

Optimasi Decision Tree menggunakan Pendekatan Ensemble Bagging untuk Klasifikasi Klinis Hipertensi

Agustina Diah Kusuma Dewi ^{1,*}, Danang Wahyu Utomo ¹

¹ Program Studi Teknik Informatika, Universitas Dian Nuswantoro, Indonesia

* Correspondence: 111202214122@mhs.dinus.ac.id

Copyright: © 2025 by the authors

Received: 27 September 2025 | Revised: 13 Oktober 2025 | Accepted: 14 November 2025 | Published: 6 Desember 2025

Abstrak

Hipertensi merupakan masalah kesehatan yang umum terjadi dan berisiko tinggi terhadap komplikasi serius, sehingga penting untuk dideteksi secara dini. Penelitian ini merupakan penelitian kuantitatif eksperimental yang bertujuan untuk meningkatkan akurasi model klasifikasi hipertensi dengan mengoptimasi algoritma *Decision Tree* menggunakan metode ensemble *Bagging*. Aspek yang diteliti meliputi pengaruh teknik penyeimbangan data terhadap kinerja model klasifikasi dalam mendeteksi risiko hipertensi. Dataset yang digunakan terdiri atas 2.655 rekam medis pasien layanan kesehatan primer. Tahapan penelitian mencakup pra-pemrosesan data dengan imputasi, normalisasi, dan balancing dengan *SMOTE*, pemodelan *Decision tree* dan *Bagging*, serta evaluasi model menggunakan *5-fold cross-validation*, *Confusion Matrix*, *ROC-AUC*, dan uji paired t-test. Hasil penelitian menunjukkan bahwa penerapan *Bagging* meningkatkan performa *Decision tree* secara signifikan ($p < 0,05$) dengan akurasi 96,88%, presisi 95,56%, recall 98,32%, f1-score 96,93%, dan AUC 0,996. Peningkatan ini menunjukkan bahwa metode *Bagging* mampu menekan variansi model dan menghasilkan prediksi yang lebih stabil pada data klinis yang tidak seimbang. Temuan ini berkontribusi pada pengembangan model klasifikasi berbasis data klinis nyata yang juga membuka peluang pengembangan *Clinical Decision Support System* (CDSS) berbasis *machine learning* untuk deteksi dini hipertensi di fasilitas kesehatan primer.

Kata kunci: *bagging*; *decision tree*; hipertensi; klasifikasi

Abstract

Hypertension is a common health problem with a high risk of serious complications, so it is important to detect it early. This study is an experimental quantitative study that aims to improve the accuracy of the hypertension classification model by optimizing the Decision tree algorithm using the Bagging ensemble method. The aspects studied include the effect of data balancing techniques on the performance of classification models in detecting the risk of hypertension. The dataset used consisted of 2,655 medical records of primary health care patients. The research stages included data pre-processing with imputation, normalization, and balancing with SMOTE, Decision Tree and Bagging modeling, and model evaluation using 5-fold cross-validation, Confusion Matrix, ROC-AUC, and paired t-test. The results showed that the application of Bagging significantly improved the performance of the Decision Tree ($p < 0.05$) with an accuracy of 96.88%, precision of 95.56%, recall of 98.32%, F1-score of 96.93%, and AUC of 0.996. This improvement indicates that the Bagging method is able to reduce model variance and produce more stable predictions on imbalanced clinical data. These findings contribute to the development of a classification model based on real clinical data, which also opens up opportunities for the development of machine learning-based Clinical Decision Support System (CDSS) for early detection of hypertension in primary health facilities.

Keywords: *bagging*; *classification*; *decision tree*; *hypertension*



PENDAHULUAN

Hipertensi (tekanan darah tinggi) merupakan kondisi ketika tekanan darah sistolik (>140 mmHg) dan diastolik (>90 mmHg) melebihi batas normal (Muslim et al., 2025). Kondisi ini disebabkan oleh gangguan pembuluh darah yang menghambat suplai oksigen (Muflih & Halimizami, 2021). Berdasarkan data WHO, sekitar 1,13 miliar orang di dunia menderita hipertensi dengan prevalensi 22%, dan Asia Tenggara menempati peringkat tertinggi (Norhasanah et al., 2024). Risiko hipertensi dipengaruhi oleh usia, jenis kelamin, genetika, obesitas, kebiasaan merokok, kurang aktivitas fisik, dan stres (Aditya et al., 2024).

Tingginya angka prevalensi dan dampak komplikasi serius, deteksi dini hipertensi menjadi langkah penting dalam upaya pencegahan dan pengendalian penyakit ini. Teknologi *machine learning* kini banyak dimanfaatkan untuk membantu tenaga medis memprediksi risiko secara akurat (Kim et al., 2022). Berbagai algoritma, seperti *Decision tree*, *Bagging*, dan model *machine learning* lainnya digunakan karena mampu mengolah data pasien yang bervariasi (Ji et al., 2022). Pada kasus hipertensi, klasifikasi berperan penting untuk memprediksi tingkat risiko berdasarkan faktor kesehatan yang terkait (Kurniawan et al., 2023).

