

ANALISIS PENGARUH JUMLAH POHON PADA KINERJA RANDOM FOREST DALAM KLASIFIKASI OBESITAS

Heri Akma Senjaya¹, Rudi Kurniawan², Bani Nurhakim³, Umi Hayati⁴.

Program Studi Teknik Informatika¹²⁴
Program Studi Manajemen Informatika³

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
heriakmasenjyaiu@gmail.com

(*) Corresponding Author : heriakmasenjyaiu@gmail.com
Published : 15 Juli 2026

Abstract—This study aims to analyze and optimize the performance of the random forest algorithm in classifying obesity status based on lifestyle, anthropometric factors, and individual habits. The main problem addressed in this research is the absence of an optimal number of decision trees ($n_estimators$) capable of producing the highest accuracy in multiclass obesity classification, as well as the need to identify the most influential features that determine the prediction outcomes. To solve this issue, a computational experimental approach was implemented by testing several variations of $n_estimators$, namely 10, 50, 100, 200, 300, and 500 trees. The research procedure begins with data preprocessing, including normalization of numerical attributes and one-hot encoding for categorical variables, followed by splitting the dataset into training and testing sets, training the model, and evaluating its performance. The evaluation phase employs accuracy, precision, recall, macro f1-score, and confusion matrix to measure the classification quality across different obesity classes. Additionally, a feature importance analysis was conducted to identify the most significant variables, such as body weight, height, age, vegetable consumption habits, physical activity, and daily water intake. The results demonstrate that the number of trees significantly affects the model's performance, with $n_estimators$ of 300 achieving the highest accuracy of 94%, while further increasing the number to 500 provides no meaningful improvement and even shows a slight decrease in effectiveness. The analysis also indicates that body weight, height, and age are the most dominant factors influencing obesity classification. These findings highlight the importance of parameter tuning in enhancing the predictive quality and computational efficiency of machine learning models. This research is also aligned with the sustainable development goals (sdgs), particularly goal 3 on ensuring healthy lives and promoting well-being, as it supports the development of data-driven decision-support systems for monitoring obesity risks. Overall, this study contributes to producing an accurate, efficient, and interpretable obesity classification model while offering valuable insights into key features that can serve as indicators in public health analysis.

Keywords : Random Forest, $n_estimators$, Obesity Classification, Feature Importance, Machine Learning.

Abstrak—Penelitian ini dilakukan untuk menganalisis dan mengoptimalkan performa algoritma *random forest* dalam mengklasifikasikan status obesitas berdasarkan data gaya hidup, antropometri, dan kebiasaan individu. Permasalahan utama yang melatarbelakangi penelitian ini adalah belum adanya nilai $n_estimators$ yang optimal untuk menghasilkan akurasi terbaik pada klasifikasi multikelas, serta perlunya identifikasi fitur-fitur yang paling berkontribusi dalam menentukan hasil prediksi. Untuk memberikan solusi terhadap masalah tersebut, penelitian ini menerapkan pendekatan eksperimen komputasional dengan menguji beberapa variasi jumlah pohon keputusan, yaitu 10, 50, 100, 200, 300, dan 500. Proses penelitian dimulai dari pengumpulan dan pengolahan data melalui tahapan pre-processing, seperti normalisasi fitur numerik dan *one-hot encoding* untuk fitur kategorikal, kemudian dilanjutkan dengan pembagian data latih dan data uji, pelatihan model, serta evaluasi kinerja model. Evaluasi dilakukan menggunakan metrik akurasi, *precision*, *recall*, *macro f1-score*, dan *confusion matrix* untuk menilai tingkat keberhasilan klasifikasi pada setiap kategori obesitas. Analisis *feature importance* juga digunakan untuk mengetahui variabel-variabel yang memiliki pengaruh paling besar, seperti berat badan, tinggi badan, usia, kebiasaan makan, aktivitas fisik, serta konsumsi air harian. Hasil penelitian menunjukkan bahwa jumlah pohon yang digunakan sangat memengaruhi performa model, di mana nilai $n_estimators$ 300 menghasilkan akurasi tertinggi sebesar 94%, sementara peningkatan hingga 500 pohon tidak memberikan peningkatan yang signifikan dan bahkan

menunjukkan sedikit penurunan efektivitas. Selain itu, hasil analisis menunjukkan bahwa fitur berat badan, tinggi badan, dan usia merupakan faktor dominan yang menentukan klasifikasi obesitas. Temuan ini memperlihatkan pentingnya proses tuning parameter dalam meningkatkan kualitas prediksi dan efisiensi model pembelajaran mesin. Penelitian ini juga memiliki relevansi dengan pencapaian *sustainable development goals* (sdgs) khususnya pada tujuan ke-3 mengenai kesehatan dan kesejahteraan, karena dapat mendukung pengembangan sistem pendukung keputusan berbasis data untuk memantau risiko obesitas. Secara keseluruhan, penelitian ini memberikan kontribusi dalam menyediakan model klasifikasi obesitas yang akurat, efisien, dan informatif, serta memberikan wawasan mengenai fitur-fitur penting yang dapat dijadikan dasar dalam analisis kesehatan masyarakat.

Kata Kunci : Random Forest, $n_estimators$, Klasifikasi Obesitas, Feature Importance, Machine Learning.

