

Klasifikasi Risiko Penyakit Metabolik dengan model Ensemble, Decision Tree Dan Logistic Regression Pada Posbindu Universitas Gadjah Mada

Muhammad Alif Taufiqurrahman*, Christian Dewa Siaga Pratama

Program Studi Magister Ilmu Komputer, Universitas Budi Luhur, Jakarta, Indonesia

Email: ¹muhalfs20@gmail.com, ²christiandewa99@gmail.com

Email Penulis Korespondensi: muhalfs20@gmail.com

Submitted 01-07-2026; Accepted 30-08-2026; Published 31-08-2026

Abstrak

Penelitian ini bertujuan untuk mengembangkan dan mengevaluasi model ensemble berbasis stacking yang mengombinasikan algoritma *Decision Tree* sebagai *base learner* dan *Logistic Regression* sebagai *meta-learner* dalam klasifikasi risiko penyakit metabolik menggunakan data Posbindu UGM. Dataset mencakup parameter kesehatan utama, antara lain tekanan darah, kadar gula darah, kolesterol, lingkar perut, dan indeks massa tubuh. Metodologi penelitian meliputi tahap pra-pemrosesan data, pembentukan label risiko berdasarkan interpretasi medis, pengembangan model klasifikasi, serta evaluasi performa menggunakan metrik accuracy, precision, recall, dan F1-score dengan pendekatan Stratified K-Fold Cross Validation. Hasil penelitian menunjukkan bahwa model ensemble (*Decision Tree* + *Logistic Regression*) mampu mengklasifikasikan risiko penyakit metabolik secara stabil dan konsisten, dengan akurasi sebesar 94,87% dan F1-score sebesar 0,9187, serta kinerja yang sebanding atau lebih baik dibandingkan model tunggal. Analisis confusion matrix menunjukkan bahwa model ensemble lebih efektif dalam mengurangi kesalahan klasifikasi pada kelas Berisiko dan Risiko Tinggi, yang secara klinis memiliki karakteristik yang saling beririsan. Selain itu, analisis kontribusi fitur menunjukkan bahwa Indeks Massa Tubuh (IMT), kolesterol, tekanan darah, dan kadar gula darah merupakan faktor dominan dalam penentuan risiko penyakit metabolik. Penelitian ini menyimpulkan bahwa pendekatan ensemble stacking tidak hanya meningkatkan stabilitas klasifikasi, tetapi juga menghasilkan model yang interpretable dan relevan secara klinis, sehingga berpotensi mendukung pengembangan sistem deteksi dini risiko penyakit metabolik berbasis Posbindu PTM. Secara lebih luas, hasil penelitian ini berkontribusi pada penguatan implementasi Posbindu berbasis teknologi di lingkungan universitas serta mendukung kebijakan nasional dalam pengendalian penyakit tidak menular melalui pemanfaatan data kesehatan komunitas secara optimal.

Kata Kunci: Posbindu; Penyakit Metabolic; Klasifikasi Risiko; *Hybrid Model*; *Decision Tree*; *Logistic Regression*; *Health Promoting University*.

Abstract

This study aims to develop and evaluate a stacking-based ensemble model that combines the *Decision Tree* algorithm as the base learner and *Logistic Regression* as the meta-learner in classifying metabolic disease risk using UGM Posbindu data. The dataset includes key health parameters, such as blood pressure, blood sugar levels, cholesterol, waist circumference, and body mass index. The research methodology includes data pre-processing, risk label formation based on medical interpretation, classification model development, and performance evaluation using accuracy, precision, recall, and F1-score metrics with a Stratified K-Fold Cross Validation approach. The results show that the ensemble model (*Decision Tree* + *Logistic Regression*) is capable of classifying metabolic disease risk in a stable and consistent manner. The confusion matrix analysis shows that the ensemble model is more effective in reducing classification errors in the Risk and High Risk classes, which clinically have overlapping characteristics. Furthermore, feature contribution analysis shows that Body Mass Index (BMI), cholesterol, blood pressure, and blood sugar levels are dominant factors in determining metabolic disease risk. This study concludes that the stacking ensemble approach not only improves classification stability but also produces interpretable and clinically relevant models, thereby potentially supporting the development of a Posbindu PTM-based early detection system for metabolic disease risk. More broadly, the results of this study contribute to strengthening the implementation of technology-based Posbindu in the university environment and support national policies in controlling non-communicable diseases through the optimal use of community health data.

Keywords: Posbindu; Metabolic Diseases; Risk Classification; *Hybrid Model*, *Decision Tree*; *Logistic Regression*; *Health Promoting University*

1. PENDAHULUAN

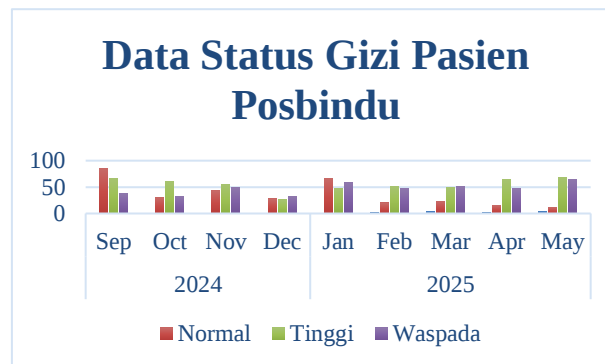
Penyakit metabolik merupakan salah satu masalah kesehatan yang semakin meningkat prevalensinya di seluruh dunia, termasuk di Indonesia. Kondisi ini mencakup berbagai gangguan seperti hipertensi, diabetes melitus tipe 2, obesitas, dan dislipidemia yang secara kolektif berkontribusi besar terhadap morbiditas dan mortalitas global[1],[2]. Di Indonesia, beban penyakit tidak menular ini terus meningkat seiring dengan perubahan gaya hidup dan pola makan masyarakat, sehingga memerlukan upaya pencegahan dan deteksi dini yang efektif. Dalam konteks ini, klasifikasi risiko penyakit metabolik menjadi sangat penting untuk mengidentifikasi individu yang berpotensi mengalami komplikasi serius sehingga intervensi dapat dilakukan secara tepat waktu dan terarah.

Universitas Gadjah Mada (UGM) sebagai perguruan tinggi negeri berbadan hukum di Indonesia memiliki komitmen kuat dalam mendukung pelaksanaan Tridarma Perguruan Tinggi, yang meliputi pendidikan, penelitian, dan pengabdian kepada masyarakat. Untuk memastikan kelancaran pelaksanaan Tridarma tersebut, kesehatan pegawai menjadi faktor krusial yang tidak dapat diabaikan. Pegawai yang sehat akan lebih produktif dan mampu memberikan pelayanan optimal dalam mendukung kegiatan akademik dan non-akademik di lingkungan universitas.



Gambar 1. Grafik Pemeriksaan Tekanan Darah Pasien Posbindu

Data hasil pemeriksaan Posbindu UGM pada periode September 2024 hingga Mei 2025 menunjukkan kecenderungan kondisi kesehatan pegawai yang perlu mendapatkan perhatian serius, khususnya terkait faktor risiko penyakit metabolik. Pada grafik tekanan darah Gambar 1, terlihat bahwa jumlah peserta dengan kategori hipertensi mendominasi pada beberapa bulan, terutama pada September 2024 dan Januari–Maret 2025. Sementara itu, peserta dengan tekanan darah normal cenderung lebih sedikit dibandingkan kategori prehipertensi dan hipertensi, menunjukkan bahwa sebagian besar pegawai berada pada kondisi tekanan darah yang tidak ideal. Hal ini mengindikasikan adanya beban risiko tinggi terhadap penyakit kardiovaskular dan perlu adanya pemantauan yang lebih ketat.



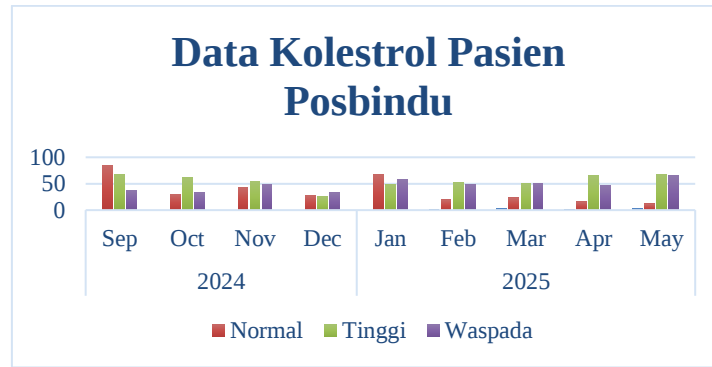
Gambar 2. Status Gizi Pasien Posbindu

Jika dilihat dari grafik status gizi Gambar 2, sebagian besar peserta berada dalam kategori Waspada, dengan tren yang relatif stabil dari bulan ke bulan. Kondisi ini menunjukkan bahwa banyak pegawai yang memiliki status gizi tidak normal dan berpotensi mengalami kelebihan berat badan. Status gizi yang tidak ideal ini berhubungan langsung dengan risiko penyakit metabolik seperti diabetes, hipertensi, dan dislipidemia. Dominasi kategori Waspada dan Tinggi menunjukkan perlunya intervensi yang lebih intensif terkait pola makan dan aktivitas fisik.