Algoritma *Decision tree* banyak digunakan dalam klasifikasi medis, namun algoritma ini memiliki kelemahan yaitu cenderung mengalami overfitting dan variance tinggi saat menghadapi data medis yang bervariasi (Iskhak et al., 2025). Untuk mengatasi hal ini, pendekatan *Ensemble Learning* seperti *Bagging* diterapkan dengan menggabungkan beberapa model seperti *Decision tree* agar hasil prediksi lebih stabil dan akurat (Xi et al., 2022). Teknik *ensemble learning* lainnya seperti *Boosting*, *Stacking*, dan *Voting* juga terbukti meningkatkan performa model (Masacgi & Rohman, 2023). *Ensemble learning* menggabungkan beberapa model untuk meningkatkan akurasi dan kestabilan prediksi (Ananda et al., 2024), termasuk dalam prediksi penyakit kronis seperti diabetes, kanker dan jantung (Anugrah et al., 2024; Ariawan, 2024; Pramudita et al., 2024), sehingga pendekatan ini relevan diterapkan juga untuk kasus hipertensi.

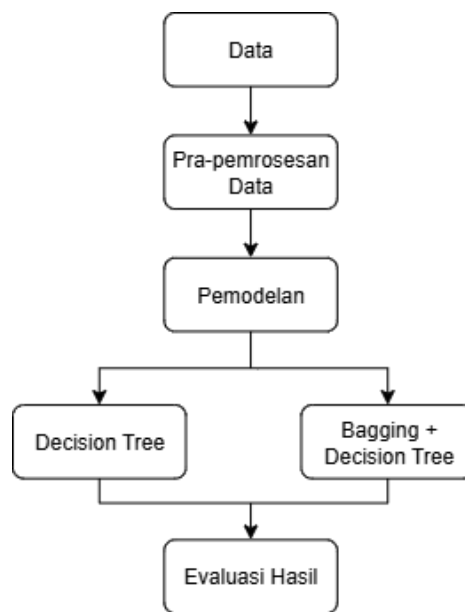
Penerapan algoritma tunggal pada penelitian sebelumnya menunjukkan akurasi rendah pada kasus hipertensi, menggunakan *Decision Tree* masing-masing 43,85% dan 50%, akibat keterbatasan data dan ketidakseimbangan kelas (Rafidah et al., 2024). Kinerja *Decision tree* cenderung tidak stabil karena variance tinggi dan overfitting (Iskhak et al., 2025), sehingga memerlukan pendekatan *ensemble* seperti *Bagging* untuk meningkatkan performa model (Anugrah et al., 2024; Haque et al., 2025). Oleh karena itu, penelitian ini mengusulkan model klasifikasi hipertensi dengan menerapkan teknik penyeimbangan data serta optimasi *Ensemble Bagging* pada *Decision tree*. Meskipun efektivitas *Bagging* telah dibuktikan pada berbagai penelitian, seperti klasifikasi kanker payudara (Pramudita et al., 2024) menunjukkan bahwa *Decision tree* sebelum dioptimasi memperoleh akurasi 94,1%, dan meningkat menjadi 94,7% setelah *Bagging*.

Namun demikian, sebagian besar penelitian sebelumnya menggunakan dataset publik atau dataset klinis yang tidak merepresentasikan data pasien di pelayanan kesehatan dasar. Penerapan *Bagging* untuk mengoptimalkan *Decision tree* pada penelitian kami menggunakan dataset klinis nyata yang penerapannya untuk klasifikasi hipertensi masih sangat terbatas. Belum banyak penelitian yang mengembangkan model klasifikasi hipertensi berbasis data klinis nyata (primer) dari fasilitas kesehatan lokal, sehingga efektivitas model pada konteks nyata masih belum banyak dibuktikan.

Penelitian ini bertujuan untuk mengembangkan model klasifikasi hipertensi menggunakan data klinis nyata dari salah satu Puskesmas di Kota Semarang. Model ini menerapkan teknik penyeimbangan data (SMOTE) serta optimasi *Decision tree* menggunakan *Bagging*. Fokus utama penelitian adalah meningkatkan akurasi, *presisi*, *recall*, dan *f1-score* dan stabilitas prediksi model, sekaligus mengevaluasi efektivitas pendekatan ensemble dalam mengurangi overfitting pada data medis yang bersifat tidak seimbang.

