killer* karena banyak penderitanya tidak menunjukkan gejala klinis yang jelas hingga terjadi komplikasi serius, seperti stroke, gagal jantung, dan penyakit ginjal kronis sehingga sulit terdeteksi (Pokharel et al., 2022). Kondisi ini semakin diperburuk oleh rendahnya kesadaran masyarakat terhadap gejala hipertensi, data menunjukkan sekitar 46% penderita tidak menyadari bahwa mereka mengidap penyakit tersebut (2023).

Di Indonesia, prevalensi hipertensi terus meningkat setiap tahunnya yang dipengaruhi oleh berbagai faktor risiko, seperti usia, indeks massa tubuh (IMT), gaya hidup serta riwayat keluarga (Yana Afrina, 2024). Situasi ini menjadikan upaya deteksi dini sangat krusial untuk mengurangi dampak buruk yang ditimbulkan (Kario et al., 2024). Akan tetapi, peningkatan jumlah pasien yang signifikan menyebabkan akumulasi data rekam medis dalam volume yang sangat besar, sehingga sulit ditangani dengan metode konvensional. Seiring dengan perkembangan teknologi, penerapan kecerdasan buatan (AI), khususnya *Machine Learning* pada bidang kesehatan dapat menjadi solusi untuk mengolah data rekam medis yang kompleks guna mendukung upaya deteksi dini penyakit hipertensi secara tepat waktu dan akurat (Narmilan et al., 2022; Mroz et al., 2024). Namun, tantangan utama dalam diagnosis klinis saat ini adalah besarnya volume data kesehatan pasien yang sulit dianalisis secara manual dengan cepat dan akurat (Badawy et al., 2024). Dengan menerapkan algoritma *Machine Learning* terbukti efektif dalam menganalisis pola dari berbagai faktor risiko untuk menghasilkan prediksi yang akurat (Zhao et al., 2021).

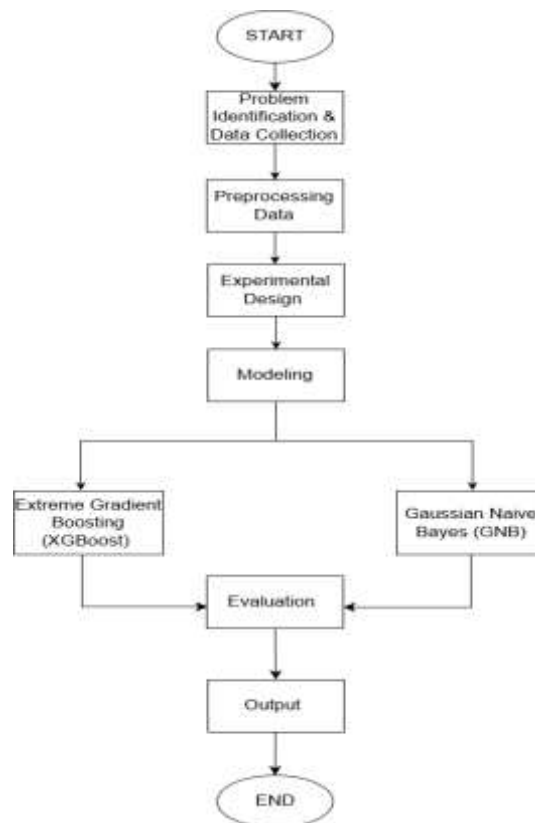
Berbagai algoritma klasifikasi telah diterapkan dalam studi bidang kesehatan untuk prediksi penyakit. Dalam studi terbaru terkait prediksi penyakit kardiovaskular, (Rayadhani & Rahardi, 2025a) melakukan komparasi antara algoritma *Random Forest*, *SVM* dan *Naive Bayes*. Hasil penelitian menunjukkan bahwa meskipun *Random Forest* dan *SVM* memiliki akurasi yang tinggi, namun keduanya membutuhkan waktu komputasi yang lebih lama. Sebaliknya, *Naive Bayes* terbukti jauh lebih efisien dari segi waktu namun dengan tingkat akurasi yang sedikit di bawah metode *ensemble*. Di sisi lain, untuk meningkatkan akurasi pada kasus data medis yang lebih kompleks *XGBoost* muncul sebagai algoritma andalan. Dalam penelitian (Maulana As'an Hamid & Egia Rosi Subhiyakto, 2025), membuktikan bahwa

XGBoost memiliki kemampuan dalam menangani data tidak seimbang (*imbalanced data*) dan meningkatkan akurasi model secara signifikan melalui mekanisme boosting, yang mengunggulkan *XGBoost* dibandingkan algoritma KNN dan SVM pada klasifikasi penyakit jantung.

Meskipun *XGBoost* menawarkan akurasi yang superior, algoritma ini memiliki kompleksitas model yang tinggi, yang sangat bertolak belakang dengan karakteristik GNB yang mengedepankan kesederhanaan probabilistik dan efisiensi komputasi (Barus et al., 2026). Upaya untuk membandingkan kedua pendekatan yang kontras ini telah dilakukan oleh (Barus et al., 2026) pada kasus prediksi penyakit diabetes. Hasilnya menempatkan *XGBoost* lebih unggul dalam akurasi, namun tetap mengakui *Naive Bayes* sebagai *baseline* model yang kompetitif dan cepat. Akan tetapi, penelitian yang membandingkan kedua algoritma ini secara langsung (*head-to-head*) pada kasus risiko hipertensi dengan fokus mengukur *trade-off* antara akurasi dan *Recall* (sensitivitas) masih perlu dieksplorasi secara mendalam. Nilai *Recall* menjadi aspek vital dalam diagnosis medis untuk meminimalkan kesalahan deteksi pada pasien yang sebenarnya sakit (*False Negative*).