Gambar 3. Tekanan Darah Pasien Posbindu

Pada grafik tekanan darah versi kategori Normal–Waspada–Tinggi Gambar 3, pola yang muncul serupa dengan temuan pada Gambar 1, yaitu jumlah peserta pada kategori Waspada dan Tinggi lebih banyak dibandingkan kelompok normal pada hampir seluruh periode. Hal ini mempertegas bahwa tekanan darah merupakan salah satu parameter yang paling banyak menunjukkan penyimpangan pada populasi pegawai yang mengikuti pemeriksaan Posbindu.



Gambar 4. Kolesterol Pasien Posbindu

Selanjutnya, grafik kolesterol Gambar 4 juga memperlihatkan fluktuasi yang cukup tinggi sepanjang periode pemeriksaan. Pada beberapa bulan, jumlah peserta dalam kategori tinggi dan waspada terlihat cukup menonjol, menandakan adanya kecenderungan gangguan metabolik lipid di kalangan pegawai. Kolesterol yang tidak stabil dapat menjadi indikator awal risiko penyakit jantung koroner maupun gangguan metabolik lainnya, sehingga pemeriksaan berkala dan analisis tren menjadi penting untuk mendukung upaya pencegahan.

Secara keseluruhan, keempat grafik tersebut memberikan gambaran bahwa kondisi kesehatan sebagian pegawai UGM cenderung berada pada kategori risiko, baik dari aspek tekanan darah, status gizi, maupun kadar kolesterol. Pola fluktuatif yang konsisten pada berbagai parameter kesehatan menunjukkan perlunya sistem deteksi dini berbasis analitik untuk mengidentifikasi individu dengan risiko tinggi secara lebih akurat. Temuan ini menjadi dasar yang kuat dalam pengembangan model klasifikasi risiko penyakit metabolik berbasis pendekatan *hybrid Decision Tree* dan *Logistic Regression*, sebagaimana dikaji dalam penelitian ini.

Sebagai bagian dari upaya promotif dan preventif dalam menjaga kesehatan civitas akademika, UGM telah mengimplementasikan program Pos Pembinaan Terpadu Penyakit Tidak Menular (Posbindu PTM) yang sudah mulai diprogram sejak tahun 2023. Program ini dilaksanakan secara rutin oleh Biro Pelayanan Kesehatan Terpadu (BPKT) dan merupakan bagian dari inisiatif *Health Promoting University (HPU)*. Kegiatan Posbindu PTM mencakup pengukuran antropometri, pemeriksaan tekanan darah, kolesterol, gula darah, asam urat, serta konsultasi gizi dan kesehatan mental[3].

Namun, meskipun data hasil pemeriksaan Posbindu telah dikumpulkan secara rutin, pemanfaatan data tersebut untuk analisis lebih lanjut, khususnya dalam mengidentifikasi risiko penyakit metabolik yang masih terbatas. Penyakit metabolik seperti diabetes melitus, hipertensi, dan dislipidemia merupakan masalah kesehatan yang signifikan dan dapat mempengaruhi produktivitas kerja pegawai. Deteksi dini terhadap risiko penyakit ini sangat penting untuk mencegah komplikasi lebih lanjut dan menurunkan beban biaya kesehatan.

Peraturan Menteri Kesehatan Nomor 71 Tahun 2015 disusun untuk merespons tingginya beban penyakit tidak menular (PTM) di masyarakat, yang berdampak signifikan terhadap angka kesakitan, kecacatan, kematian, serta pembiayaan kesehatan. Tujuan utama dari peraturan ini adalah melindungi masyarakat dari risiko PTM, meningkatkan kualitas hidup, serta memberikan kepastian hukum dalam pelaksanaan penanggulangan PTM yang komprehensif, efisien, efektif, dan berkelanjutan [4].

Penelitian sebelumnya telah menunjukkan bahwa penggabungan beberapa algoritma machine learning mampu meningkatkan akurasi prediksi risiko penyakit metabolik. [5] mengembangkan model berbasis *Decision Tree* dan *Logistic Regression* untuk prediksi diabetes menggunakan data klinis dan demografis, dan menunjukkan bahwa model gabungan tersebut menghasilkan kinerja yang lebih baik dibandingkan model tunggal. Pendekatan ini menegaskan bahwa kombinasi model klasifikasi yang saling melengkapi mampu meningkatkan stabilitas dan akurasi prediksi risiko kesehatan. Namun demikian, penelitian tersebut masih menggunakan pendekatan hybrid secara umum dan belum secara eksplisit menerapkan mekanisme ensemble berbasis stacking, serta berfokus pada prediksi diabetes klinis, bukan pada klasifikasi risiko dalam konteks skrining kesehatan populasi seperti Posbindu.

Model tunggal seperti *Decision Tree* atau *Logistic Regression* memiliki keterbatasan. DT rentan overfitting pada data bervariasi tinggi, sedangkan LR kurang mampu memodelkan hubungan non-linear antar variabel. Oleh karena itu, penelitian ini mengembangkan DT sebagai base classifier untuk mempartisi data dan menangkap pola non-linear, dengan LR sebagai meta-model untuk menghitung probabilitas risiko secara interpretable[6],[7]. Kombinasi ini telah terbukti pada penelitian kesehatan sebelumnya mampu meningkatkan akurasi, recall, dan interpretabilitas dibandingkan model tunggal, sehingga relevan diterapkan pada data Posbindu UGM yang memiliki kompleksitas variabel kesehatan beragam. Hasil dari penelitian ini dapat menjadi dasar dalam pengambilan keputusan terkait intervensi kesehatan yang lebih tepat sasaran, sehingga mendukung tercapainya tujuan UGM sebagai Health Promoting University dan mendukung pelaksanaan Tridarma Perguruan Tinggi secara optimal[8].

Penelitian ini bertujuan untuk mengembangkan model ensemble (stacking) yang menggabungkan algoritma *Decision Tree* sebagai *base learner* dan *Logistic Regression* sebagai *meta-learner* dalam mengklasifikasikan risiko penyakit metabolik berdasarkan data Posbindu pegawai Universitas Gadjah Mada. Model yang dikembangkan diharapkan mampu memanfaatkan data pemeriksaan kesehatan secara optimal sehingga dapat menghasilkan klasifikasi risiko yang lebih akurat dibandingkan penggunaan model tunggal. Selain membangun model, penelitian ini juga bertujuan

mengevaluasi performa model ensemble melalui metode Stratified K-Fold Cross Validation dengan menggunakan metrik accuracy, precision, recall, F1-score, dan ROC-AUC, sehingga dapat diketahui tingkat efektivitas model dalam mengidentifikasi individu yang berisiko mengalami penyakit metabolik.

Selanjutnya, penelitian ini bertujuan mengidentifikasi variabel kesehatan yang memiliki kontribusi paling signifikan terhadap prediksi risiko penyakit metabolik melalui analisis feature importance pada Decision Tree dan koefisien pada Logistic Regression, sehingga dapat diperoleh informasi mengenai faktor-faktor risiko dominan yang perlu menjadi perhatian dalam upaya pencegahan. Hasil penelitian ini juga diharapkan dapat menjadi dasar dalam merumuskan rekomendasi pengembangan sistem deteksi dini berbasis data untuk mendukung pengambilan keputusan pada program Health Promoting University di Universitas Gadjah Mada. Dengan demikian, model yang dihasilkan tidak hanya memberikan kontribusi pada pengembangan metode klasifikasi di bidang *machine learning*, tetapi juga mendukung peningkatan kualitas layanan kesehatan preventif melalui pemanfaatan data Posbindu secara lebih efektif dan berkelanjutan.

2. METODOLOGI PENELITIAN

2.1 Metode Penelitian

Penelitian ini menggunakan metode penelitian kuantitatif dengan pendekatan eksperimen untuk mengembangkan dan mengevaluasi model *ensemble (stacking)* dalam klasifikasi risiko penyakit metabolik berdasarkan data Posbindu pegawai Universitas Gadjah Mada (UGM). Pendekatan kuantitatif dipilih karena penelitian memanfaatkan data numerik hasil pemeriksaan kesehatan yang dianalisis menggunakan algoritma *machine learning*. Melalui pendekatan eksperimen, dilakukan pengujian terhadap model yang dibangun untuk mengetahui kemampuan klasifikasi serta membandingkan performanya dengan model tunggal.

Model *ensemble* yang dikembangkan menggabungkan algoritma Decision Tree sebagai *base learner* dan Logistic Regression sebagai *meta learner*. Decision Tree digunakan untuk mempelajari pola hubungan yang bersifat nonlinier antarvariabel kesehatan, sedangkan Logistic Regression digunakan untuk mengombinasikan hasil prediksi sehingga menghasilkan klasifikasi yang lebih stabil dan akurat [9]. Seluruh proses penelitian meliputi pengumpulan data, prapengolahan data, pembangunan model, evaluasi model, hingga interpretasi hasil klasifikasi.