METODE

Penelitian ini menggunakan 2.655 data klinis nyata rekam medis elektronik (RME) pasien hipertensi dari salah satu Puskesmas di Kota Semarang, dikumpulkan antara Juli hingga September 2025. Data dipilih karena merepresentasikan kondisi pasien nyata dengan keragaman sosial dan demografi yang sesuai kondisi populasi Indonesia. Komposisi data menunjukkan 72% perempuan dan 28% laki-laki, dengan rentang usia 15–92 tahun. Distribusi kelas tidak seimbang, yaitu 78% tidak hipertensi dan 22% hipertensi. Penelitian ini telah memperoleh izin etik dari Dinas Kesehatan Kota Semarang.



Gambar 1. Alur Penelitian

Gambar 1 menampilkan alur penelitian mulai dari tahap pra-pemrosesan data, penerapan algoritma *Decision tree* dan optimasi dengan *Bagging*, hingga evaluasi menggunakan *Confusion Matrix* dan *ROC*. Pada tahap pra-pemrosesan dilakukan penghapusan atribut yang tidak relevan, pembersihan data, dan data biner dikonversi ke numerik. *Missing value* diimputasi menggunakan modus untuk data kategorikal dan median untuk data numerik, karena sebagian besar nilai hilang akibat pasien tidak menjalani pemeriksaan tertentu (Babiloni-Chust et al., 2022). Ketidakseimbangan kelas diatasi dengan *SMOTE* ($k_neighbors=5$, $random_state=42$). Setelah *class balancing*, data dibagi menjadi 80% data latih dan 20% data uji menggunakan *stratified split*. Normalisasi dilakukan menggunakan *StandardScaler* agar setiap fitur memiliki rata-rata 0 dan deviasi standar 1.

Decision tree merupakan algoritma klasifikasi yang menentukan atribut terbaik berdasarkan kondisi fitur (Anand et al., 2025). Pada penelitian ini, model menggunakan pengaturan default Scikit-learn 1.6.1 dengan kriteria Gini Impurity, $max_depth=None$, $min_samples_split=2$, dan $random_state=42$. Bagging merupakan teknik ensemble learning, dimana penerapan Bagging dengan $n_estimators=100$ digunakan untuk meningkatkan akurasi dan mengurangi overfitting. Validasi dilakukan menggunakan 5-fold cross-validation, dan peningkatan performa diuji dengan paired t-test ($p < 0,05$).

Evaluasi model menggunakan *Confusion Matrix* serta metrik *ROC-AUC*. Selain itu, digunakan *Balanced Accuracy* untuk memberikan penilaian yang lebih adil pada data yang tidak seimbang. Seluruh eksperimen dijalankan menggunakan Python 3.12.12 di Google Colab dengan library utama Pandas 2.2.2, NumPy 2.0.2, Scikit-learn 1.6.1, Matplotlib 3.10, dan Seaborn 0.13.2.

HASIL DAN PEMBAHASAN

Hasil

Hasil menunjukkan jika *Bagging* meningkatkan performa *Decision tree* tunggal secara signifikan setelah data dibersihkan, dimapping, dan diseimbangkan. Sebelum pra-pemrosesan, data memiliki kombinasi atribut numerik dan biner, jadi untuk melanjutkan ke proses-proses selanjutnya akan di lakukan terlebih dahulu proses *mapping* untuk mengubah data yang berbentuk biner menjadi data numerik.

Tabel 1. Data Asli sebelum pra-pemrosesan data

n o	jenis _kela _min	um ur	bb	tb	bmi	status_ bmi	td_s istol ik	td_di astoli k	konsums i_obat_h t	gula _dar _ah	beresi ko_hip ertensi
1	L	55	50	160	19,5	Normal	90	70	Tidak	100	Tidak
2	P	80	36	145	17,5	Kurus	154	66	Ya	214	Ya

Berdasarkan data asli yang ada pada tabel 1, terdapat kolom yang tidak diperlukan, maka dilakukan juga proses *drop columns* pada kolom *no* dan *status_bmi*, sebelum dilakukan *mapping*. Pada tabel 1 terlihat pada class *jenis_kelamin*, *konsumsi_obat_ht* dan *beresiko_hipertensi* terdapat data binary. Maka data pada variabel kategorikal diubah menjadi bentuk biner, di antaranya *jenis_kelamin* (0 = Laki-laki, 1 = Perempuan), *konsumsi_obat_ht* (0 = Tidak, 1 = Ya), dan *beresiko_hipertensi* (0 = Tidak, 1 = Ya). Proses ini memungkinkan data diolah oleh algoritma *machine learning*.