Berdasarkan kesenjangan penelitian tersebut, penelitian ini bertujuan untuk menganalisis dan membandingkan kinerja algoritma *XGBoost* dan *Gaussian Naive Bayes* (GNB) dalam klasifikasi risiko hipertensi. Penelitian ini berfokus pada evaluasi performa kedua algoritma dalam menangani dataset medis untuk menentukan model terbaik. Untuk memastikan keandalan model dan menghindari bias dalam evaluasi, penelitian ini menerapkan teknik validasi *K-Fold Cross Validation* dengan nilai $K = 10$. Hasil dari penelitian ini diharapkan dapat memberikan rekomendasi algoritma yang tidak hanya akurat, tetapi juga memiliki sensitivitas (*Recall*) yang tinggi serta efisien untuk diterapkan pada sistem deteksi dini di fasilitas kesehatan.

2. METODE

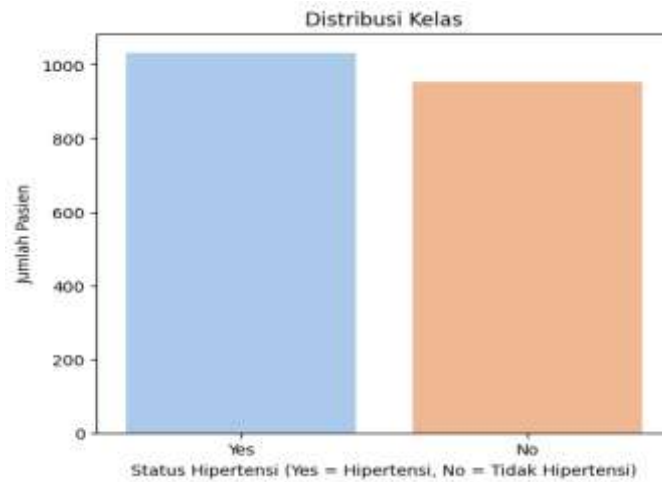


Gambar 1. Flowchart Metodologi Penelitian (*Research Methodology Flowchart*).

Tahapan penelitian ini dirancang secara sistematis pada Gambar 1 untuk membandingkan kinerja algoritma *XGBoost* dan *Gaussian Naive Bayes* (GNB) dalam klasifikasi risiko hipertensi.

Pengumpulan Data (*Data Collection*)

Data yang digunakan adalah data sekunder berupa rekam medis elektronik (*Electronic Medical Record*) yang diperoleh dari repositori publik Kaggle. Dataset ini terdiri dari 1.985 data pasien yang mencakup 9 atribut klinis, termasuk label target biner (Hipertensi dan Tidak Hipertensi). Sebelum memasuki tahap pembagian data (*splitting data*) menjadi data latih (*training data*) dan data uji (*testing data*), distribusi kelas pada keseluruhan dataset dianalisis terlebih dahulu untuk mengidentifikasi potensi ketidakseimbangan (*imbalance*) yang dapat mempengaruhi kinerja algoritma. Berdasarkan hasil eksplorasi awal, menunjukkan bahwa distribusi kelas antara kategori Tidak Hipertensi dan Hipertensi relatif seimbang secara alami. Jumlah sampel pada kelas Hipertensi (Label Yes) tercatat sebanyak 1032 data, sedangkan kelas Tidak Hipertensi (Label No) berjumlah 953 data.



Gambar 2. Visualiasi Distribusi Kelas pada Dataset Awal.

Gambar 2 memperlihatkan perbandingan jumlah sampel yang cukup seimbang antara pasien Hipertensi dan Tidak Hipertensi. Mengingat rasio ketidakseimbangan yang sangat rendah yaitu 0.92, penelitian ini memutuskan untuk tidak menerapkan teknik *resampling* seperti SMOTE. Keputusan ini diambil untuk menjaga orisinalitas karakteristik data medis dan menghindari potensi *noise* atau bias yang mungkin timbul dari pembuatan data buatan. Distribusi alami ini dinilai sudah representatif bagi algoritma *XGBoost* dan GNB untuk digunakan dalam proses pembagian data latih dan uji serta validasi model.