2.2 Sumber dan Pengumpulan Data

2.2.1 Sumber Data

Data yang digunakan merupakan data sekunder yang berasal dari hasil pemeriksaan kesehatan program Pos Pembinaan Terpadu Penyakit Tidak Menular (Posbindu PTM) yang diselenggarakan oleh Biro Pelayanan Kesehatan Terpadu (BPKT) Universitas Gadjah Mada. Dataset mencakup hasil pemeriksaan pegawai selama periode tahun 2024–2025 yang terdiri atas parameter tekanan darah, kadar gula darah, kolesterol total, asam urat, berat badan, tinggi badan, indeks massa tubuh (IMT), lingkaran perut, serta informasi demografis yang mendukung proses analisis.

Sebelum digunakan dalam penelitian, seluruh data telah dianonimkan dengan menghapus identitas pribadi responden sehingga memenuhi prinsip kerahasiaan dan etika penelitian. Penggunaan data dilakukan setelah memperoleh izin dari pihak pengelola data sesuai dengan ketentuan yang berlaku.

2.2.2 Pengumpulan Data

Pengumpulan data dilakukan melalui dokumentasi hasil pemeriksaan kesehatan Posbindu PTM yang telah tersimpan dalam basis data BPKT UGM. Selanjutnya dilakukan proses seleksi terhadap data yang sesuai dengan kebutuhan penelitian berdasarkan variabel yang digunakan. Seluruh data kemudian diperiksa kualitasnya sebelum memasuki tahap prapengolahan agar data yang dianalisis memiliki tingkat validitas dan konsistensi yang baik.

2.3 Instrumen Penelitian

Instrumen penelitian terdiri atas instrumen pengukuran kesehatan dan instrumen analisis data. Keandalan pengukuran tekanan darah dan penggunaan perangkat digital dalam pemantauan kesehatan telah menjadi perhatian penting dalam penelitian klinis dan komunitas [10], [11], [12]. Instrumen pengukuran kesehatan meliputi tensimeter digital, glukometer, alat ukur kolesterol, alat ukur asam urat, timbangan digital, stadiometer, serta pita pengukur lingkaran perut yang digunakan dalam kegiatan Posbindu PTM. Seluruh alat telah dikalibrasi sesuai standar pelayanan kesehatan.

Instrumen analisis data menggunakan perangkat lunak Python dengan pustaka Pandas, NumPy, Scikit-learn, Matplotlib, dan Jupyter Notebook. Perangkat lunak tersebut digunakan untuk melakukan proses *preprocessing*, pembangunan model *stacking*, pelatihan model, validasi, serta evaluasi performa menggunakan berbagai metrik klasifikasi.

2.4 Tahapan Penelitian

2.4.1 Data Selection (Seleksi Data)

Tahap awal dilakukan dengan memilih data yang relevan terhadap tujuan penelitian. Variabel yang digunakan meliputi tekanan darah sistolik dan diastolik, kadar gula darah, kolesterol total, indeks massa tubuh, lingkaran perut, usia, serta jenis kelamin [13]. Variabel yang tidak berhubungan langsung dengan proses klasifikasi seperti nama, nomor identitas, dan informasi administratif dihapus dari dataset.

2.4.2 Data Preprocessing (Prapengolahan Data)

Tahap ini bertujuan meningkatkan kualitas data sebelum digunakan dalam proses pelatihan model. Proses yang dilakukan meliputi penghapusan data duplikat, penanganan *missing value*, penghapusan nilai pencilan yang tidak logis secara medis, normalisasi data numerik, pengkodean data kategorikal (*encoding*), serta pemilihan fitur (*feature selection*) yang memiliki pengaruh terhadap klasifikasi risiko penyakit metabolik.

2.4.3 Data Transformation (Transformasi Data)

Data yang telah dibersihkan kemudian ditransformasikan agar sesuai dengan kebutuhan algoritma *machine learning*. Tahapan ini meliputi normalisasi menggunakan *StandardScaler*, pembentukan variabel turunan seperti kategori IMT dan kategori tekanan darah, serta pembagian dataset menjadi data latih dan data uji menggunakan metode *Stratified K-Fold Cross Validation*.

2.4.4 Pemodelan Stacking Ensemble

Pada tahap ini dibangun model *stacking ensemble* dengan menggunakan Decision Tree sebagai *base learner* dan Logistic Regression sebagai *meta learner*. Decision Tree terlebih dahulu menghasilkan prediksi terhadap data latih, kemudian hasil prediksi tersebut dijadikan masukan bagi Logistic Regression untuk menghasilkan klasifikasi akhir risiko penyakit metabolik.

2.4.5 Evaluasi Model

Evaluasi dilakukan untuk mengetahui kemampuan model dalam mengklasifikasikan risiko penyakit metabolik. Metrik evaluasi yang digunakan meliputi *Accuracy*, *Precision*, *Recall*, *F1-Score*, dan *Receiver Operating Characteristic–Area Under Curve (ROC-AUC)*. Selain itu, performa model *stacking* dibandingkan dengan model Decision Tree dan Logistic Regression secara terpisah.

2.4.6 Implementasi Model

Model terbaik yang diperoleh selanjutnya digunakan untuk melakukan simulasi klasifikasi terhadap data baru sehingga dapat memberikan prediksi tingkat risiko penyakit metabolik. Hasil klasifikasi diharapkan dapat dimanfaatkan sebagai dasar pengambilan keputusan dalam program deteksi dini di lingkungan Universitas Gadjah Mada.

2.4.7 Validasi Model

Validasi dilakukan menggunakan metode *Stratified K-Fold Cross Validation* untuk memastikan model memiliki kemampuan generalisasi yang baik terhadap data yang belum pernah digunakan pada proses pelatihan. Hasil validasi digunakan untuk menilai konsistensi dan keandalan model sebelum diterapkan pada sistem pendukung keputusan.

2.5 Teknik Analisis Data

2.5.1 Arsitektur Model Stacking

Model yang dikembangkan menggunakan pendekatan *stacking ensemble*, yaitu mengombinasikan Decision Tree sebagai *base learner* dan Logistic Regression sebagai *meta learner*. Pendekatan ini bertujuan meningkatkan akurasi klasifikasi dengan memanfaatkan kelebihan masing-masing algoritma.

2.5.2 Decision Tree sebagai Base Learner

Decision Tree digunakan untuk membentuk aturan klasifikasi berdasarkan hubungan antarvariabel kesehatan. Algoritma ini menghasilkan struktur pohon keputusan yang mampu menangkap pola data nonlinier sehingga efektif dalam proses klasifikasi awal[6],[14].

2.5.3 Logistic Regression sebagai Meta Learner

Logistic Regression menerima keluaran dari Decision Tree sebagai masukan untuk menghasilkan prediksi akhir. Penggunaan algoritma ini memberikan estimasi probabilitas yang mudah diinterpretasikan dan meningkatkan stabilitas hasil klasifikasi[7],[15].

2.5.4 Pengukuran Kinerja Model

Kinerja model dievaluasi menggunakan *Confusion Matrix*, *Accuracy*, *Precision*, *Recall*, *F1-Score*, dan *ROC-AUC*. Hasil evaluasi kemudian dibandingkan antara model *stacking* dan model tunggal untuk mengetahui efektivitas pendekatan yang diusulkan dalam mengklasifikasikan risiko penyakit metabolik.

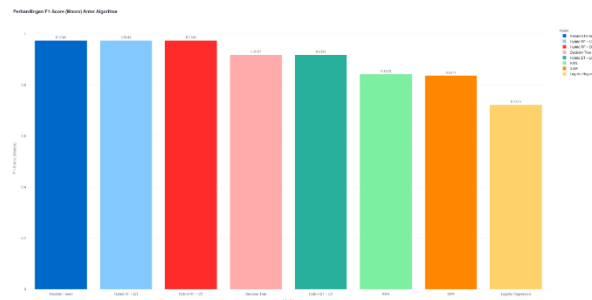
Berdasarkan tahapan penelitian yang telah dirancang, mulai dari pengumpulan data, prapengolahan data, transformasi data, pembangunan model *stacking ensemble*, hingga evaluasi dan validasi model, diharapkan penelitian ini mampu menghasilkan model klasifikasi risiko penyakit metabolik yang akurat, andal, dan memiliki kemampuan generalisasi yang baik. Penerapan metode *stacking* yang mengombinasikan Decision Tree sebagai *base learner* dan Logistic Regression sebagai *meta learner* diharapkan dapat meningkatkan performa klasifikasi dibandingkan model tunggal serta memberikan informasi yang bermanfaat dalam mendukung deteksi dini penyakit metabolik. Dengan demikian, hasil penelitian ini diharapkan dapat menjadi dasar pengembangan sistem pendukung keputusan berbasis data pada program Posbindu Universitas Gadjah Mada serta memberikan kontribusi terhadap pemanfaatan *machine learning* dalam upaya promotif dan preventif di bidang kesehatan.