Tabel 2. Data setelah *mapping* dan *drop columns*

jenis_kel amin	um ur	bb	tb	bmi	td_sis tolik	td_dias tolik	konsumsi _obat_ht	gula_ darah	beresiko_h ipertensi
0	55	50	160	19,5	90	70	0	100	0
1	80	36	145	17,5	154	66	1	214	1

Data pada Tabel 2 telah melalui proses pembersihan dengan menghapus atribut tidak diperlukan dan mengubah variabel kategorikal menjadi biner, sehingga seluruh data berbentuk numerik. Selanjutnya dilakukan pengecekan dan imputasi missing value menggunakan median untuk data numerik dan modus untuk data kategorikal agar dataset lengkap dan siap digunakan dalam pelatihan model.

Tabel 3. Cek *missing value*

Kolom	Nilai
jenis kelamin	1
umur	1
bb	1
tb	1
bmi	1
td_sistolik	1
td_distolik	1
konsumsi obat ht	7
gula darah	1058
beresiko_hipertensi	1

Tabel 3 terlihat pada kolom umur, bb, tb, bmi, td_sistolik, td_diastolik, gula_darah dilakukan terdapat *missing value* yang akan diimputasi menggunakan median. Pada kolom jenis_kelamin, konsumsi_obat_ht, dan beresiko_hipertensi dilakukan imputasi modus. Dengan pengisian seperti ini, data jadi lebih utuh dan tidak menimbulkan kesalahan saat proses analisis nanti.

Tabel 4. Setelah imputasi median & modus

Kolom	Nilai
jenis_kelamin	0
umur	0
bb	0
tb	0
bmi	0
td_sistolik	0
td_distolik	0
konsumsi_obat_ht	0
gula_darah	0
beresiko_hipertensi	0

Hasil pada tabel 4 menunjukkan hasil dari tahap imputasi, sudah tidak ada lagi *missing value* pada seluruh atribut. Dataset akhir berisi 2.655 pasien berusia 15–92 tahun (mean 47,7; median 50; SD 17,5) dengan rasio 72% laki-laki dan 28% perempuan. Rata-rata tekanan darah 133/76 mmHg, menunjukkan distribusi data yang representatif terhadap kondisi klinis hipertensi. Tahap selanjutnya dilakukan *class balancing* karena distribusi data tidak seimbang. Dari 2.655 data hipertensi, kategori Ya (1) sebagai kelas minoritas diseimbangkan dengan Tidak (0) menggunakan teknik oversampling *SMOTE*.

Tabel 5. Sebelum *class balancing*

	Tidak (0)	Ya (1)
Jumlah Data	2080	575

Sementara itu hasil pada tabel 5 menunjukkan bahwa terdapat ketidakseimbangan jumlah data antara kategori beresiko hipertensi Tidak(0) dan Ya(1). Proses *class balancing* dilakukan untuk menyeimbangkan jumlah data tiap kelas agar model tidak bias dan mengurangi risiko *overfitting*. Setelah penyeimbangan, model diharapkan menghasilkan prediksi yang lebih adil dan akurat.

Tabel 6. Setelah *class balancing*

	Tidak (0)	Ya (1)
Jumlah Data	2080	2080

Setelah dilakukan proses balancing menggunakan *SMOTE*, hasilnya bisa dilihat di tabel 6 menunjukkan proporsi data kelas berisiko hipertensi (1) dan tidak berisiko (0) menjadi seimbang (2080:2080). Dengan begitu, data sudah siap digunakan untuk pelatihan model tanpa perlu khawatir soal ketidakseimbangan kelas. Tabel 7 merupakan tahapan normalisasi agar semua fitur memiliki skala yang sama, supaya model bisa menghitung dengan lebih akurat dan hasil prediksi jadi lebih baik. Normalisasi dilakukan dengan metode *StandardScaler*. Metode ini mengubah semua nilai menjadi standar, artinya setiap fitur memiliki rata-rata nol dan deviasi satu. Hal ini penting karena beberapa fitur seperti berat badan dan tekanan darah punya

satuan yang berbeda, dan kalau tidak dinormalisasi bisa mempengaruhi hasil perhitungan model. Setelah proses ini, data jadi lebih seragam dan mudah diolah oleh algoritma seperti yang ditunjukkan pada tabel 8, dimana semua nilai sudah dalam bentuk yang sudah ditransformasi.