Pra-pemrosesan Data (*Preprocessing Data*)

Tahapan ini bertujuan untuk meningkatkan kualitas data sebelum diproses ke tahap pemodelan. Proses ini mencakup tiga langkah utama, yaitu: 1) Pembersihan Data (*Data Cleaning*): Melakukan pemeriksaan terhadap nilai yang hilang (*missing values*) dan penghapusan data duplikat. Langkah ini memastikan bahwa model dilatih menggunakan data yang bersih dan konsisten, sehingga meminimalkan bias pada hasil prediksi. 2) Transformasi Data (*Label Encoding*): Mengubah atribut kategorikal (seperti Gender) menjadi format numerik (seperti, Wanita = 0, Pria = 1) agar dapat diproses oleh algoritma matematika untuk memproses variabel kualitatif. 3) Normalisasi Fitur (*Feature Scaling*): Teknik *Min-Max Scaler* diterapkan untuk mentransformasi nilai fitur numerik ke dalam rentang 0 hingga 1 (Rayadhani & Rahardi, 2025). Normalisasi ini diimplementasikan secara spesifik pada alur kerja (*pipeline*) algoritma GNB untuk mengatasi sensitivitas terhadap variasi skala data serta menjaga stabilitas performa model. Sedangkan untuk algoritma *XGBoost*, fitur digunakan dalam skala aslinya karena algoritma berbasis pohon keputusan ini tidak dipengaruhi oleh perbedaan skala antar atribut.

Skenario Eksperimen

Skenario pengujian dirancang dalam tiga tahapan utama untuk menjamin validitas hasil, yaitu: 1) Pembagian Data (*Data Splitting* → Dataset dibagi menggunakan teknik *Stratified Random Sampling* dengan rasio 80:20 (80% untuk pelatihan, 20% untuk pengujian). Teknik ini memastikan proporsi kelas Hipertensi dan Tidak Hipertensi pada data latih dan data uji tetap konsisten dengan distribusi aslinya. 2) Validasi Model → Penelitian ini menerapkan metode *10-Fold Cross-Validation* pada proses pelatihan. Data latih dibagi menjadi 10 bagian untuk divalidasi secara berulang untuk mencegah *overfitting* dan menguji stabilitas performa model. Metode ini dipilih untuk memastikan objektivitas hasil evaluasi model karena seluruh data memiliki kesempatan yang sama untuk menjadi data latih (*training data*) maupun data uji (*testing data*) (Surur et al., 2025). 3) Strategi Pemodelan (*Cost-Sensitive*) → GNB digunakan sebagai *baseline*, sedangkan XGBoost menerapkan strategi *Cost-Sensitive* (*scale_pos_weight* > 1). Konfigurasi ini memberikan prioritas lebih tinggi padaantisipasi kesalahan kelas positif untuk meminimalkan *False Negative* dan memaksimalkan *Recall*.

Arsitektur Pemodelan (*Modeling*)

Penelitian ini membandingkan kinerja dua algoritma *Machine Learning* dengan karakteristik pendekatan yang berbeda, yaitu berbasis *ensemble learning* dan probabilistik.

Extreme Gradient Boosting (XGBoost)

XGBoost merupakan algoritma *Ensemble Learning* berbasis *Gradient Boosting Decision Tree* (GBDT) yang memanfaatkan kumpulan prediktor untuk meningkatkan akurasi prediksi secara signifikan (Dhanka & Maini, 2025). Algoritma ini bekerja dengan membangun pohon keputusan secara sekuensial, di mana setiap iterasi bertujuan untuk belajar dari kesalahan sebelumnya dan memperbarui *residual error*, sehingga menghasilkan model yang semakin presisi (Jinbo et al., 2025). Keunggulan utama XGBoost terletak pada efisiensinya dalam menangani data yang kompleks dan non-linier, serta kemampuannya yang terbukti handal dalam mengatasi masalah ketidakseimbangan data (*imbalanced data*) pada kasus klasifikasi (Dong et al., 2024). Dengan meminimalkan fungsi *loss* pada setiap tahapan, XGBoost mampu mencapai performa prediksi tinggi dengan tetap menjaga efisiensi komputasi yang optimal (Afifah et al., 2021).

Secara matematis, optimasi model tersebut dicapai dengan meminimalkan fungsi tujuan (*objective function*) yang terdiri dari komponen *loss function* dan regularisasi untuk mencegah *overfitting*. Hal tersebut dinyatakan dalam persamaan berikut:

$$Obj(\theta) = \sum_{i=1}^n L(y_i, \hat{y}_i) + \sum_{k=1}^k \Omega(f_k) \quad (1)$$

Keterangan:

Persaman (1) terdiri dari fungsi *loss* $L(y_i, \hat{y}_i)$ yang mengukur selisih prediksi antara nilai prediksi ($\hat{y}_i$) dan aktual (y_i). Sedangkan $\Omega(f_k)$ adalah komponen regularisasi yang mengontrol kompleksitas model untuk mencegah *overfitting*.

Gaussian Naive Bayes (GNB)

Berbeda dengan *XGBoost* yang berbasis pohon keputusan, *Gaussian Naive Bayes* (GNB) merupakan metode klasifikasi probabilistik yang didasarkan pada *Teorema Bayes* dengan asumsi independensi antar fitur (Wang et al., 2025). Varian ini dirancang khusus untuk menangani dataset dengan fitur kontinu, dan secara teoritis memiliki tingkat kesalahan terkecil dengan memanfaatkan seluruh atribut dalam proses prediksi (Benghazouani et al., 2025). Meskipun algoritma ini menggunakan asumsi independensi yang kuat sehingga terkadang kurang mempertimbangkan relevansi fitur yang tidak signifikan, GNB mampu menunjukkan efisiensi klasifikasi yang sangat stabil dan berkinerja baik, terutama pada dataset berskala kecil (Zhang, 2025). Algoritma ini terus dikembangkan dalam sistem kecerdasan buatan (AI) medis untuk memprediksi gangguan kesehatan dengan membandingkan parameter probabilitas tertentu (Anbazhagan & Rangaswamy, 2025). Mengingat karakteristik data yang bersifat kontinu, probabilitas *likelihood* untuk setiap fitur dihitung menggunakan fungsi densitas probabilitas *Gaussian* yang dirumuskan sebagai berikut:

$$P(x_i | y) = \frac{1}{\sqrt{2\pi} \sigma_{y^2}} \exp \left(-\frac{(x_i - \mu_y)^2}{2 \sigma_{y^2}} \right) \quad (2)$$