3. HASIL DAN PEMBAHASAN

3.1 Hasil Penelitian Algoritma Klasifikasi Tunggal

Pada tahap awal pengujian, penelitian ini mengevaluasi performa beberapa algoritma klasifikasi tunggal, yaitu *Decision Tree*, *Logistic Regression*, Random Forest, Support Vector Machine (SVM), dan K-Nearest Neighbor (KNN). Evaluasi

dilakukan menggunakan data uji (test set) untuk mengukur kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya. Metrik evaluasi yang digunakan meliputi Accuracy, Precision (Macro), Recall (Macro), dan F1-Score (Macro), yang dipilih untuk memberikan gambaran performa model secara menyeluruh pada kondisi distribusi kelas yang tidak seimbang.



Gambar 5. Hasil Testing Perbandingan Algoritma

Hasil pengujian pada Gambar 5 menunjukkan bahwa *Random Forest* memperoleh performa terbaik di antara algoritma tunggal yang diuji, ditandai dengan nilai akurasi dan F1-Score yang paling tinggi serta stabil. Keunggulan ini dapat dijelaskan oleh karakteristik *Random Forest* sebagai metode ensemble berbasis pohon keputusan yang mampu menangkap hubungan non-linear antar fitur kesehatan serta mengurangi risiko *over fitting* melalui mekanisme bagging dan pemilihan fitur secara acak. Dengan demikian, *Random Forest* dapat berfungsi sebagai baseline kuat dalam mengevaluasi efektivitas pendekatan *hybrid* pada tahap selanjutnya.

Algoritma *Decision Tree* juga menunjukkan performa yang relatif baik, terutama dari sisi interpretabilitas model. Struktur pohon keputusan yang dihasilkan memungkinkan identifikasi aturan klasifikasi berbasis parameter kesehatan yang mudah dipahami secara klinis. Namun demikian, performa *Decision Tree* tunggal masih berada di bawah *Random Forest*, yang mengindikasikan bahwa model ini lebih sensitif terhadap variasi data dan berpotensi mengalami *over fitting* ketika dihadapkan pada data dengan kompleksitas tinggi dan distribusi kelas yang tidak seimbang.

Sebaliknya, *Logistic Regression* menunjukkan performa yang lebih rendah dibandingkan algoritma lainnya. Temuan ini mengindikasikan bahwa hubungan antar parameter kesehatan dalam data Posbindu tidak sepenuhnya bersifat linier, sehingga pendekatan regresi linier kurang optimal dalam memodelkan kompleksitas risiko penyakit metabolik. Meskipun demikian, keunggulan *Logistic Regression* terletak pada stabilitas dan interpretasi probabilistik, yang menjadikannya relevan untuk digunakan sebagai komponen dalam model *hybrid*.

Secara keseluruhan, hasil pengujian algoritma klasifikasi tunggal menunjukkan bahwa tidak terdapat satu algoritma tunggal yang unggul dalam seluruh aspek, baik dari sisi akurasi, stabilitas, maupun interpretabilitas. Temuan ini memperkuat dasar pemilihan pendekatan model *hybrid*, yang bertujuan mengombinasikan keunggulan model berbasis aturan seperti *Decision Tree* dengan kemampuan prediksi probabilistik dan stabilitas model statistik seperti *Logistic Regression*.

3.2 Hasil Penelitian Model Hybrid

3.2.1 Konsep dan Implementasi Model Hybrid

Model *hybrid* dalam penelitian ini dikembangkan dengan mengombinasikan dua algoritma klasifikasi menggunakan pendekatan stacking ensemble. Pada pendekatan ini, satu atau lebih model dasar (*base learner*) digunakan untuk mempelajari pola awal pada data dan menghasilkan prediksi atau probabilitas kelas, yang selanjutnya digunakan sebagai masukan bagi model tingkat lanjut (*meta learner*). Strategi ini bertujuan untuk mengintegrasikan keunggulan masing-masing algoritma, sehingga mampu meningkatkan stabilitas dan kualitas prediksi dibandingkan penggunaan model tunggal.

Dalam implementasinya, model dasar dilatih menggunakan data latih untuk menangkap pola hubungan antar variabel kesehatan, baik yang bersifat linier maupun non-linier. Output dari model dasar kemudian diteruskan ke model meta untuk menghasilkan keputusan klasifikasi akhir. Pendekatan ini memungkinkan proses pengambilan keputusan yang lebih adaptif, terutama pada dataset dengan karakteristik multivariabel dan distribusi kelas yang tidak seimbang, seperti data pemeriksaan kesehatan Posbindu.

Berdasarkan skema tersebut, penelitian ini menguji beberapa kombinasi model *hybrid* sebagai berikut:

- Hybrid Decision Tree + Logistic Regression (DT + LR)*
- Hybrid Random Forest + Decision Tree (RF + DT)*
- Hybrid Random Forest + Logistic Regression (RF + LR)*

Berdasarkan pendekatan stacking ensemble yang digunakan, penelitian ini menguji beberapa kombinasi model *hybrid* untuk mengevaluasi efektivitas penggabungan algoritma dengan karakteristik pembelajaran yang berbeda. Kombinasi *Decision Tree + Logistic Regression (DT + LR)* dipilih untuk mengintegrasikan keunggulan *Decision Tree* dalam membentuk aturan klasifikasi yang mudah diinterpretasikan dengan kemampuan *Logistic Regression* dalam

menghasilkan prediksi probabilistik yang stabil. Sementara itu, kombinasi Random Forest + *Decision Tree* (RF + DT) dan Random Forest + *Logistic Regression* (RF + LR) digunakan sebagai pembanding untuk menilai pengaruh model ensemble yang lebih kompleks dalam skema *hybrid*.

Pengujian terhadap ketiga kombinasi model *hybrid* tersebut bertujuan untuk menilai sejauh mana pendekatan *hybrid* mampu meningkatkan stabilitas dan kualitas klasifikasi dibandingkan algoritma tunggal, khususnya pada data kesehatan dengan distribusi kelas yang tidak seimbang. Hasil evaluasi performa dari seluruh model, baik tunggal maupun *hybrid*, selanjutnya dibandingkan menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*.

3.3 Performa Model Hybrid

Berdasarkan hasil pengujian, model *hybrid* menunjukkan performa yang kompetitif dibandingkan dengan algoritma klasifikasi tunggal. Model *Hybrid* DT + LR mampu mempertahankan tingkat akurasi yang setara dengan *Decision Tree* tunggal, dengan keseimbangan *precision* dan *recall* yang relatif baik pada seluruh kelas risiko. Hasil ini menunjukkan bahwa integrasi *Logistic Regression* sebagai meta learner tidak menurunkan performa klasifikasi, sekaligus memberikan keuntungan dari sisi stabilitas dan interpretasi probabilistik.

Model *Hybrid* RF + LR dan *Hybrid* RF + DT menunjukkan performa yang sangat mendekati, bahkan identik, dengan Random Forest tunggal. Temuan ini mengindikasikan bahwa Random Forest sebagai model dasar telah memiliki kemampuan klasifikasi yang sangat kuat dalam memodelkan hubungan kompleks antar variabel kesehatan. Oleh karena itu, penambahan model meta pada Random Forest tidak selalu menghasilkan peningkatan performa yang signifikan secara numerik, khususnya pada metrik akurasi dan F1-Score.

Hasil ini menegaskan bahwa tidak semua kombinasi *hybrid* secara otomatis menghasilkan perbaikan kinerja, dan efektivitas pendekatan *hybrid* sangat bergantung pada karakteristik data serta kekuatan model dasar yang digunakan. Dalam konteks penelitian ini, model *Hybrid* DT + LR tetap relevan karena mampu menggabungkan keunggulan *Decision Tree* dalam membentuk aturan klasifikasi yang mudah dipahami dengan kemampuan *Logistic Regression* dalam menghasilkan estimasi probabilitas risiko yang stabil dan konsisten.

Dengan demikian, meskipun beberapa model tunggal dan ensemble seperti Random Forest menunjukkan performa yang sangat tinggi, pendekatan *hybrid* DT + LR tetap memiliki nilai tambah dari sisi interpretabilitas, stabilitas prediksi, dan kesesuaian dengan tujuan penelitian, yaitu mendukung sistem deteksi dini risiko penyakit metabolik yang dapat dipahami dan dimanfaatkan oleh pemangku kepentingan non-teknis di lingkungan Posbindu.

3.4 Perbandingan Model Tunggal dan Model Hybrid

Perbandingan performa antara model klasifikasi tunggal dan model *hybrid* dilakukan untuk mengevaluasi efektivitas masing-masing pendekatan dalam mengklasifikasikan risiko penyakit metabolik berdasarkan data Posbindu. Hasil perbandingan disajikan dalam bentuk tabel ringkasan dan grafik visual pada dashboard, sehingga memungkinkan analisis performa model secara komprehensif dari berbagai metrik evaluasi.