Tabel 7. Data sebelum *StandardScaler*

Atribut	Nilai
Jenis kelamin	1
umur	50,0
bb	60,0
tb	155,0
bmi	24,3
td sistolik	124,0
td diastolik	77,0
konsumsi obat ht	0
gula darah	113,0

Tabel 8. Data sesudah *standardscaler*

Atribut	Nilai
jenis_kelamin	0,611799
umur	-0,047376
bb	-0,056034
tb	-0,099439
bmi	-0,061221
td sistolik	-0,072681
td diastolik	-0,216474
konsumsi obat ht	-0,814862
gula darah	-0,186952

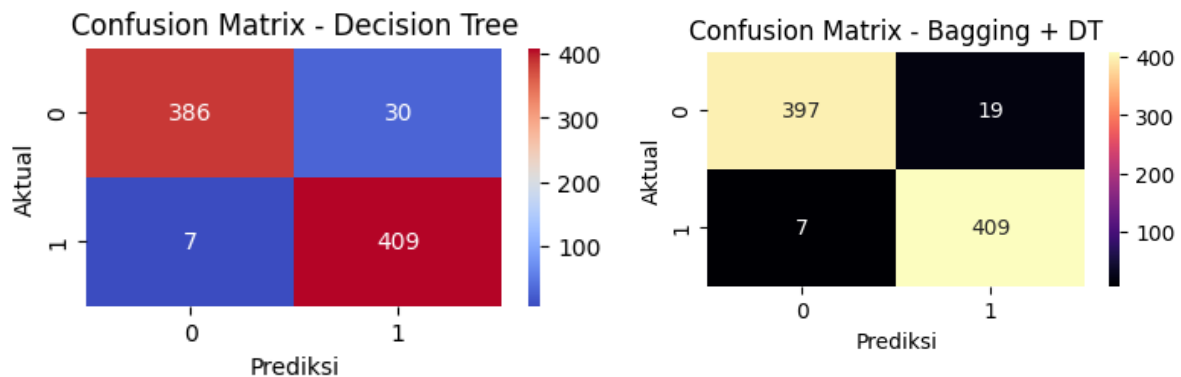
Tabel 9. Hasil Perbandingan Kinerja Model *Decision tree* dan *Bagging + Decision tree*

Algoritma	Akurasi	Presisi	Recall	F1-score	Balanced Accuracy
<i>Decision tree</i>	95,55%	93,17%	98,32%	95,68%	95,55%
<i>Bagging + Decision tree</i>	96,88%	95,56%	98,32%	96,93%	96,88%

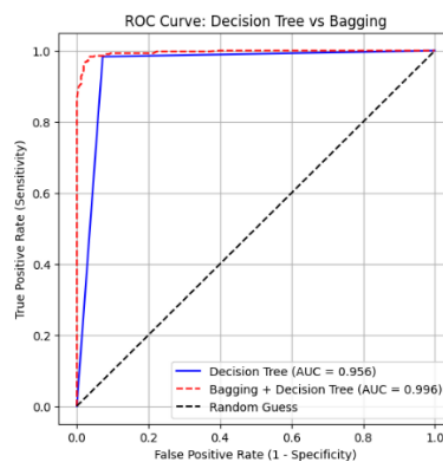
Hasil ditampilkan pada tabel 9 menunjukkan bahwa model *Decision tree* tunggal memiliki akurasi 95,55%, sedangkan model gabungan *Bagging + Decision tree* meningkat menjadi 96,88%. Artinya, metode gabungan ini membuat prediksi jadi lebih tepat dan lebih konsisten. Nilai *balanced accuracy* meningkat dari 0,9555 menjadi 0,9688, menunjukkan kinerja model yang semakin baik. Evaluasi juga dilakukan menggunakan *Confusion Matrix* untuk melihat seberapa sering model benar dan salah dalam mengklasifikasikan data.

Sementara itu, hasil pada gambar 2 menunjukkan bahwa metode *Bagging* meningkatkan akurasi model dengan menurunkan false positive dari 30 ke 19 dan false negative dari 11 ke 7. Nilai *Balanced Accuracy* naik dari 0,9555 menjadi 0,9688, dan hasil ROC-AUC menegaskan efektivitas *Bagging + Decision tree* dalam membedakan pasien berisiko hipertensi dan tidak. Selanjutnya hasil pada gambar 3 menunjukkan bahwa *Bagging + Decision tree* ($AUC = 0,996$) memiliki kemampuan klasifikasi lebih baik dibanding *Decision tree* tunggal ($AUC = 0,956$). Pada ambang 0,5, model mencapai TPR 0,983 dan FPR 0,045. Nilai AUC 0.996 menunjukkan

kemampuan diskriminatif model yang hampir sempurna dalam membedakan pasien hipertensi dan tidak, menegaskan reliabilitas model *Bagging* + DT untuk aplikasi klinis.



Gambar 2. *Confusion matrix*



Gambar 3. *Kurva roc*

Tabel 10. Hasil 5-fold cross-validation dan uji *paired t-test*

Model	Mean Akurasi	Std Deviasi	t-statistic	p-value	keterangan
<i>Decision tree</i>	0,9546	0,0096	-	-	Model dasar (<i>baseline</i>) tanpa ensemble
<i>Bagging + Decision tree</i>	0,9654	0,0061	4,8809	0,0082	Peningkatan akurasi terbukti signifikan secara statistik ($p < 0.05$).