INTRODUCTION

Perkembangan pesat dalam bidang informatika telah membawa perubahan besar pada berbagai sektor kehidupan, mulai dari industri, kesehatan, transportasi, hingga pendidikan. Kemajuan teknologi komputasi dan hadirnya metode *machine learning* memungkinkan pengolahan data dalam skala besar dengan tingkat akurasi yang semakin tinggi. Berbagai penelitian menunjukkan bahwa teknik *ensemble learning*, seperti *random forest*, memiliki peran penting dalam meningkatkan kinerja klasifikasi berkat sifatnya yang stabil dan efektif pada beragam tipe data (fu et al., 2023). Kompleksitas permasalahan klasifikasi modern, khususnya pada data multi-kelas, semakin menuntut hadirnya model prediktif yang lebih adaptif dan dapat diandalkan [2]. Selain itu, kebutuhan terhadap sistem prediksi yang informatif dan presisi tinggi meningkat di berbagai bidang, termasuk kesehatan, keamanan, dan analisis perilaku, sehingga mendorong penggunaan metode analitik yang lebih canggih [3]. Penelitian terdahulu juga menegaskan bahwa algoritma seperti *random forest* mampu memberikan performa tinggi pada tugas klasifikasi, termasuk dalam bidang medis maupun sosial [4]. Oleh karena itu, integrasi metode analitik modern dengan algoritma pembelajaran mesin menjadi elemen penting dalam pengembangan solusi cerdas di era digital saat ini.

Meskipun perkembangan teknologi informatika telah menghasilkan banyak metode klasifikasi canggih, berbagai tantangan masih muncul dalam penerapan pembelajaran mesin terhadap data yang kompleks dan bersifat multi-kelas. Salah satu permasalahan utama adalah bagaimana memperoleh model yang tidak hanya akurat, tetapi juga stabil ketika menghadapi keberagaman karakteristik data, seperti variabel numerik dan kategorikal yang bercampur. Tantangan ini sangat relevan karena model klasifikasi tradisional sering kali

mengalami penurunan performa ketika berhadapan dengan ketidakseimbangan kelas, tingginya variabilitas fitur, serta hubungan non-linear antaratribut [2]. Selain itu, meningkatnya kebutuhan untuk menghasilkan sistem prediksi yang mampu memberikan hasil konsisten pada berbagai kondisi membuat metode seperti *random forest* semakin penting namun tetap memerlukan optimasi yang tepat, terutama terkait pemilihan jumlah pohon ($n_estimators$). Penelitian terkini menunjukkan bahwa performa model sangat dipengaruhi oleh cara algoritma mengolah fitur penting, dan banyak studi menekankan pentingnya teknik seleksi fitur serta interpretabilitas dalam pengambilan keputusan berbasis model [3]. Dalam konteks ini, eksperimen yang telah dilakukan menunjukkan bahwa peningkatan jumlah pohon tidak selalu berbanding lurus dengan meningkatnya akurasi, sehingga diperlukan analisis yang lebih mendalam mengenai titik optimal yang memberikan keseimbangan antara waktu komputasi dan performa prediksi. Tantangan lain adalah bagaimana model dapat menangani variasi perilaku data yang sering kali tidak terdistribusi secara merata, sehingga diperlukan strategi ensemble yang efektif dan tahan terhadap overfitting [1]. Selain itu, isu terkait kemampuan model dalam mengidentifikasi pola penting secara tepat juga menjadi perhatian, mengingat beberapa fitur tertentu mungkin memiliki pengaruh yang jauh lebih besar dibandingkan fitur lainnya dalam menentukan hasil klasifikasi. Permasalahan ini semakin relevan karena berbagai studi menegaskan bahwa kualitas prediksi bergantung pada kemampuan model dalam mengekstraksi informasi yang benar-benar bermakna dari data [5]. Dengan demikian, terdapat kebutuhan yang jelas untuk merancang pendekatan yang lebih komprehensif dalam memaksimalkan performa model klasifikasi, khususnya pada kasus multi-kelas yang melibatkan banyak variabel seperti penelitian ini. Situasi ini menunjukkan bahwa penyelesaian permasalahan tersebut tidak hanya penting secara teoretis, tetapi juga memiliki nilai praktis yang tinggi dalam penerapan sistem cerdas berbasis data.