Keterangan :

Di mana $P(x_i | y)$ adalah probabilitas fitur x_i pada kelas y , μ_y adalah rata-rata (*mean*) dari fitur pada kelas y , dan σ_{y^2} adalah *varians* dari fitur tersebut.

#proposal ocik

Evaluasi (Evaluation)

Kinerja kedua model dievaluasi menggunakan *Confusion Matrix* dengan empat komponen dasar, yaitu: *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Berdasarkan komponen tersebut, kinerja model diukur menggunakan lima metrik utama, sebagai berikut:

Akurasi (*Accuracy*) : Mengukur tingkat ketepatan prediksi secara keseluruhan.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (3)$$

Presisi (*Precision*) : Mengukur seberapa akurat model saat memprediksi kelas positif. Metrik ini menjadi indikator prioritas untuk keselamatan pasien guna meminimalkan *False Positive*.

$$Precision = \frac{TP}{TP + FP} \quad (4)$$

Sensitivitas (*Recall*) : Mengukur kemampuan model dalam mendeteksi seluruh pasien yang benar-benar sakit. Metrik ini menjadi indikator prioritas untuk keselamatan pasien guna meminimalkan *False Negative*.

$$Recall = \frac{TP}{TP + FN} \quad (5)$$

F1-Score : Rata-rata harmonis antara *precision* dan *recall*, dengan memberikan gambaran performa yang lebih seimbang dibandingkan akurasi.

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (6)$$

Area Under Curve (AUC-ROC) : Kurva ROC (*Receiver Operating Characteristic*) memvisualisasikan *trade-off* antara *True Positive Rate* dan *False Positive Rate*. Sedangkan, nilai AUC digunakan untuk mengukur kemampuan model dalam membedakan antar kelas secara keseluruhan. Nilai AUC mendekati 1 mengindikasikan performa klasifikasi yang optimal.

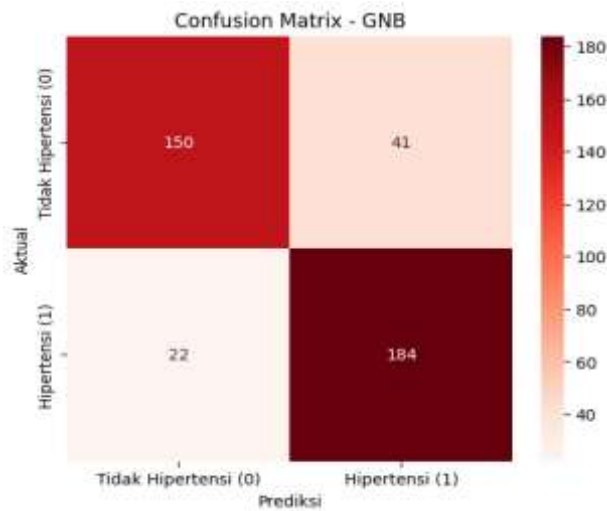
Selain metrik performa di atas, evaluasi juga mencakup analisis terhadap *Feature Importance* (khusus pada model *XGBoost*) untuk memberikan interpretasi medis mengenai atribut klinis mana yang paling berkontribusi terhadap risiko hipertensi.

3. HASIL DAN PEMBAHASAN

Bagian ini memaparkan data kuantitatif yang diperoleh dari eksperimen klasifikasi risiko hipertensi menggunakan algoritma *Gaussian Naive Bayes* (GNB) dan *Extreme Gradient Boosting* (XGBoost).

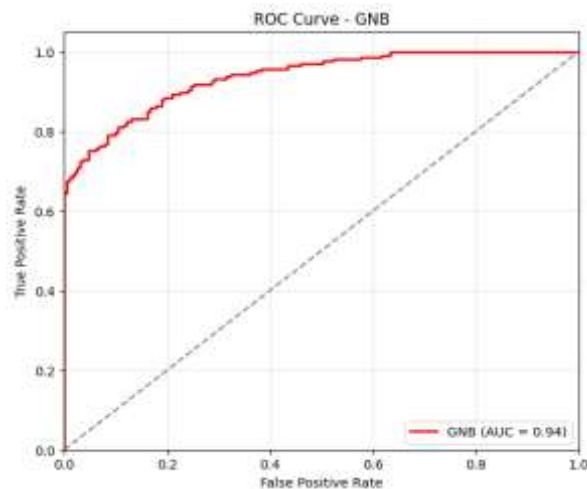
Hasil Pengujian Model *Gaussian Naive Bayes* (GNB)

Pengujian model GNB dilakukan pada data uji (*testing data*) setelah melalui proses *preprocessing*. Rincian hasil klasifikasi model GNB divisualisasikan melalui *Confusion Matrix* pada Gambar 3. Berdasarkan pengujian terhadap 397 sampel data, model GNB mampu mengidentifikasi kelas dengan benar (TP & TN) sebanyak 334 kali. Sementara itu, kesalahan klasifikasi (FP & FN) terjadi pada 63 data sisanya.