Hasil pada tabel 10 menunjukkan bahwa model *Bagging* dan *Decision tree* memiliki rata-rata akurasi yang lebih tinggi (0.9654) dibandingkan model dasar *Decision tree* (0.9546). Ini menunjukkan peningkatan kinerja setelah penerapan *Bagging*. Hal ini membuktikan bahwa peningkatan akurasi yang diamati dari model *Decision Tree* ke *Bagging + Decision Tree* adalah signifikan secara statistik. Artinya, peningkatan tersebut bukan terjadi secara kebetulan.

Pembahasan

Berdasarkan hasil pada penelitian ini, menunjukkan bahwa optimasi *Decision tree* menggunakan *Bagging* memberikan peningkatan performa yang signifikan secara statistik ($p < 0,05$), mengubah akurasi dari 95,55% menjadi 96,88% dan presisi dari 93,17% menjadi 95,56%, dengan recall yang tetap tinggi di 98,32%. Dalam penelitian ini kami juga

mendapatkan hasil dari *Balanced Accuracy Decision tree* 0,9555 dan *Bagging* 0,9688. Hal ini menunjukkan kemampuan prediksi yang lebih pas dengan kondisi data sebenarnya. Peningkatan ini terjadi karena *Bagging* mampu menekan variansi pada *Decision tree* yang cenderung mengalami overfitting, sehingga prediksi menjadi lebih stabil, terutama pada data klinis yang bersifat tidak seimbang (Agustriya et al., 2024).

Hasil temuan kami jika dibandingkan dengan penelitian sebelumnya, hasil ini konsisten dengan penelitian Hwang et al. (2024) dan Khan et al. (2023) yang menyatakan bahwa metode ensemble learning efektif meningkatkan akurasi pada data medis imbalanced. Namun, penelitian ini menunjukkan performa lebih tinggi dengan peningkatan akurasi sekitar 1,39%, melampaui rata-rata kenaikan 0,5–1,0% yang dilakukan pada penelitian sebelumnya, dimana hasil akurasi 96,88% ini lebih tinggi dibanding penelitian pada klasifikasi hipertensi menggunakan *decision tree* tunggal sebesar 43,85% (Rafidah et al., 2024). Sementara itu Pramudita et al. (2024) pada klasifikasi kanker payudara menghasilkan peningkatan akurasi dari 94,1% menjadi 94,7% setelah penerapan *Bagging*, dan Anugrah et al. (2024) mendapat peningkatan signifikan pada klasifikasi diabetes menggunakan teknik yang sama.

Penelitian ini sejalan dengan penelitian Sifat & Kibria (2024) yang menyimpulkan bahwa metode ensemble seperti *Bagging* dan *Boosting* mampu mengungguli model tunggal, tetapi juga menunjukkan kontribusi baru pada konteks data klinis nyata yang sebelumnya jarang diuji. Peningkatan akurasi menunjukkan bahwa *Bagging* lebih efektif untuk data medis yang imbalanced dan mengandung noise. Hasil ini membuktikan bahwa model ensemble mampu menghasilkan klasifikasi yang lebih akurat dan reliabel dalam mendukung sistem pengambilan keputusan medis.

Selain peningkatan akurasi, model yang dikembangkan juga menunjukkan kinerja sensitivitas (*recall*) yang sangat tinggi, yaitu 98,32%, menandakan kemampuan model dalam mendeteksi hampir semua pasien hipertensi. Hal ini sangat penting dalam konteks deteksi dini dan pencegahan komplikasi hipertensi. Chai et al. (2022) menunjukkan bahwa model ensemble cenderung lebih konsisten dan sensitif dalam mengklasifikasikan kasus risiko hipertensi dibanding algoritma tunggal. Peningkatan presisi dari 93,17% menjadi 95,56% menunjukkan berkurangnya kesalahan klasifikasi: *false positive* turun dari 30 ke 19 dan *false negative* dari 11 ke 7. Dengan demikian, risiko over-diagnosis dan under-diagnosis dapat diminimalkan. Nilai AUC yang meningkat dari 0,956 menjadi 0,996 juga menunjukkan kemampuan diskriminatif model yang sangat baik dalam membedakan pasien berisiko hipertensi.

Penelitian ini menjadi langkah awal dalam penerapan algoritma *ensemble learning* pada dataset berbasis data klinis dengan data klinis nyata yang sebelumnya belum pernah diklasifikasikan. Peningkatan performa model telah divalidasi melalui uji statistik yang signifikan ($p < 0,05$). Meski demikian, penelitian ini masih terbatas pada satu teknik *ensemble*, yakni *Bagging*, serta menggunakan data dari satu wilayah dan periode waktu tertentu. Oleh karena itu, studi lanjutan disarankan untuk mengeksplorasi metode *ensemble* lain seperti *Random Forest* atau *XGBoost* dan menguji model pada data dari wilayah berbeda maupun data real-time berbasis RME untuk meningkatkan stabilitas serta generalisasi model. Hasil penelitian ini juga membuka peluang pengembangan CDSS berbasis *machine learning* yang lebih akurat dan objektif bagi layanan kesehatan primer.