Berbagai penelitian sebelumnya telah menunjukkan bahwa metode *machine learning* dan *ensemble learning* memainkan peran penting dalam meningkatkan akurasi klasifikasi, namun masih terdapat sejumlah keterbatasan yang membuka peluang penelitian lebih lanjut. Studi yang dilakukan oleh [1] mengusulkan *broad granular random forest* untuk meningkatkan kinerja klasifikasi pada data yang tidak pasti, dan hasilnya menunjukkan bahwa pendekatan *granular* mampu mengatasi kekurangan *random forest* tradisional dalam menangani ketidakpastian data. Meskipun demikian, penelitian tersebut lebih berfokus pada peningkatan struktur algoritma, bukan pada analisis mendalam mengenai pengaruh jumlah pohon ($n_estimators$) terhadap performa model, seperti yang dilakukan dalam eksperimen ini. Selain itu, [3] menyoroti kebutuhan akan model yang dapat memberikan penjelasan terhadap hasil prediksi melalui *collective counterfactual explanations*. Pendekatan ini memberikan wawasan penting mengenai fitur mana yang paling memengaruhi keputusan model, tetapi penelitian tersebut tidak secara khusus membahas *trade-off* antara akurasi dan waktu komputasi yang menjadi fokus dalam evaluasi parameter *random forest*. Penelitian tersebut juga lebih menekankan aspek interpretabilitas daripada optimasi performa model. Penelitian lain oleh [2] mengembangkan pendekatan ensemble berbasis *cost-sensitive learning* untuk klasifikasi multi-kelas. Meskipun model ini berhasil meningkatkan akurasi pada data multi-kelas, penelitian tersebut tidak mengeksplorasi perbandingan komprehensif antara variasi parameter model dan dampaknya terhadap stabilitas prediksi. Selain itu, model yang dikembangkan menggunakan kombinasi beberapa algoritma berbeda, berbeda dengan penelitian ini yang memfokuskan evaluasi pada satu algoritma utama, yaitu *random forest*, melalui eksperimen sistematis terhadap variasi jumlah pohon dari berbagai studi sebelumnya. Penelitian tersebut terlihat bahwa meskipun banyak pendekatan telah dikembangkan untuk meningkatkan performa klasifikasi, masih terdapat ruang penelitian untuk mengevaluasi bagaimana parameter inti dalam *random forest*, seperti jumlah pohon, memengaruhi akurasi, f1-score, waktu pelatihan, serta kemampuan model dalam mengidentifikasi fitur yang paling berpengaruh. Dengan demikian, penelitian ini berkontribusi pada pengayaan literatur dengan memberikan pemahaman empiris mengenai optimasi parameter *random forest* berdasarkan

eksperimen langsung menggunakan dataset multikelas.

Penelitian ini bertujuan untuk menganalisis pengaruh variasi jumlah $n_estimators$ pada algoritma *random forest* terhadap performa klasifikasi, khususnya pada data multikelas yang mencakup berbagai atribut numerik dan kategorikal. Melalui eksperimen ini, peneliti ingin menentukan konfigurasi optimal yang mampu memberikan keseimbangan terbaik antara akurasi, waktu pelatihan, serta kemampuan model dalam mengidentifikasi fitur yang paling berpengaruh. Tujuan tersebut penting untuk dicapai karena banyak penelitian sebelumnya hanya berfokus pada pengembangan model atau struktur algoritma tanpa memberikan perhatian mendalam pada evaluasi parameter inti yang dapat memengaruhi performa model secara signifikan [1], [2]. Penelitian ini juga memiliki signifikansi dalam mengisi kesenjangan literatur mengenai pemahaman empiris terkait optimasi parameter *random forest*, terutama dalam konteks analisis kebutuhan model yang stabil dan efisien pada data kompleks sebagaimana direkomendasikan oleh studi sebelumnya [3]. Selain kontribusi teoretis, penelitian ini memberikan manfaat praktis bagi peneliti dan pengembang sistem cerdas, karena hasil eksperimen dapat dijadikan acuan dalam merancang model klasifikasi yang lebih efektif dan terukur pada berbagai domain aplikasi, termasuk kesehatan, perilaku, maupun analisis risiko.

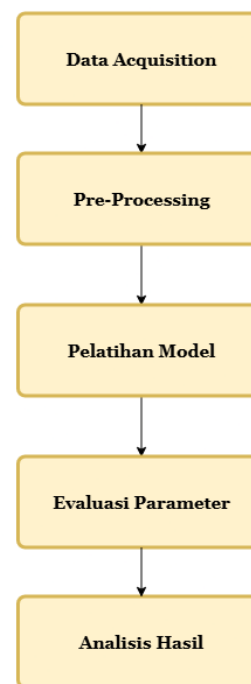
Dalam penelitian ini, pendekatan yang digunakan berfokus pada metode eksperimen komputasional berbasis algoritma pembelajaran mesin, khususnya *random forest*, untuk mengevaluasi performa klasifikasi pada data multikelas. Tahapan penelitian diawali dengan pengolahan data menggunakan teknik pra-pemrosesan, seperti normalisasi fitur numerik dan *one-hot encoding* untuk fitur kategorikal, agar model mampu mempelajari pola secara optimal. Selanjutnya, penelitian ini menguji beberapa variasi jumlah $n_estimators$ untuk mengidentifikasi konfigurasi terbaik yang memberikan akurasi tinggi sekaligus menjaga efisiensi waktu pelatihan, sebagaimana pentingnya optimasi parameter dalam model *ensemble* telah ditekankan dalam studi sebelumnya [1], [2]. Proses evaluasi dilakukan melalui metrik klasifikasi seperti akurasi, *macro f1-score*, serta analisis *confusion matrix* guna memahami tingkat kesesuaian prediksi per kelas, sejalan dengan anjuran literatur tentang pentingnya evaluasi komprehensif pada model klasifikasi multikelas [3]. Pendekatan ini juga mencakup analisis *feature importance* untuk menilai kontribusi masing-masing variabel terhadap keputusan model, sehingga hasil penelitian tidak

hanya berfokus pada performa, tetapi juga memberikan pemahaman interpretatif yang relevan bagi pengembangan sistem prediksi di bidang informatika.