Gambar 3. *Confusion Matrix GNB.*

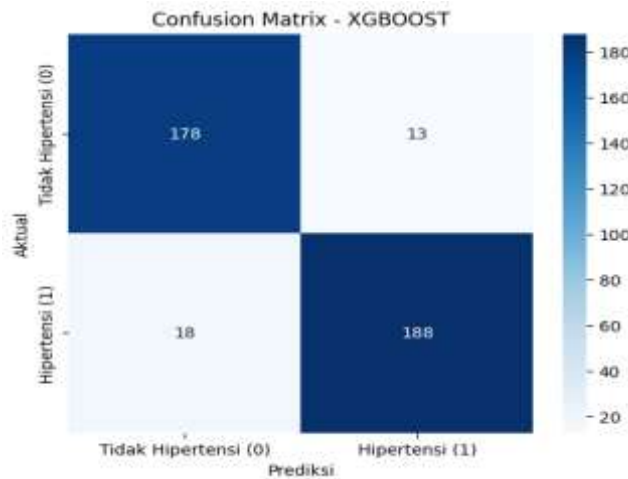
Selain itu, performa model dalam memisahkan kelas positif dan negatif dengan nilai *Area Under Curve* (AUC) sebesar 0.94 ditunjukkan melalui kurva ROC pada Gambar 4.



Gambar 4. *ROC Curver GNB.*

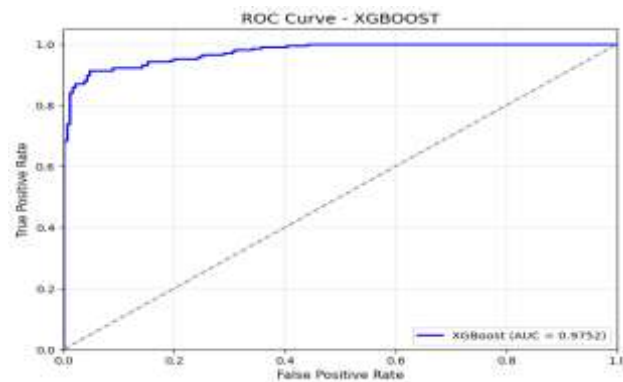
Hasil Pengujian Model *XGBoost*

Berbeda dengan GNB yang memerlukan normalisasi (*Min-Max Scaler*), pengujian model *XGBoost* dilakukan secara langsung pada data uji menggunakan konfigurasi model yang telah ditentukan. Visualisasi performa model *XGBoost* divisualisasikan melalui *Confusion Matrix* pada Gambar 5. Dari total 397 sampel data uji yang dievaluasi, model *XGBoost* berhasil mengidentifikasi kelas dengan benar (TP & TN) sebanyak 366 kali. Di sisi lain, terlihat bahwa jumlah kesalahan klasifikasi (FP & FN) menurun menjadi 31 dibandingkan model GNB.



Gambar 5. Confusion Matrix XGBoost.

Selain itu, kurva ROC XGBoost pada Gambar 6 menunjukkan garis lebih mendekati sudut kiri atas dibandingkan GNB, dengan nilai AUC mencapai 0.97.



Gambar 6. ROC Curve XGBoost.

Komparasi Kinerja Model

Perbandingan langsung (*head-to-head*) antara performa XGBoost dan GNB pada data uji disajikan pada Gambar 7.



Gambar 7. Grafik Perbandingan performa GNB vs XGBoost.

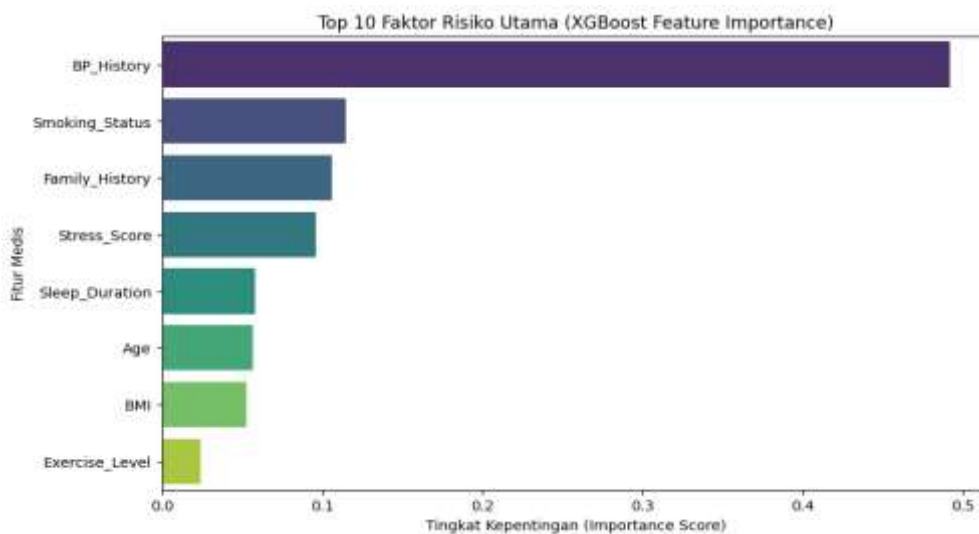
Metrik Evaluasi	XGBoost	GNB	Selisih (XGB - GNB)
0 Akurasi	92.19%	84.13%	+8.06%
1 Presisi	93.53%	81.78%	+11.75%
2 Recall	91.26%	89.32%	+1.94%
3 F1-Score	92.38%	85.38%	+7.00%
4 AUC Score	0.9752	0.9367	+0.0386