SIMPULAN

Penelitian ini menunjukkan bahwa penggunaan *bagging* pada *decision tree* mampu meningkatkan kinerja klasifikasi risiko hipertensi, dengan akurasi 96,88%, presisi 95,56%, *recall* 98,32%, *f1-score* 96,93%, dan AUC 0,996, menjadikan model lebih stabil dan andal. Nilai *recall* yang tinggi mendukung deteksi hipertensi secara dini, sedangkan presisi yang meningkat membantu mengurangi kemungkinan salah diagnosis dan intervensi medis yang tidak perlu. Dengan hasil ini, dapat menjadi hal baru dalam penerapan *ensemble learning*

bagging pada data medis nyata menunjukkan bahwa penelitian ini mengisi kekosongan literatur penerapan *ensemble bagging* pada klasifikasi hipertensi dengan data klinis nyata yang sebelumnya belum banyak dilakukan dan dapat berpotensi integrasi model ke CDSS untuk mendukung skrining di layanan kesehatan primer. Penelitian ini dibatasi oleh penggunaan dataset dari satu fasilitas kesehatan dan belum dibandingkan dengan metode *ensemble* lain, sehingga penelitian lanjutan disarankan untuk menguji model pada dataset yang lebih luas dan membandingkan berbagai metode *ensemble* serta implementasinya dalam CDSS secara nyata.

REFERENSI

- Aditya, M. F. R., Lutvi, N., & Indahyanti, U. (2024). Prediksi Penyakit Hipertensi Menggunakan Metode Decison Tree dan Random Forest. *Jurnal Ilmiah Komputasi*, 23(1), 9–16. <https://doi.org/10.32409/jikstik.23.1.3503>
- Agustriya, M., & Ula, M. (2024). Analisis Kinerja Algoritma Klasifikasi Naïve Bayes Menggunakan Genetic Algorithm dan Bagging untuk Data Publik Risiko Transaksi Kartu Kredit. *JUSTIN (Jurnal Sistem dan Teknologi Informasi)*, 12(3), 584-591. <https://doi.org/10.26418/justin.v12i3.80136>
- Anand, V., Khajuria, A., Pachauri, R. K., & Gupta, V. (2025). Optimized machine learning based comparative analysis of predictive models for classification of kidney tumors. *Scientific Reports*, 15(1), 1–13. <https://doi.org/10.1038/s41598-025-15414-w>
- Ananda, I. K., Fanani, A. Z., Setiawan, D., & Wicaksono, D. F. (2024). Penerapan Random Oversampling dan Algoritma Boosting untuk Memprediksi Kualitas Buah Jeruk. *Edumatic: Jurnal Pendidikan Informatika*, 8(1), 282–289. <https://doi.org/10.29408/edumatic.v8i1.25836>
- Anugrah, M. I., Zeniarja, J., & Setiawan, D. S. (2024). Peningkatan Performa Model Hard Voting Classifier dengan Teknik Oversampling ADASYN pada Penyakit Diabetes. *Edumatic: Jurnal Pendidikan Informatika*, 8(1), 290–299. <https://doi.org/10.29408/edumatic.v8i1.25838>
- Ariawan, A. (2024). Optimasi Prediksi Gagal Jantung dengan Teknik Ensemble Bagging Pada Neural Network. *Kesatria: Jurnal Penerapan Sistem Informasi (Komputer dan Manajemen)*, 5(3), 1000-1009.
- Babiloni-Chust, I., Dos Santos, R. S., Medina-Gali, R. M., Perez-Serna, A. A., Encinar, J. A., Martinez-Pinna, J., ... & Nadal, A. (2022). G protein-coupled estrogen receptor activation by bisphenol-A disrupts the protection from apoptosis conferred by the estrogen receptors ER α and ER β in pancreatic beta cells. *Environment international*, 164, 107250. <https://doi.org/10.1016/j.envint.2022.107250>
- Chai, S. S., Goh, K. L., Cheah, W. L., Chang, Y. H. R., & Ng, G. W. (2022). Hypertension prediction in adolescents using anthropometric measurements: do machine learning models perform equally well?. *Applied Sciences*, 12(3), 1600. <https://doi.org/10.3390/app12031600>
- Haque, R., Wang, C., & Pala, N. (2025). An Ensemble-Based AI Approach for Continuous Blood Pressure Estimation in Health Monitoring Applications. *Sensors*, 25(15), 4574. <https://doi.org/10.3390/s25154574>
- Hwang, S. H., Lee, H., Lee, J. H., Lee, M., Koyanagi, A., Smith, L., ... & Lee, J. (2024). Machine Learning–Based Prediction for Incident Hypertension Based on Regular Health Checkup Data: Derivation and Validation in 2 Independent Nationwide Cohorts in South Korea and Japan. *Journal of Medical Internet Research*, 26, e52794. <https://doi.org/10.3390/app12031600>
- Ji, W., Zhang, Y., Cheng, Y., Wang, Y., & Zhou, Y. (2022). Development and validation of prediction models for hypertension risks: A cross-sectional study based on 4,287,407 participants. *Frontiers in Cardiovascular Medicine*, 9, 928948.