Dalam penelitian ini, pendekatan yang digunakan berfokus pada metode eksperimen komputasional berbasis algoritma pembelajaran mesin, khususnya *random forest*, untuk mengevaluasi performa klasifikasi pada data multikelas. Tahapan penelitian diawali dengan pengolahan data menggunakan teknik pra-pemrosesan, seperti normalisasi fitur numerik dan *one-hot encoding* untuk fitur kategorikal, agar model mampu mempelajari pola secara optimal. Selanjutnya, penelitian ini menguji beberapa variasi jumlah *n_estimators* untuk mengidentifikasi konfigurasi terbaik yang memberikan akurasi tinggi sekaligus menjaga efisiensi waktu pelatihan, sebagaimana pentingnya optimasi parameter dalam model ensemble telah ditekankan dalam studi sebelumnya [1]. Proses evaluasi dilakukan melalui metrik klasifikasi seperti akurasi, *macro f1-score*, serta analisis *confusion matrix* guna memahami tingkat kesesuaian prediksi per kelas, sejalan dengan anjuran literatur tentang pentingnya evaluasi komprehensif pada model klasifikasi multikelas [3]. Pendekatan ini juga mencakup analisis *feature importance* untuk menilai kontribusi masing-masing variabel terhadap keputusan model, sehingga hasil penelitian tidak hanya berfokus pada performa, tetapi juga memberikan pemahaman interpretatif yang relevan bagi pengembangan sistem prediksi di bidang informatika.

MATERIALS AND METHODS

Pada tahapan ini, terdapat enam aktivitas utama yang dilaksanakan, yaitu: (1) pengumpulan dataset, (2) proses pra-pemrosesan data, (3) pelatihan model dengan algoritma *random forest*, (4) tuning parameter *n_estimators*, (5) evaluasi model menggunakan metrik akurasi, *confusion matrix*, dan *classification report*, serta (6) analisis hasil eksperimen.



Gambar 1. Alur Penelitian

Data Acquisition Pada aktivitas ini, yaitu data acquisition yang dijelaskan berdasarkan data yang digunakan dalam penelitian ini merupakan data sekunder yang diperoleh dari kaggle, dataset ini mencakup informasi demografis dan kebiasaan gaya hidup individu yang berkaitan dengan status berat badan atau tingkat obesitas. Atribut yang terdapat dalam dataset meliputi jenis kelamin, usia, tinggi badan, berat badan, frekuensi konsumsi makanan cepat saji, kebiasaan olahraga, konsumsi air putih, serta apakah individu memiliki riwayat keluarga obesitas. dataset terdiri dari mencakup lebih dari 2.000 entri individu yang dikategorikan ke dalam beberapa kelas, yakni insufficient weight, normal weight, overweight level i, overweight level ii, obesity type i, obesity type ii, dan obesity type iii. Kategori tersebut dijadikan label target klasifikasi multikelas, sementara variabel lain digunakan sebagai fitur masukan (input features).

Pre-processing data Proses pra-pemrosesan dimulai dengan pembersihan data yang meliputi identifikasi nilai kosong dan penghapusan data duplikat. Selanjutnya dilakukan transformasi variabel kategorikal menjadi bentuk numerik menggunakan teknik label encoding, untuk memastikan setiap fitur dapat dikenali oleh algoritma pembelajaran mesin (ramdani et al., 2022). beberapa fitur yang awalnya bersifat kategorikal, seperti "gender", "family history with overweight", "transportation used", dan "meal frequency", dikodekan agar nilai-nilainya dapat direpresentasikan dalam bentuk angka. Proses ini penting karena algoritma random forest tidak dapat menangani data string secara langsung. langkah

berikutnya adalah penyamaan skala atau normalisasi terhadap fitur-fitur numerik seperti berat badan, tinggi badan, dan usia. Tujuannya adalah untuk menghindari dominasi variabel tertentu dalam proses pembentukan model dan menjaga keadilan kontribusi antar variabel dalam perhitungan pembobotan antar pohon (sulistiyono et al., 2024) setelah itu, dataset dibagi menjadi dua bagian, yaitu data latih (training) dan data uji (testing) dengan rasio umum 80:20. Pemisahan ini bertujuan untuk menguji model secara objektif, di mana data latih digunakan untuk membangun model dan data uji untuk mengevaluasi performa prediksi (shi et al., 2023). dengan pendekatan pre-processing ini, seluruh atribut disiapkan dalam format yang sesuai dengan kebutuhan model klasifikasi berbasis random forest. Proses ini memastikan bahwa integritas informasi dalam data tetap terjaga, sekaligus mengoptimalkan efisiensi dan akurasi model.