Berdasarkan hasil pengujian, model berbasis ensemble dan *hybrid* secara umum menunjukkan performa yang lebih stabil dibandingkan model tunggal sederhana. Stabilitas ini terlihat dari keseimbangan nilai *precision* dan *recall* antar kelas risiko, terutama pada kondisi distribusi data yang tidak seimbang. Model ensemble mampu mengurangi sensitivitas terhadap variasi data karena menggabungkan beberapa proses pengambilan keputusan, sehingga menghasilkan prediksi yang lebih konsisten.

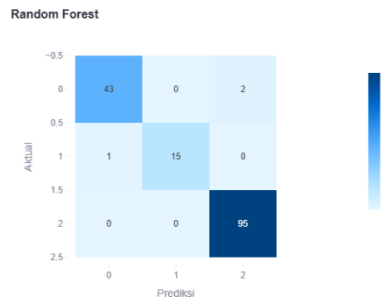
Hasil analisis juga menunjukkan bahwa model *hybrid* tidak selalu unggul secara absolut dibandingkan model tunggal terbaik. Namun demikian, pendekatan *hybrid* terbukti mampu menjaga keseimbangan performa antar kelas, khususnya dalam membedakan kelas Berisiko dan Risiko Tinggi yang memiliki karakteristik medis saling beririsan. Temuan ini mengindikasikan bahwa keunggulan model *hybrid* lebih terletak pada konsistensi dan ketahanan model terhadap ketidakseimbangan data, bukan semata-mata pada peningkatan nilai akurasi.

Di sisi lain, Random Forest tetap menunjukkan performa terbaik secara keseluruhan, baik sebagai model tunggal maupun sebagai bagian dari model *hybrid*. Hal ini menegaskan bahwa Random Forest memiliki kemampuan yang sangat kuat dalam memodelkan hubungan non-linear antar variabel kesehatan. Namun, keunggulan tersebut juga diiringi dengan kompleksitas model yang lebih tinggi dan tingkat interpretabilitas yang relatif lebih rendah dibandingkan model berbasis aturan dan probabilistik.

Secara keseluruhan, hasil perbandingan ini menunjukkan bahwa pemilihan model klasifikasi tidak hanya ditentukan oleh kompleksitas algoritma atau nilai akurasi tertinggi, tetapi juga oleh kesesuaian model dengan karakteristik data dan tujuan penggunaan. Dalam konteks sistem deteksi dini risiko penyakit metabolik di Posbindu, model *hybrid*—khususnya kombinasi *Decision Tree* dan *Logistic Regression*—memiliki keunggulan dalam hal keseimbangan performa, stabilitas prediksi, serta kemudahan interpretasi, sehingga lebih sesuai untuk mendukung pengambilan keputusan kesehatan berbasis data.

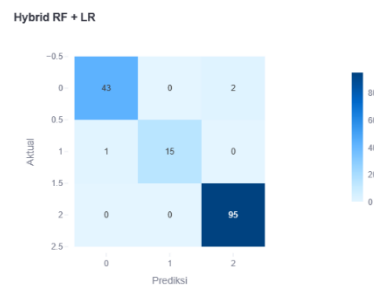
3.5 Analisis Confusion Matrix

Confusion matrix digunakan untuk menganalisis secara lebih mendalam pola kesalahan klasifikasi yang dihasilkan oleh setiap model, khususnya dalam membedakan kelas Normal, Berisiko, dan Risiko Tinggi. Melalui *confusion matrix*, dapat diketahui tidak hanya tingkat akurasi model, tetapi juga jenis kesalahan yang paling sering terjadi pada masing-masing kelas risiko. *Confusion matrix* pada model *Random Forest* Gambar 4 menunjukkan bahwa sebagian besar data berhasil diklasifikasikan dengan benar, khususnya pada kelas Risiko Tinggi yang memiliki jumlah prediksi benar paling dominan. Meskipun demikian, masih terdapat beberapa kesalahan klasifikasi, terutama pada kelas Normal yang diprediksi sebagai Risiko Tinggi serta satu data kelas Berisiko yang diprediksi sebagai Normal. Kesalahan ini mengindikasikan bahwa meskipun *Random Forest* memiliki kemampuan generalisasi yang baik, model masih mengalami kesulitan dalam membedakan individu dengan nilai indikator kesehatan yang berada di sekitar batas ambang klinis, seperti gula darah dan tekanan darah pra-hipertensi.



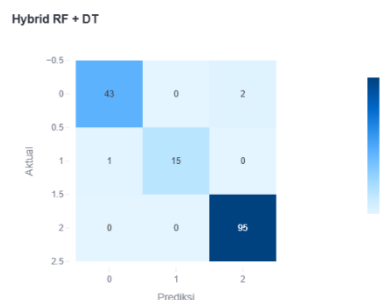
Gambar 6. Confusion Matrix Random Forest

Confusion matrix pada model Random Forest pada Gambar 6 yang dikombinasikan dengan *Logistic Regression* menunjukkan peningkatan stabilitas klasifikasi dibandingkan model tunggal. Mayoritas prediksi berada tepat pada diagonal utama, menandakan tingkat klasifikasi yang sangat baik pada seluruh kelas risiko. Kesalahan klasifikasi yang muncul relatif minimal dan terutama terjadi pada kelas Normal yang diprediksi sebagai Risiko Tinggi. Hal ini menunjukkan bahwa penambahan *Logistic Regression* membantu memperhalus keputusan akhir model, khususnya dalam menangani data borderline, sehingga mengurangi ambiguitas pada kelas risiko transisi.



Gambar 7. Confusion Matrix Random Forest + *Logistic Regression*

Model Random Forest yang dikombinasikan dengan *Decision Tree* pada gambar 7 menunjukkan performa klasifikasi yang sangat konsisten, dengan sebagian besar data terklasifikasi dengan benar pada setiap kelas risiko. Hampir seluruh data Risiko Tinggi berhasil diidentifikasi secara tepat, dan kesalahan klasifikasi yang terjadi sangat terbatas. Pola ini menunjukkan bahwa penggabungan dua algoritma berbasis pohon memperkuat kemampuan model dalam menangkap pola non-linear antar fitur kesehatan, sehingga sangat efektif digunakan untuk kebutuhan skrining risiko penyakit metabolik.



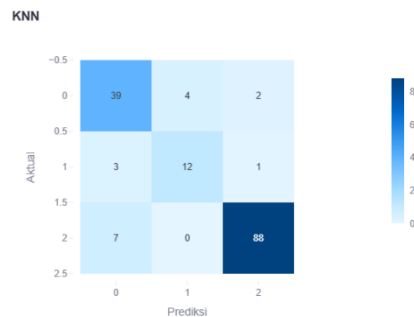
Gambar 8. Confusion Matrix Random Forest + *Decision Tree*

Confusion matrix pada model *Decision Tree* tunggal pada Gambar 8 memperlihatkan tingkat kesalahan klasifikasi yang lebih tinggi dibandingkan model ensemble. Beberapa data kelas Normal salah diklasifikasikan sebagai Berisiko dan Risiko Tinggi, serta terdapat data Risiko Tinggi yang diprediksi sebagai Normal. Kondisi ini menunjukkan keterbatasan *Decision Tree* dalam menggeneralisasi data dengan karakteristik yang saling beririsan serta kecenderungan model terhadap *overfitting*, sehingga performanya kurang stabil ketika digunakan sebagai model tunggal.



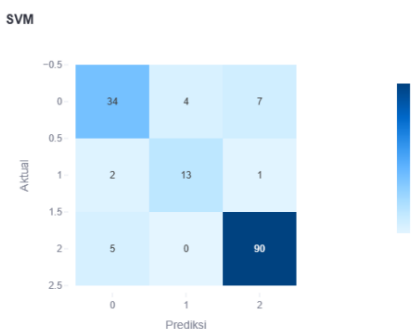
Gambar 9. Confusion Matrix *Decision Tree*

Model *K-Nearest Neighbor (KNN)* pada Gambar 9 menunjukkan pola kesalahan klasifikasi yang cukup signifikan, khususnya pada kelas Risiko Tinggi yang beberapa kali diprediksi sebagai Normal. Selain itu, kesalahan juga terjadi pada kelas Berisiko yang salah diklasifikasikan sebagai Normal. Hal ini mengindikasikan bahwa KNN kurang optimal dalam menangani data Posbindu yang memiliki perbedaan skala fitur dan karakteristik medis yang tumpang tindih, karena metode ini sangat bergantung pada kedekatan jarak antar data.