Gambar 8. Selisih Performa XGBoost vs GNB.

Gambar 7 dan 8 menunjukkan dominasi *XGBoost* di seluruh metrik, dengan peningkatan signifikan pada Presisi (+11.75%) dan Akurasi (+8.06%). Strategi *Cost-Sensitive* juga terbukti efektif menjaga keunggulan *Recall* (+1.94%) dibandingkan dengan *baseline*, dengan nilai AUC mencapai 0.9752.

Identifikasi Faktor Risiko Utama (*Feature Importance*)

Berdasarkan model *XGBoost*, identifikasi fitur yang paling berpengaruh terhadap prediksi ditampilkan pada Gambar 9.



Gambar 9. Top 10 Faktor Risiko Utama Pada Prediksi Hipertensi.

Gambar 9 menunjukkan Riwayat Tekanan Darah (*BP_History*) sebagai fitur paling dominan (skor ~0.5), diikuti dengan Status Merokok (*Smoking_Status*) dan Riwayat Keluarga (*Family_History*). Hasil ini membuktikan bahwa model sangat bergantung pada riwayat klinis pasien.

Pembahasan

Analisis Komparatif Kinerja Algoritma

Hasil pengujian menunjukkan bahwa algoritma *Extreme Gradient Boosting* (*XGBoost*) secara konsistensi mengungguli *Gaussian Naive Bayes* (GNB), dengan keunggulan paling menonjol terletak pada Presisi yang mencapai 93.53%. Hal ini mengindikasikan kemampuan model yang sangat baik dalam meminimalkan *False Positive* sehingga menghasilkan tingkat

kepercayaan prediksi yang tinggi. Keunggulan ini diatribusikan pada arsitektur *Gradient Boosting* yang mampu menangkap pola non-linear kompleks, temuan ini sejalan dengan studi Yadav dkk. (2024) terkait diagnosis penyakit jantung yang membuktikan XGBoost menjadi model dengan performa terbaik mengungguli metode probabilistik klasik (Yadav et al., 2024)

Selain itu, penerapan strategi *Cost-Sensitive Learning* terbukti efektif menjaga Recall tetap tinggi (91.26%). Biasanya, peningkatan presisi akan menurunkan recall (*precision-recall trade-off*), namun dengan memberikan bobot penalti lebih besar (*scale_pos_weight*) pada kesalahan kelas positif, model berhasil menyeimbangkan keduanya. Kombinasi antara ketepatan diagnosis (Presisi tinggi) dan sensitivitas deteksi (*Recall* tinggi) menjadikan XGBoost sangat reliabel sebagai instrumen awal skrining medis.

Interpretasi Medis Faktor Risiko (Clinical Relevance)

Analisis *Feature Importance* menetapkan atribut Riwayat Tekanan Darah (*BP_History*) sebagai determinan utama. Temuan ini sangat relevan secara medis dan sejalan dengan pedoman klinis JNC-8 (*Joint National Committee*), yang menyatakan bahwa riwayat elevasi tekanan darah sebelumnya adalah indikator fundamental dalam diagnosis dan manajemen hipertensi kronis (Mancia et al., 2023). Fitur lain yang berpengaruh signifikan adalah Status Merokok (*Smoking_Status*) dan Riwayat Keluarga (*Family_History*). Secara patofisiologis, kandungan nikotin dalam rokok diketahui menstimulasi sistem saraf simpatis untuk memproduksi katekolamin, yang menyebabkan vasokonstriksi (penyempitan pembuluh darah), serta peningkatan beban kerja jantung secara signifikan (Münzel et al., 2025).

Sementara itu, dominasi atribut riwayat keluarga memvalidasi bahwa hipertensi memiliki komponen genetik yang kuat (*hereditary*), di mana individu dengan riwayat keluarga positif memiliki predisposisi yang jauh lebih tinggi akibat interaksi poligenik yang kompleks (Takase et al., 2025).

Keterbatasan Penelitian

Meskipun menghasilkan performa yang baik, penelitian ini memiliki beberapa keterbatasan yang perlu diperhatikan. Pertama, penggunaan dataset sekunder dari repositori publik (Kaggle) berpotensi memunculkan bias demografis, mengingat karakteristik populasi global mungkin tidak sepenuhnya merepresentasikan heterogenitas pasien di Indonesia, sehingga uji validitas model perlu diuji lebih lanjut. Kedua, pemodelan saat ini terbatas pada 9 atribut klinis fundamental. Penambahan atribut klinis lain seperti kadar gula darah, profil lipid lengkap, serta rekaman aktivitas fisik yang mendetail direkomendasikan untuk studi mendatang guna meningkatkan presisi deteksi dini.