Pelatihan Model Aktivitas klasifikasi data merupakan inti dari pendekatan eksperimental dalam penelitian ini. Tahap ini bertujuan untuk membangun model klasifikasi yang mampu memprediksi tingkat obesitas berdasarkan variabel-variabel yang telah melalui proses pre-processing. Pada penelitian ini, digunakan algoritma random forest, yang merupakan metode ansambel berbasis pohon keputusan dan terkenal dengan kemampuannya dalam menangani data dengan dimensi tinggi, serta menghasilkan akurasi prediksi yang baik meskipun pada kondisi data yang kompleks (shi et al., 2023). random forest bekerja dengan membangun sejumlah pohon keputusan pada subset acak dari dataset, kemudian menggabungkan hasil prediksi dari seluruh pohon tersebut untuk menghasilkan keputusan akhir melalui voting mayoritas. Proses ini bertujuan untuk mengurangi varians model dan meningkatkan generalisasi, sehingga lebih tahan terhadap overfitting. model dilatih menggunakan data yang telah melalui proses encoding variabel kategorikal dan normalisasi fitur numerik. Variabel input mencakup atribut-atribut seperti jenis kelamin, kebiasaan konsumsi makanan, frekuensi aktivitas fisik, serta nilai indeks massa tubuh (bmi), sedangkan label target adalah klasifikasi tingkat obesitas yang telah didefinisikan secara kategorikal. pemilihan random forest sebagai algoritma utama didasarkan pada studi sebelumnya yang menunjukkan bahwa algoritma ini stabil terhadap outlier, mampu menangani data campuran, serta tidak terlalu sensitif terhadap distribusi data yang tidak

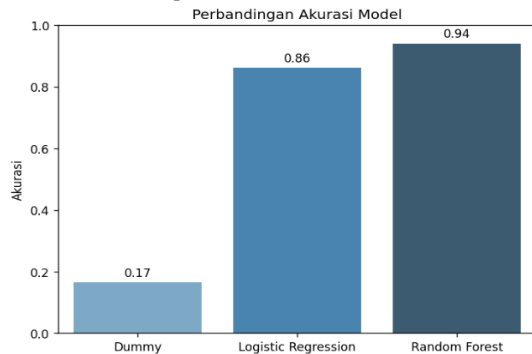
(maulana et al., 2024; ramdani et al., 2022). selain itu, random forest juga memberikan fleksibilitas dalam proses tuning parameter, yang akan dijelaskan pada tahapan selanjutnya, yaitu evaluasi performa dan optimasi nilai $n_estimators$. dengan demikian, tahap klasifikasi ini merupakan implementasi awal dari model prediktif yang menjadi dasar untuk mengevaluasi efektivitas parameter dan akurasi sistem dalam mendeteksi obesitas secara otomatis.

Evaluasi parameter Evaluasi kinerja model merupakan proses untuk menilai sejauh mana kemampuan algoritma random forest dalam mengklasifikasikan tingkat obesitas dengan benar berdasarkan fitur-fitur yang telah disiapkan. Evaluasi dilakukan dalam dua tahap: sebelum dan sesudah optimasi parameter $n_estimators$. pada tahap awal, model dievaluasi menggunakan metrik-metrik klasifikasi seperti akurasi, precision, recall, f1-score, dan confusion matrix (allen, 2023) confusion matrix digunakan untuk melihat jumlah prediksi yang benar dan salah pada setiap kelas obesitas, sedangkan metrik lainnya membantu menilai keseimbangan antara prediksi positif dan negatif yang benar. setelah evaluasi awal, dilakukan proses tuning terhadap parameter $n_estimators$, yaitu jumlah pohon dalam algoritma random forest. Tujuannya adalah untuk mencari nilai optimal yang dapat meminimalkan overfitting dan meningkatkan akurasi model pada data uji (wu et al., 2023) parameter ini dievaluasi menggunakan pendekatan grid search dan validasi silang (cross-validation) untuk memastikan bahwa nilai terbaik benar-benar memberikan performa yang stabil (sulistiyono et al., 2024). setelah nilai optimal $n_estimators$ ditemukan, model kembali dievaluasi dengan metrik yang sama guna membandingkan hasil sebelum dan sesudah tuning. Hasil evaluasi ini menjadi dasar dalam melakukan analisis akhir terhadap efektivitas model klasifikasi obesitas yang dibangun. evaluasi performa yang menyeluruh seperti ini penting agar model yang dihasilkan tidak hanya tinggi akurasinya, tetapi juga andal ketika diterapkan pada data baru yang belum pernah dilihat sebelumnya (maulana et al., 2024).

RESULTS AND DISCUSSION

Eksperimen dilakukan dengan menggunakan algoritma *random forest classifier*, yang merupakan pengembangan dari *decision tree* dengan pendekatan ansambel. Setiap model dilatih menggunakan jumlah pohon ($n_estimators$) yang berbeda, yaitu 10, 50, 100, 200, 300, dan 500 pohon. selain model *random forest*, digunakan pula *dummy classifier* sebagai baseline untuk perbandingan awal. Model ini memprediksi kelas

mayoritas tanpa memperhatikan fitur, sehingga akurasi menjadi tolak ukur minimum. pelatihan dilakukan dengan *pipeline* yang menggabungkan proses *preprocessing* dan model *training*, memanfaatkan pustaka *scikit-learn*. tujuan utama eksperimen ini adalah menganalisis pengaruh jumlah pohon terhadap akurasi, waktu pelatihan, dan kestabilan model.