- <https://doi.org/10.3389/fcvm.2022.928948>
- Iskhak, Z. M., Darmanto, E., & Muzid, S. (2025). Integrasi Sistem Pakar Forward Chaining dan Decision Tree untuk Deteksi Hama berbasis WhatsApp. *Edumatic: Jurnal Pendidikan Informatika*, 9(2), 560-569. <https://doi.org/10.29408/edumatic.v9i2.31189>
- Khan, A. A., Chaudhari, O., & Chandra, R. (2024). A review of ensemble learning and data augmentation models for class imbalanced problems: Combination, implementation and evaluation. *Expert Systems with Applications*, 244, 122778. <https://doi.org/10.1016/j.eswa.2023.122778>
- Kim, H., Hwang, S., Lee, S., & Kim, Y. (2022). Classification and prediction on hypertension with blood pressure determinants in a deep learning algorithm. *International journal of environmental research and public health*, 19(22), 15301. <https://doi.org/10.3390/ijerph192215301>
- Kurniawan, R., Utomo, B., Siregar, K. N., Ramli, K., Besral, Suhatri, R. J., & Pratiwi, O. A. (2023). Hypertension prediction using machine learning algorithm among Indonesian adults. *IAES International Journal of Artificial Intelligence*, 12(2), 776-784. <https://doi.org/10.11591/ijai.v12.i2.pp776-784>
- Masacgi, G. N., & Rohman, M. S. (2023). Optimasi Model Algoritma Klasifikasi menggunakan Metode Bagging pada Stunting Balita. *Edumatic: Jurnal Pendidikan Informatika*, 7(2), 455-464. <https://doi.org/10.29408/edumatic.v7i2.23812>
- Muflih, M., & Halimizami, H. (2021). Hubungan Tingkat Pengetahuan Dan Gaya Hidup Dengan Upaya Pencegahan Stroke Pada Penderita Hipertensi Di Puskesmas Desa Binjai Medan. *Indonesian Trust Health Journal*, 4(2), 463-471. <https://doi.org/10.37104/ithj.v4i2.79>
- Muslim, Z., Saktia, L., & Pudiarifanti, N. (2025). The level of compliance in the use of drugs based on the antihypertensive and the patient's classification of hypertension: A cross-sectional study. *Riset Informasi Kesehatan*, 14(2), 240. <https://doi.org/10.30644/rik.v14i2.1001>
- Norhasanah, N., Susanti, N., & Normila, N. (2024). The Pengaruh edukasi gizi seimbang dalam mencegah obesitas terhadap pengetahuan dan sikap remaja. *Jurnal Ilmu Kesehatan Bhakti Husada: Health Sciences Journal*, 15(02), 461-472. <https://doi.org/10.34305/jikbh.v15i02.1170>
- Pramudita, R., Muis, S., Safitri, N., & Shafirawati, F. (2024). Optimasi Algoritma Machine Learning Menggunakan Teknik Bagging Pada Klasifikasi Diagnosis Kanker Payudara. *Tematik*, 11(1), 128-134. <https://doi.org/10.38204/tematik.v11i1.1928>
- Rafidah, A. R., Nurmallasari, M., Hosizah, H., & Larasati, D. N. (2024). Penerapan Decision Tree dan Neural Network untuk Prediksi Severity Level Pada Kasus Hipertensi di RSUD Khidmat Sehat Afiat (KISA) Depok. *J-REMI: Jurnal Rekam Medik Dan Informasi Kesehatan*, 6(1), 9-18. <https://doi.org/10.25047/j-remi.v6i1.4963>
- Sifat, I. K., & Kibria, M. K. (2024). Optimizing hypertension prediction using ensemble learning approaches. *PLoS One*, 19(12), e0315865. <https://doi.org/10.1371/journal.pone.0315865>
- Xi, Y., Wang, H., & Sun, N. (2022). Machine learning outperforms traditional logistic regression and offers new possibilities for cardiovascular risk prediction: A study involving 143,043 Chinese patients with hypertension. *Frontiers in cardiovascular medicine*, 9, 1025705. <https://doi.org/10.3389/fcvm.2022.1025705>