4. KESIMPULAN

Penelitian ini membuktikan bahwa algoritma *Extreme Gradient Boosting (XGBoost)* secara konsisten mengungguli *Gaussian Naive Bayes (GNB)* dalam klasifikasi risiko hipertensi. *XGBoost* mencatatkan performa superior dengan Akurasi 92.19% dan AUC 0.9752, melampaui GNB yang hanya mencapai Akurasi 84.13%. Keunggulan utama model ini terletak pada penerapan strategi *Cost-Sensitive Learning* yang terbukti efektif menyeimbangkan *trade-off* performa dengan menghasilkan Recall 91.26% dan Presisi 93.53%.

Selain keunggulan statistik, model juga menunjukkan validitas klinis yang kuat melalui *Feature Importance*. Identifikasi Riwayat Tekanan Darah, Status Merokok, dan Riwayat Keluarga memvalidasi relevansi klinis model terhadap pedoman medis global. Dengan demikian, *XGBoost* direkomendasikan sebagai metode yang lebih andal dan akurat untuk sistem deteksi disini hipertensi dibandingkan pendekatan probabilistik klasik.

DAFTAR PUSTAKA

- Afifah, K., Yulita, I. N., & Sarathan, I. (2021). Sentiment Analysis on Telemedicine App Reviews using XGBoost Classifier. *2021 International Conference on Artificial Intelligence and Big Data Analytics*, 22–27. <https://doi.org/10.1109/ICAIBDA53487.2021.9689735>
- Anbazhagan, T., & Rangaswamy, B. (2025). Early prediction of CKD from time series data using adaptive PSO optimized echo state networks. *Scientific Reports*, 15(1), 6966. <https://doi.org/10.1038/s41598-025-91028-6>
- Badawy, M., Ramadan, N., & Hefny, H. A. (2024). Big data analytics in healthcare: Data sources, tools, challenges, and opportunities. *Journal of Electrical Systems and Information Technology*, 11(1), 63. <https://doi.org/10.1186/s43067-024-00190-w>
- Barus, H. P., Robet, R., & Tarigan, F. A. (2026). Comparison of XGBoost and Naive Bayes Models in Type 2 Diabetes Prediction with RFE Feature Selection. *Sinkron*, 10(1), 352–360. <https://doi.org/10.33395/sinkron.v10i1.15509>
- Benghazouani, S., Nouh, S., & Zakrani, A. (2025). Optimizing breast cancer diagnosis: Harnessing the power of nature-inspired metaheuristics for feature selection with soft voting classifiers. *International Journal of Cognitive Computing in Engineering*, 6, 1–20. <https://doi.org/10.1016/j.ijcce.2024.09.005>
- Boateng, E. B., & Ampofo, A. G. (2023). A glimpse into the future: Modelling global prevalence of hypertension. *BMC Public Health*, 23(1), 1906. <https://doi.org/10.1186/s12889-023-16662-z>
- Dhanka, S., & Maini, S. (2025). A hybridization of XGBoost *Machine Learning* model by Optuna hyperparameter tuning suite for cardiovascular disease classification with significant effect of outliers and heterogeneous training datasets. *International Journal of Cardiology*, 420, 132757. <https://doi.org/10.1016/j.ijcard.2024.132757>