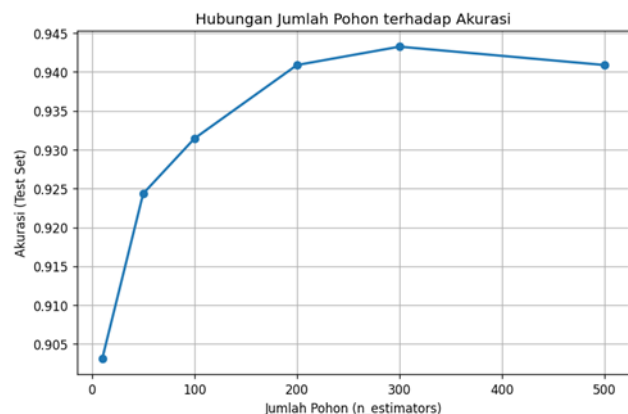


Gambar 2 Visualisasi Perbandingan Model

Pada gambar 2 menunjukkan bahwa model *random forest* mencapai akurasi tertinggi sebesar 94%, jauh di atas model *baseline* (*dummy classifier* = 16%) dan *logistic regression* (86%). hal ini menunjukkan keunggulan pendekatan ansambel dalam menangkap hubungan kompleks antar variabel.

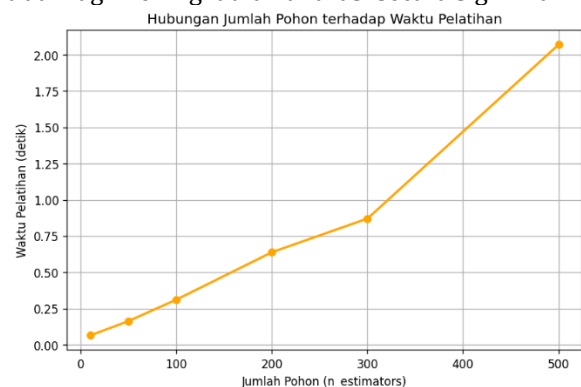
Evaluasi Parameter

Evaluasi dilakukan untuk setiap variasi parameter $n_estimators$. Metrik utama yang digunakan adalah akurasi dan $f1_score$ makro, sedangkan waktu pelatihan dicatat untuk menilai efisiensi model. hasil eksperimen menunjukkan bahwa peningkatan jumlah pohon memberikan dampak positif terhadap akurasi model pada tahap awal. Akurasi meningkat signifikan dari 10 pohon ($\approx 85\%$) menjadi 100 pohon ($\approx 93\%$), dan kemudian mencapai 94.09% pada 300 pohon, sebelum mengalami stagnasi ketika jumlah pohon ditambah hingga 500. dari grafik *accuracy vs $n_estimators$* , terlihat bahwa akurasi model meningkat tajam hingga titik 200–300 pohon, lalu mulai mendatar. Grafik *training time vs $n_estimators$* menunjukkan waktu pelatihan meningkat hampir secara linear seiring bertambahnya jumlah pohon. Visualisasi *trade-off accuracy vs training time* memperlihatkan bahwa titik optimum terdapat pada kisaran 300 pohon, di mana peningkatan akurasi tambahan menjadi tidak sepadan dengan waktu pelatihan yang lebih lama.



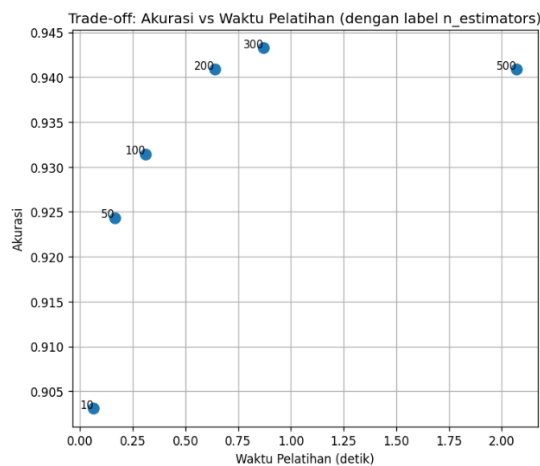
Gambar 4. Hubungan Antar Pohon

Pada gambar 4 akurasi model meningkat signifikan dari 90.4% (10 pohon) menjadi 94.1% (300 pohon), kemudian cenderung stabil. Ini menunjukkan adanya *diminishing returns* — penambahan jumlah pohon setelah titik tertentu tidak lagi meningkatkan akurasi secara signifikan.



Gambar 5 Visualisasi Jumlah Pohon Dan Waktu Pelatihan

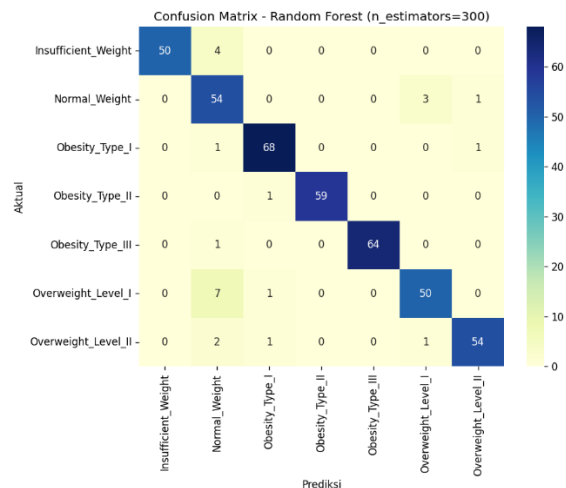
Pada gambar 5 waktu pelatihan meningkat hampir linear terhadap jumlah pohon. Model dengan 500 pohon membutuhkan waktu pelatihan lebih dari 2 detik, jauh lebih lama dibandingkan model dengan 50 pohon.