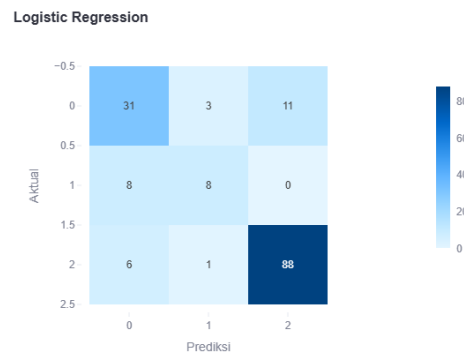
Gambar 10. Confusion Matrix KNN

Confusion matrix pada model *Support Vector Machine (SVM)* pada Gambar 10 menunjukkan bahwa meskipun model mampu mengklasifikasikan sebagian besar data dengan benar, masih terdapat kesalahan yang cukup nyata antara kelas Normal dan Risiko Tinggi. Beberapa data Normal diprediksi sebagai Risiko Tinggi dan sebaliknya, yang mengindikasikan bahwa hyperplane yang dibentuk oleh SVM belum sepenuhnya optimal dalam memisahkan kelas risiko metabolik. Hal ini disebabkan oleh tumpang tindih nilai fitur kesehatan pada kelompok risiko transisi.



Gambar 11. Confusion Matrix SVM

Model Logistic Regression pada Gambar 11 menunjukkan tingkat kesalahan klasifikasi yang relatif lebih tinggi dibandingkan model lainnya, terutama pada kelas Normal dan Berisiko. Sejumlah data Normal diprediksi sebagai Risiko Tinggi, dan beberapa data Berisiko salah diklasifikasikan sebagai Normal. Pola ini menunjukkan keterbatasan *Logistic Regression* dalam menangkap hubungan non-linear antar indikator kesehatan, sehingga model ini kurang efektif digunakan sebagai model tunggal untuk klasifikasi risiko penyakit metabolik berbasis data Posbindu.



Gambar 12. Confusion Matrix *Logistic Regression*

Berdasarkan hasil visualisasi confusion matrix, kesalahan klasifikasi paling dominan terjadi antara kelas Berisiko dan Risiko Tinggi. Kondisi ini dapat dipahami karena kedua kelas tersebut memiliki karakteristik medis yang saling beririsan, terutama pada parameter kesehatan yang berada di sekitar batas ambang klinis, seperti tekanan darah pra-hipertensi, gula darah pra-diabetes, serta nilai indeks massa tubuh dan lingkar perut yang berada pada kategori peralihan. Akibatnya, model cenderung mengalami kesulitan dalam membedakan individu yang berada pada fase transisi risiko metabolik. Sebaliknya, kelas Normal relatif lebih mudah diklasifikasikan dengan benar oleh hampir seluruh model. Hal ini disebabkan oleh nilai parameter kesehatan pada kelas Normal yang berada dalam rentang rujukan klinis yang jelas dan konsisten, sehingga membentuk pola yang lebih terpisah dibandingkan dua kelas risiko lainnya. Temuan ini menunjukkan bahwa model memiliki sensitivitas yang baik dalam mengidentifikasi individu tanpa faktor risiko metabolik yang signifikan. Model berbasis ensemble menunjukkan pola kesalahan yang lebih seimbang antar kelas dibandingkan beberapa model tunggal. Kesalahan prediksi pada model ensemble cenderung tersebar secara proporsional dan tidak terfokus pada satu kelas tertentu, yang mengindikasikan kemampuan model dalam menangkap variasi data yang lebih kompleks. Kondisi ini menjadi penting dalam konteks deteksi dini, karena kesalahan klasifikasi ekstrem misalnya mengklasifikasikan individu Risiko Tinggi sebagai Normal dapat diminimalkan. Secara keseluruhan, analisis confusion matrix memperkuat temuan bahwa pendekatan *hybrid* tidak hanya dinilai dari nilai akurasi global, tetapi juga dari kualitas distribusi kesalahan klasifikasi. Dalam konteks layanan kesehatan preventif seperti Posbindu, kemampuan model dalam membedakan kelas Berisiko dan Risiko Tinggi secara lebih hati-hati memiliki nilai klinis yang signifikan, karena berkaitan langsung dengan penentuan prioritas rujukan dan intervensi kesehatan.

3.6 Analisis ROC Curve dan AUC

Receiver Operating Characteristic (ROC) Curve digunakan untuk mengevaluasi kemampuan masing-masing model dalam membedakan kelas risiko kesehatan melalui pendekatan one-vs-rest. Pendekatan ini memungkinkan analisis performa model pada setiap kelas risiko secara terpisah, yaitu Normal, Berisiko, dan Risiko Tinggi. Evaluasi ROC dilakukan menggunakan data uji (test set) untuk menghindari bias optimistis yang dapat muncul apabila evaluasi dilakukan pada data pelatihan.

Berdasarkan visualisasi ROC Curve yang disajikan dalam bentuk grid multi-model, terlihat pada Gambar 4 bahwa *Random Forest* dan sebagian besar model *hybrid* memiliki nilai *Area Under Curve (AUC)* yang sangat tinggi pada seluruh kelas risiko, dengan kurva ROC mendekati sudut kiri atas. Hal ini menunjukkan bahwa model-model tersebut memiliki kemampuan pemisahan kelas yang sangat baik, baik dalam mendeteksi individu tanpa risiko maupun individu dengan risiko metabolik yang lebih tinggi.



Gambar 13. ROC Curve dan AUC Random Forest

Model *Hybrid RF + LR* pada Gambar 13 dan *Hybrid RF + DT* pada Gambar 4. menunjukkan pola kurva ROC yang hampir identik dengan *Random Forest* tunggal, yang mengindikasikan bahwa *Random Forest* sebagai *base learner*

telah mampu menangkap struktur non-linear data secara optimal. Oleh karena itu, penambahan meta learner pada model berbasis Random Forest tidak memberikan peningkatan AUC yang signifikan. Temuan ini memperkuat hasil pada subbab sebelumnya bahwa model *hybrid* tidak selalu menghasilkan peningkatan performa numerik, terutama ketika model dasar sudah sangat kuat.

Hybrid RF + LR



(a) ROC Curve dan AUC Hybrid RF + LR

Hybrid RF + DT



(b) ROC Curve dan AUC Hybrid RF + DT

Gambar 14. (a) ROC Curve dan AUC Hybrid RF + LR dan (b) ROC Curve dan AUC Hybrid RF + DT

Pada model *Decision Tree* pada Gambar 14 dan *Hybrid DT + LR* nilai AUC tetap berada pada kategori tinggi untuk seluruh kelas risiko, meskipun sedikit lebih rendah dibandingkan Random Forest. Namun demikian, pola kurva ROC pada model *hybrid DT + LR* pada (b) menunjukkan peningkatan stabilitas pemisahan kelas dibandingkan *Decision Tree* tunggal, khususnya pada kelas Berisiko dan Risiko Tinggi. Hal ini mengindikasikan bahwa *Logistic Regression* sebagai meta learner berperan dalam memperhalus keputusan klasifikasi melalui estimasi probabilistik, terutama pada area ambang batas kelas.

Hybrid DT + LR



(a) ROC Curve dan AUC Hybrid DT + LR

Decision Tree



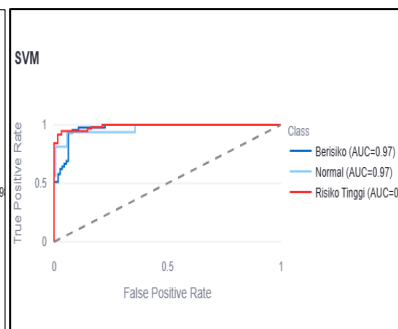
(b) ROC Curve dan AUC Decision Tree

Gambar 15. (a) ROC Curve dan AUC Hybrid DT + LR dan (b) ROC Curve dan AUC Decision Tree

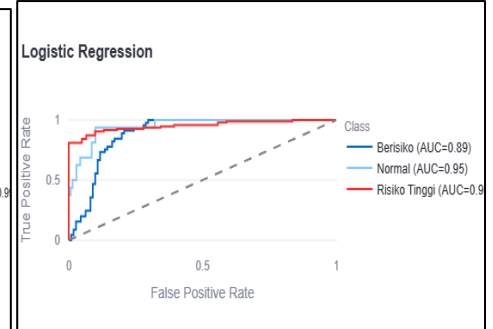
Sebaliknya, model KNN pada Gambar 15, menunjukkan nilai AUC yang relatif lebih rendah pada beberapa kelas, khususnya pada kelas Berisiko. Kurva ROC yang lebih landai pada bagian awal grafik mencerminkan sensitivitas model yang lebih terbatas dalam membedakan kelas risiko menengah, yang secara klinis memang memiliki karakteristik yang saling beririsan. Temuan ini sejalan dengan hasil confusion matrix, di mana kesalahan klasifikasi paling sering terjadi pada transisi antara kelas Berisiko dan Risiko Tinggi.