- Dong, T., Oronti, I. B., Sinha, S., Freitas, A., Zhai, B., Chan, J., Fudulu, D. P., Caputo, M., & Angelini, G. D. (2024). Enhancing Cardiovascular Risk Prediction: Development of an Advanced Xgboost Model with Hospital-Level Random Effects. *Bioengineering*, 11(10), 1039. <https://doi.org/10.3390/bioengineering11101039>
- Jinbo, Z., Yufu, L., & Haitao, M. (2025). Handling missing data of using the XGBoost-based multiple imputation by chained equations regression method. *Frontiers in Artificial Intelligence*, 8, 1553220. <https://doi.org/10.3389/frai.2025.1553220>
- Kario, K., Okura, A., Hoshide, S., & Mogi, M. (2024). The WHO Global report 2023 on hypertension warning the emerging hypertension burden in globe and its treatment strategy. *Hypertension Research*, 47(5), 1099–1102. <https://doi.org/10.1038/s41440-024-01622-w>
- Mancia, G., Kreutz, R., Brunström, M., Burnier, M., Grassi, G., Januszewicz, A., Muiesan, M. L., Tsioufis, K., Agabiti-Rosei, E., Algharably, E. A. E., Azizi, M., Benetos, A., Borghi, C., Hitij, J. B., Cifkova, R., Coca, A., Cornelissen, V., Cruickshank, J. K., Cunha, P. G., ... Kjeldsen, S. E. (2023). 2023 ESH Guidelines for the management of arterial hypertension The Task Force for the management of arterial hypertension of the European Society of Hypertension: Endorsed by the International Society of Hypertension (ISH) and the European Renal Association (ERA). *Journal of Hypertension*, 41(12), 1874–2071. <https://doi.org/10.1097/HJH.00000000000003480>
- Maulana As'an Hamid & Egia Rosi Subhiyakto. (2025). Performance Comparison of Random Forest, SVM, and XGBoost Algorithms with SMOTE for Stunting Prediction. *Journal of Applied Informatics and Computing*, 9(4), 1163–1169. <https://doi.org/10.30871/jaic.v9i4.9701>
- Mroz, T., Griffin, M., Cartabuke, R., Laffin, L., Russo-Alvarez, G., Thomas, G., Smedira, N., Meese, T., Shost, M., & Habboub, G. (2024). Predicting hypertension control using *Machine Learning*. *PLOS ONE*, 19(3), e0299932. <https://doi.org/10.1371/journal.pone.0299932>
- Münzel, T., Crea, F., Rajagopalan, S., & Lüscher, T. (2025). Nicotine and the cardiovascular system: Unmasking a global public health threat. *European Heart Journal*, ehaf1010. <https://doi.org/10.1093/eurheartj/ehaf1010>
- Narmilan, A., Gonzalez, F., Salgadoe, A., & Powell, K. (2022). Detection of White Leaf Disease in Sugarcane Using *Machine Learning* Techniques over UAV Multispectral Images. *Drones*, 6(9), 230. <https://doi.org/10.3390/drones6090230>
- Pokharel, Y., Karmacharya, B. M., & Neupane, D. (2022). Hypertension—A Silent Killer Without Global Bounds. *Journal of the American College of Cardiology*, 80(8), 818–820. <https://doi.org/10.1016/j.jacc.2022.05.043>
- Rayadhani, W. A., & Rahardi, M. (2025a). Comparative Analysis of Random Forest, SVM, and Naive Bayes for Cardiovascular Disease Prediction. *Journal of Applied Informatics and Computing*, 9(6), 3234–3243. <https://doi.org/10.30871/jaic.v9i6.11451>
- Rayadhani, W. A., & Rahardi, M. (2025b). Comparative Analysis of Random Forest, SVM, and Naive Bayes for Cardiovascular Disease Prediction. *Journal of Applied Informatics and Computing*, 9(6), 3234–3243. <https://doi.org/10.30871/jaic.v9i6.11451>

- Septian, E., Khaefi, M. R., Athoillah, A., Aisyah, D. N., Hardhantyo, M., Rahman, F. M., & Manikam, L. (2025). Prediction of Personalised Hypertension Using *Machine Learning* in Indonesian Population. *Journal of Medical Systems*, 49(1), 137. <https://doi.org/10.1007/s10916-025-02253-5>
- Surur, M., Santoso, N. A., & Santoso, B. A. (2025). Klasifikasi Keterlambatan Pembayaran SPP Santri Menggunakan Algoritma K-Nearest Neighbor di Pesantren Al Fajar Tegal. *RIGGS: Journal of Artificial Intelligence and Digital Business*, 4(3), 873–883. <https://doi.org/10.31004/riggs.v4i3.2117>
- Takase, M., Hirata, T., Nakaya, N., Kogure, M., Hatanaka, R., Nakaya, K., Chiba, I., Tokioka, S., Nochioka, K., Nakamura, T., Tsuchiya, N., Metoki, H., Satoh, M., Narita, A., Obara, T., Ishikuro, M., Ohseto, H., Takahashi, I., Kobayashi, T., ... the ToMMo investigators. (2025). Associations of family history of hypertension, genetic, and lifestyle risks with incident hypertension. *Hypertension Research*, 48(10), 2606–2617. <https://doi.org/10.1038/s41440-025-02314-9>
- Wang, W., Yan, L., Liu, F., & Li, Y. (2025). Improving Gaussian Naive Bayes classification on imbalanced data through coordinate-based minority feature mining. *PeerJ Computer Science*, 11, e3003. <https://doi.org/10.7717/peerj-cs.3003>
- World Health Organization. (n.d.). *Global report on hypertension: The race against a silent killer*. Retrieved <https://www.who.int/teams/noncommunicable-diseases/hypertension-report>
- World Health Organization. (2023). *Hypertension—Fact Sheet* [Fact Sheet / Webpage]. <https://www.who.int/news-room/fact-sheets/detail/hypertension>
- Yadav, J., Nair, A. M., George, J., & Alapatt, B. P. (2024). Predictive Modelling of Heart Disease: Exploring *Machine Learning* Classification Algorithms. *2024 International Conference on Distributed Computing and Optimization Techniques (ICDCOT)*, 1–7. <https://doi.org/10.1109/ICDCOT61034.2024.10515398>
- Yana Afrina, D. A. S. (2024). *Determinant Kejadian Hipertensi: Studi Literatur Review*. 3. <https://journal.mandiracendikia.com/index.php/JIK-MC/article/view/814/634>
- Zhang, L. (2025). Features extraction based on Naive Bayes algorithm and TF-IDF for news classification. *PLOS One*, 20(7), e0327347. <https://doi.org/10.1371/journal.pone.0327347>
- Zhao, H., Zhang, X., Xu, Y., Gao, L., Ma, Z., Sun, Y., & Wang, W. (2021). Predicting the Risk of Hypertension Based on Several Easy-to-Collect Risk Factors: A *Machine Learning* Method. *Frontiers in Public Health*, 9, 619429. <https://doi.org/10.3389/fpubh.2021.619429>