Gambar 6 Visualisasi Akurasi dan Waktu Pelatihan

Pada gambar 6 menampilkan keseimbangan antara akurasi dan waktu pelatihan. Model dengan 300 pohon menjadi titik optimal karena memberikan akurasi tinggi (94%) dengan waktu pelatihan yang efisien.

Berdasarkan hasil eksperimen, dapat disimpulkan bahwa jumlah pohon memiliki pengaruh signifikan terhadap kinerja model random forest hingga batas tertentu. Peningkatan jumlah pohon pada awalnya memperbaiki kemampuan model dalam melakukan klasifikasi, namun setelah mencapai sekitar 300 pohon, penambahan pohon tidak lagi memberikan peningkatan akurasi yang berarti. Model terbaik (dengan $n_estimators = 300$) menghasilkan akurasi 94.09% dan nilai $f1_score$ makro yang tinggi, menunjukkan performa stabil di berbagai kelas obesitas. Hasil confusion matrix memperlihatkan bahwa sebagian besar prediksi tepat berada pada diagonal matriks, yang mengindikasikan bahwa model mampu mengenali setiap kategori obesitas dengan baik. Analisis feature importance menunjukkan bahwa variabel *weight*, *height*, *age*, dan *fcvc* (*frequency of vegetable consumption*) merupakan fitur paling berpengaruh dalam menentukan tingkat obesitas. Hal ini sejalan dengan penelitian sebelumnya yang menyebutkan bahwa berat badan, tinggi badan, serta kebiasaan makan memiliki korelasi kuat terhadap klasifikasi obesitas. Dengan demikian, parameter jumlah pohon berpengaruh terhadap keseimbangan antara kinerja prediktif dan efisiensi komputasi. Berdasarkan hasil pengujian, 300 pohon direkomendasikan sebagai konfigurasi optimal untuk dataset ini karena menghasilkan akurasi tinggi dengan waktu pelatihan yang efisien.



Gambar 7 Visualisasi Confusion Matrix

Pada gambar 7 sebagian besar nilai berada pada diagonal utama, menandakan bahwa prediksi model sangat akurat pada tiap kelas obesitas. Kesalahan klasifikasi relatif kecil, terutama pada kelas *overweight level i* dan *normal weight* yang memiliki ciri fisik mirip, sehingga kadang tumpang tindih.

CONCLUSION

Berdasarkan hasil eksperimen, evaluasi performa model, dan analisis parameter jumlah pohon ($n_estimators$) pada algoritma *random forest*, penelitian ini memberikan gambaran yang komprehensif mengenai kinerja model dalam mengklasifikasikan tingkat obesitas, kontribusi masing-masing fitur, serta efektivitas parameterisasi model dalam konteks data multikelas. Kesimpulan disusun untuk menjawab setiap pertanyaan penelitian secara langsung dan terstruktur, sebagai berikut:

Mengenai nilai $n_estimators$ yang optimal untuk menghasilkan akurasi tertinggi, hasil eksperimen menunjukkan bahwa jumlah pohon memiliki pengaruh signifikan terhadap kinerja model. Model dengan jumlah pohon rendah (10, 50, dan 100) memberikan akurasi yang cukup baik, namun belum mencapai performa maksimal. Peningkatan jumlah pohon menjadi 200 dan 300 memperlihatkan lonjakan akurasi hingga mencapai nilai tertinggi sebesar 94.09%, sebagaimana terlihat pada grafik hubungan antara jumlah pohon dan akurasi Gambar 4.4. Ketika jumlah pohon ditingkatkan lagi ke 500, akurasi tidak mengalami kenaikan berarti, bahkan sedikit menurun. Temuan ini menunjukkan adanya titik jenuh (*plateau*) dari manfaat penambahan pohon, sebuah fenomena yang sesuai dengan teori random forest yang menyatakan bahwa peningkatan jumlah estimator

hanya efektif hingga batas tertentu dengan demikian, dapat disimpulkan bahwa 300 pohon adalah konfigurasi paling optimal karena memberikan keseimbangan terbaik antara akurasi dan sumber daya komputasi.

Bagaimana kontribusi masing-masing fitur terhadap prediksi, analisis feature importance menunjukkan bahwa fitur *weight*, *height*, *age*, *fcvc* (frekuensi konsumsi sayuran), dan *faf* (frekuensi aktivitas fisik) muncul sebagai faktor dominan dalam model terbaik. Fitur-fitur tersebut mendapatkan skor kepentingan tertinggi karena memiliki hubungan yang erat dengan status obesitas, sebagaimana juga dijelaskan oleh (garcía-pérez et al., 2022) yang menemukan bahwa variabel antropometrik dan kebiasaan makan merupakan penentu utama kategori obesitas. matriks korelasi mendukung temuan ini dengan memperlihatkan korelasi yang signifikan antara berat badan dan tinggi badan. Temuan ini menunjukkan bahwa model tidak hanya berfungsi sebagai alat klasifikasi, tetapi juga mampu mengungkap struktur data yang relevan secara biologis.