(a) ROC Curve dan AUC KNN



(b) ROC Curve dan AUC SVM



(c) ROC Curve dan AUC Logistic Regression

Gambar 16. (a) ROC Curve dan AUC KNN, (b) ROC Curve dan AUC SVM dan (c) ROC Curve dan AUC Logistic Regression

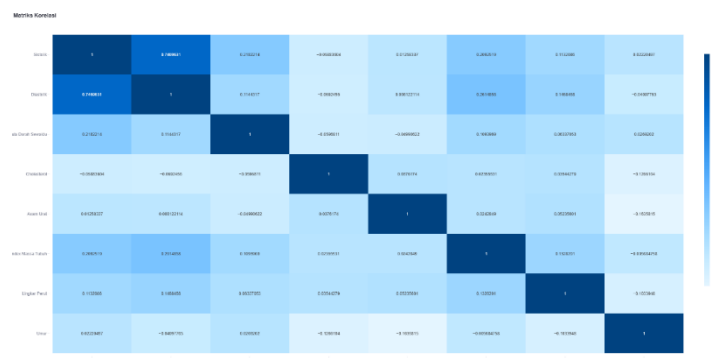
Secara keseluruhan, analisis ROC Curve dan AUC menunjukkan bahwa nilai AUC yang tinggi menandakan kemampuan pemisahan kelas yang baik, namun tidak dapat dijadikan satu-satunya dasar dalam pemilihan model. Dalam konteks sistem deteksi dini risiko penyakit metabolik, aspek keseimbangan sensitivitas antar kelas, stabilitas prediksi, serta kemudahan interpretasi hasil menjadi pertimbangan yang sama pentingnya. Oleh karena itu, meskipun Random Forest menunjukkan performa AUC tertinggi, model *hybrid Decision Tree + Logistic Regression* tetap relevan karena mampu memberikan performa yang konsisten dengan interpretasi yang lebih mudah diterjemahkan ke dalam pengambilan keputusan kesehatan preventif di lingkungan Posbindu.

3.7 Analisis Fitur Dominan dalam Klasifikasi Risiko

Analisis fitur dominan dilakukan untuk mengidentifikasi variabel kesehatan yang paling berpengaruh dalam proses klasifikasi risiko penyakit metabolik. Identifikasi ini didasarkan pada struktur model klasifikasi, kontribusi fitur terhadap keputusan model, serta analisis hubungan antar variabel menggunakan matriks korelasi. Pendekatan ini bertujuan untuk memastikan bahwa model yang dikembangkan tidak hanya unggul secara statistik, tetapi juga selaras dengan pengetahuan medis dan praktik klinis.

Berdasarkan hasil pemodelan, beberapa fitur terbukti memiliki pengaruh dominan terhadap penentuan kelas risiko, yaitu Indeks Massa Tubuh (IMT), gula darah sewaktu, tekanan darah (sistolik dan diastolik), serta lingkaran perut. Temuan ini sejalan dengan penelitian terkait faktor risiko sindrom metabolik dan pemodelan risiko kesehatan [1], [2]. Fitur-fitur tersebut secara konsisten muncul sebagai variabel pembeda utama dalam model berbasis pohon keputusan maupun model *hybrid*, khususnya dalam membedakan kelas Berisiko dan Risiko Tinggi. Temuan ini menunjukkan bahwa indikator antropometri dan metabolik memiliki peran penting dalam menentukan tingkat risiko penyakit metabolik.

Analisis matriks korelasi antar fitur, sebagaimana ditunjukkan pada Gambar 16, memperlihatkan adanya hubungan yang cukup kuat antara tekanan darah sistolik dan diastolik, yang mencerminkan konsistensi fisiologis antar parameter tekanan darah. Selain itu, IMT dan lingkaran perut menunjukkan korelasi positif dengan beberapa parameter metabolik lainnya, meskipun tidak selalu dalam tingkat korelasi yang tinggi. Kondisi ini mengindikasikan bahwa risiko penyakit metabolik dipengaruhi oleh kombinasi multi-parameter, bukan oleh satu variabel tunggal.

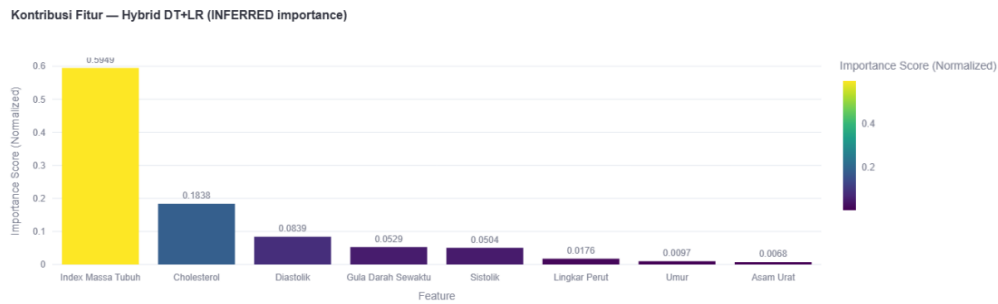


Gambar 17. Korelasi Fitur yang berpengaruh

Meskipun sebagian besar nilai korelasi antar fitur berada pada kategori rendah hingga sedang, model klasifikasi tetap mampu memanfaatkan interaksi non-linear antar variabel, khususnya pada algoritma berbasis pohon dan ensemble. Hal ini menjelaskan mengapa model seperti Random Forest dan model *hybrid* mampu mencapai performa yang tinggi, meskipun hubungan linier antar fitur tidak selalu kuat. Dengan kata lain, rendahnya korelasi linier tidak serta-merta menurunkan kemampuan model dalam mengidentifikasi pola risiko yang kompleks.

Kesesuaian antara fitur dominan yang diidentifikasi oleh model dengan teori medis menunjukkan bahwa model yang dikembangkan memiliki validitas klinis. Parameter seperti IMT, lingkaran perut, tekanan darah, dan kadar gula darah merupakan komponen utama dalam penilaian risiko sindrom metabolik menurut berbagai pedoman kesehatan. Oleh karena itu, hasil penelitian ini memperkuat bahwa pendekatan data mining yang digunakan mampu merepresentasikan kondisi kesehatan secara realistis dan dapat mendukung proses deteksi dini risiko penyakit metabolik berbasis data Posbindu.

Analisis feature importance pada model *Hybrid Decision Tree + Logistic Regression (DT + LR)* dilakukan untuk memahami kontribusi relatif setiap fitur kesehatan terhadap keputusan klasifikasi risiko penyakit metabolik. Karena model *hybrid* ini dikembangkan menggunakan pendekatan stacking, di mana model meta tidak secara langsung mempelajari fitur asli, maka feature importance pada model *hybrid* tidak dihitung secara langsung, melainkan diturunkan (inferred) dari kontribusi fitur pada model dasar.



Gambar 18. Analisis Inferred Feature Importance pada Model *Hybrid Decision Tree + Logistic Regression*

Pendekatan inferred feature importance dilakukan dengan mengintegrasikan nilai feature importance dari *Decision Tree* dan bobot koefisien dari *Logistic Regression* yang telah dinormalisasi. Dengan cara ini, kontribusi fitur pada model *hybrid* mencerminkan peran fitur-fitur tersebut dalam membentuk keputusan akhir melalui kombinasi aturan pohon keputusan dan estimasi probabilistik.

Berdasarkan hasil analisis yang ditunjukkan pada Gambar 22, Indeks Massa Tubuh (IMT) merupakan fitur dengan kontribusi paling dominan dalam model *hybrid*, dengan nilai importance ternormalisasi tertinggi. Hal ini menunjukkan bahwa IMT berperan sebagai indikator utama dalam membedakan tingkat risiko penyakit metabolik, khususnya dalam pemisahan awal antara kelas Normal, Berisiko, dan Risiko Tinggi.

Fitur kolesterol, tekanan darah diastolik, serta gula darah sewaktu menempati urutan berikutnya dengan kontribusi yang lebih rendah namun tetap signifikan. Fitur-fitur ini berfungsi sebagai faktor pendukung yang memperhalus keputusan klasifikasi, terutama pada individu yang berada pada zona transisi risiko metabolik. Sementara itu, tekanan darah sistolik memiliki kontribusi yang relatif moderat, yang menunjukkan perannya dalam memperkuat klasifikasi pada kondisi tertentu. Sebaliknya, fitur lingkar perut, usia, dan asam urat menunjukkan nilai importance yang relatif rendah dalam model *ensemble*. Hal ini tidak mengindikasikan bahwa fitur-fitur tersebut tidak relevan secara klinis, melainkan menunjukkan bahwa kontribusi informasinya telah terwakili oleh fitur lain yang lebih dominan, khususnya IMT dan parameter metabolik utama yang disajikan pada Tabel 1. Selain itu, rendahnya kontribusi fitur-fitur tersebut juga dapat dipengaruhi oleh korelasi antar variabel serta struktur keputusan yang terbentuk pada model dasar.