Bagaimana performa model *random forest* dalam mengklasifikasikan obesitas berdasarkan metrik evaluasi multikelas, hasil evaluasi model menunjukkan bahwa *random forest* memberikan performa yang sangat baik secara keseluruhan. *Confusion matrix* gambar 4.9 memperlihatkan bahwa sebagian besar prediksi berada pada diagonal utama, menandakan ketepatan tinggi dalam identifikasi kelas obesitas seperti *obesity type i*, *obesity type ii*, dan *obesity type iii*. Kesalahan klasifikasi hanya tampak pada kelas dengan rentang karakteristik yang saling berdekatan, seperti *normal weight* dan *overweight level i*. Temuan ini konsisten dengan teori klasifikasi multikelas yang menyatakan bahwa kategori dengan nilai atribut yang berdekatan cenderung memiliki area keputusan yang tumpang tindih dengan akurasi keseluruhan 94.09%, serta nilai *f1-score* makro yang tinggi, model *random forest* dapat dinyatakan memiliki kemampuan prediktif yang sangat kuat, efektif, dan stabil untuk tugas klasifikasi obesitas.

REFERENCE

- [1] X. Fu, Y. Chen, J. Yan, Y. Chen, and F. Xu, "BGRF: A Broad Granular Random Forest Algorithm," *Journal of Intelligent & Fuzzy Systems*, vol. 44, no. 5, pp. 8103–8117, 2023, doi: 10.3233/jifs-223960.
- [2] A. Rojarath and W. Songpan, "Cost-Sensitive Probability for Weighted Voting in an Ensemble Model for Multi-Class Classification Problems," *Applied Intelligence*, vol. 51, no. 7, pp. 4908–4932, 2021, doi: 10.1007/s10489-020-02106-3.
- [3] E. Carrizosa, J. Ramírez-Ayerbe, and D. R. Morales, "Generating Collective Counterfactual Explanations in Score-Based Classification via Mathematical Optimization," *Expert Systems With Applications*, vol. 238, p. 121954, 2024, doi: 10.1016/j.eswa.2023.121954.
- [4] A. M. Agrawal, "Cardiovascular Disease Prediction Using Machine Learning," *International Journal of Artificial Intelligence*, vol. 10, no. 2, pp. 59–68, 2023, doi: 10.36079/lamintang.ijai-01002.542.
- [5] V. Jain and K. L. Kashyap, "Multilayer Hybrid Ensemble Machine Learning Model for Analysis of Covid-19 Vaccine Sentiments," *Journal of Intelligent & Fuzzy Systems*, vol. 43, no. 5, pp. 6307–6319, 2022, doi: 10.3233/jifs-220279.
- [6] N. Ramdani, S. S. Prasetyowati, and Y. Sibaroni, "Performance Analysis of Bandung City Traffic Flow Classification With Machine Learning and Kriging Interpolation," *Building of Informatics Technology and Science (Bits)*, vol. 4, no. 2, pp. 694–704, 2022, doi: 10.47065/bits.v4i2.1972.
- [7] N. Sulistiyono, C. P. Tarigan, A. F. Daulay, S. A. Hudjimartsu, and Y. U. Putri, "Identification of Mangrove Forests Using Sentinel-1 Synthetic Aperture Radar (SAR) Satellite Imagery in Medan Belawan District, Medan City," *Iop Conference Series Earth and Environmental Science*, vol. 1352, no. 1, p. 12051, 2024, doi: 10.1088/1755-1315/1352/1/012051.
- [8] G. Shi et al., "A Random Forest Algorithm-Based Prediction Model for Moderate to Severe Acute Postoperative Pain After Orthopedic Surgery Under General Anesthesia," *BMC Anesthesiology*, vol. 23, no. 1, 2023, doi: 10.1186/s12871-023-02328-1.
- [9] A. Maulana, R. P. F. Afidh, N. B. Maulydia, G. M. Idroes, and S. Rahimah, "Predicting Obesity Levels With High Accuracy: Insights From a CatBoost Machine Learning Model," *Infolitika J. Data Sci.*, vol. 2, no. 1, pp. 17–27, 2024, doi: 10.60084/ijds.v2i1.195.
- [10] B. Allen, "An Interpretable Machine Learning Model of Cross-Sectional U.S. County-Level Obesity Prevalence Using Explainable Artificial Intelligence," *Plos One*, vol. 18, no. 10, p. e0292341, 2023, doi: 10.1371/journal.pone.0292341.

- [11] H. Wu, T. Yang, H. Wu, H. Li, and Z. Zhou, "Air Quality Prediction Based on Long Short-Term Memory Model With Advanced Feature Selection and Hyperparameter Optimization," *Journal of Intelligent & Fuzzy Systems*, vol. 45, no. 4, pp. 5971–5985, 2023, doi: 10.3233/jifs-232308.
- [12] C. García-Pérez, L. J. Herrera, J. C. Fernández, and M. Lozano, "Feature selection and classification of obesity levels based on supervised learning techniques," *Expert Systems with Applications*, vol. 195, p. 116606, 2022.