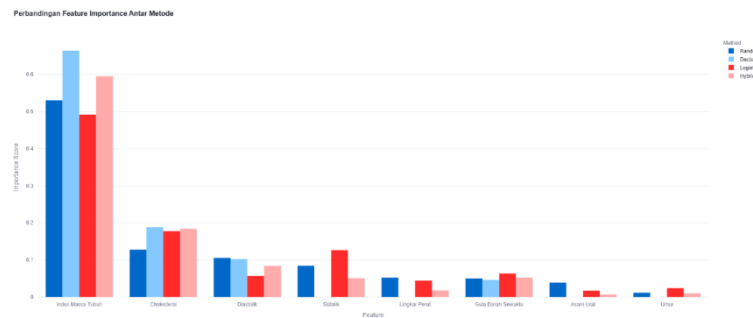
Tabel 1. Analisis Inferred Feature Importance pada Model *Hybrid Decision Tree + Logistic Regression*

Feature	Importance
Index Massa Tubuh	0.594877
Kolesterol	0.18383
Diastolik	0.083896
Gula Darah Sewaktu	0.052859
Sistolik	0.050416
Lingkar Perut	0.017623
Umur	0.009666
Asam Urat	0.006833

Secara keseluruhan, hasil inferred feature importance pada model *hybrid* menunjukkan bahwa keputusan klasifikasi risiko penyakit metabolik tetap didasarkan pada indikator kesehatan utama yang diakui secara medis. Kesesuaian antara fitur dominan hasil pemodelan dan teori medis memperkuat bahwa model *hybrid* yang dikembangkan tidak hanya memiliki performa yang baik secara statistik, tetapi juga relevan dan dapat diinterpretasikan secara klinis, sehingga mendukung penerapannya dalam sistem deteksi dini risiko penyakit metabolik berbasis data Posbindu.

Selain analisis fitur dominan pada masing-masing model, penelitian ini juga melakukan perbandingan feature importance antar metode klasifikasi untuk mengidentifikasi konsistensi kontribusi fitur pada berbagai pendekatan pemodelan. Perbandingan ini mencakup Random Forest, *Decision Tree*, *Logistic Regression*, serta model *hybrid Decision Tree + Logistic Regression (DT + LR)* dengan nilai importance yang telah dinormalisasi.

Berdasarkan visualisasi perbandingan yang ditunjukkan pada Gambar 3.19, Indeks Massa Tubuh (IMT) secara konsisten muncul sebagai fitur dengan kontribusi paling dominan pada seluruh metode yang digunakan. Dominasi IMT pada berbagai algoritma menunjukkan bahwa indikator antropometri ini merupakan faktor utama dalam penentuan risiko penyakit metabolik, terlepas dari perbedaan mekanisme pembelajaran model. Temuan ini memperkuat validitas hasil pemodelan sekaligus menunjukkan keselarasan dengan teori medis terkait sindrom metabolik.



Gambar 19. Perbandingan Feature Importance Antar Metode Klasifikasi

Fitur kolesterol dan tekanan darah diastolik juga menunjukkan kontribusi yang relatif stabil pada sebagian besar metode, meskipun dengan tingkat kepentingan yang bervariasi. Variasi ini mencerminkan perbedaan cara masing-masing algoritma dalam memanfaatkan informasi fitur, di mana model berbasis pohon cenderung menekankan pemisahan berbasis ambang nilai, sedangkan *Logistic Regression* memanfaatkan hubungan linier antar variabel.

Sebaliknya, fitur seperti lingkaran perut, asam urat, dan usia menunjukkan nilai importance yang lebih rendah dan cenderung bervariasi antar metode. Hal ini mengindikasikan bahwa kontribusi fitur-fitur tersebut bersifat kontekstual dan sangat bergantung pada kombinasi fitur lain yang digunakan dalam model. Pada model *hybrid* DT + LR, fitur-fitur dengan importance rendah tetap berkontribusi secara tidak langsung melalui proses penggabungan keputusan dari model dasar, meskipun tidak menjadi penentu utama klasifikasi.

Perbandingan ini menunjukkan bahwa meskipun nilai importance absolut antar metode berbeda, pola dominasi fitur utama relatif konsisten. Konsistensi tersebut menegaskan bahwa keputusan klasifikasi risiko penyakit metabolik dalam penelitian ini tidak bergantung pada satu algoritma tertentu, melainkan pada indikator kesehatan yang secara medis memang relevan. Dengan demikian, model *hybrid* yang dikembangkan mampu mempertahankan interpretabilitas dan validitas klinis melalui pemanfaatan fitur-fitur dominan yang sama dengan model tunggal dan ensemble.

3.8 Pembahasan

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa penerapan pendekatan data mining dan machine learning pada data pemeriksaan kesehatan Posbindu mampu memberikan kontribusi yang signifikan dalam mendukung deteksi dini risiko penyakit metabolik. Melalui tahapan *Knowledge Discovery in Databases* (KDD) yang diterapkan secara sistematis, penelitian ini berhasil mengolah data kesehatan mentah menjadi informasi yang bermakna untuk proses klasifikasi risiko kesehatan.

Hasil evaluasi menunjukkan bahwa algoritma klasifikasi memiliki kemampuan yang berbeda-beda dalam memodelkan karakteristik data Posbindu yang bersifat multivariabel dan tidak seimbang. Model berbasis ensemble dan *hybrid* secara umum menunjukkan performa yang lebih stabil dibandingkan model tunggal sederhana, terutama dalam menjaga keseimbangan precision dan recall antar kelas risiko. Temuan ini mengindikasikan bahwa pendekatan kombinasi model dapat membantu mengurangi bias terhadap kelas mayoritas dan meningkatkan keandalan prediksi pada kelas risiko yang lebih kritis.

Model *hybrid Decision Tree + Logistic Regression* yang dikembangkan dalam penelitian ini memberikan alternatif pendekatan yang seimbang antara performa klasifikasi, stabilitas prediksi, dan interpretabilitas model. Meskipun model ensemble seperti Random Forest menunjukkan performa numerik yang sangat tinggi, model *hybrid* DT + LR tetap relevan karena mampu menghasilkan keputusan klasifikasi yang lebih mudah dipahami dan dijelaskan, khususnya dalam konteks layanan kesehatan preventif yang melibatkan pemangku kepentingan non-teknis.

Analisis confusion matrix, ROC Curve, serta feature importance menunjukkan bahwa kesalahan klasifikasi paling sering terjadi pada kelas Berisiko dan Risiko Tinggi, yang secara klinis memang memiliki karakteristik yang saling beririsan. Hal ini menegaskan bahwa klasifikasi risiko penyakit metabolik merupakan permasalahan yang kompleks dan tidak dapat direduksi menjadi keputusan biner sederhana. Oleh karena itu, kemampuan model dalam menjaga kehati-hatian pada zona transisi risiko menjadi aspek penting dalam mendukung pengambilan keputusan kesehatan.

Pemanfaatan dashboard interaktif sebagai media visualisasi hasil penelitian memberikan nilai tambah yang signifikan dalam proses interpretasi dan komunikasi hasil. Visualisasi performa model, distribusi risiko, serta kontribusi fitur kesehatan memungkinkan hasil penelitian dipahami tidak hanya oleh peneliti, tetapi juga oleh tenaga kesehatan dan pengambil kebijakan. Dengan pendekatan ini, hasil pemodelan tidak berhenti pada aspek teknis, tetapi dapat ditransformasikan menjadi informasi yang aplikatif dalam praktik layanan kesehatan.

Dengan demikian, penelitian ini menunjukkan bahwa integrasi data Posbindu dengan pendekatan machine learning dan visualisasi interaktif memiliki potensi untuk dikembangkan lebih lanjut sebagai sistem pendukung keputusan (Decision Support System) dalam layanan kesehatan preventif. Pemanfaatan Posbindu sebagai sarana deteksi dini faktor risiko penyakit tidak menular juga didukung oleh penelitian sebelumnya [16], [17].

Pengembangan lanjutan dapat diarahkan pada pemanfaatan data longitudinal, integrasi dengan sistem informasi kesehatan, serta penyempurnaan model prediksi untuk mendukung upaya pencegahan penyakit metabolik secara berkelanjutan.

Berdasarkan hasil perbandingan antara pelabelan risiko kesehatan yang dilakukan secara manual oleh pemeriksa dengan hasil klasifikasi menggunakan sistem yang dikembangkan, dari total 938 data yang diberi label secara manual, terdapat 861 data yang menunjukkan kesesuaian dengan hasil prediksi sistem. Hasil ini merepresentasikan tingkat kecocokan sebesar sekitar 91,8%, yang menunjukkan bahwa sistem mampu mereplikasi keputusan pemeriksa dengan tingkat akurasi yang tinggi.

Perbedaan klasifikasi pada sebagian kecil data yang tersisa dapat dijelaskan melalui analisis confusion matrix, di mana kesalahan paling dominan terjadi pada kelas Berisiko dan Risiko Tinggi, khususnya pada data dengan nilai indikator kesehatan yang berada di sekitar batas ambang klinis, seperti tekanan darah pra-hipertensi dan gula darah pra-diabetes. Pola confusion matrix pada model ensemble dan hybrid menunjukkan mayoritas prediksi berada pada diagonal utama, menandakan stabilitas dan konsistensi klasifikasi yang baik serta distribusi kesalahan yang lebih seimbang dibandingkan model tunggal. Dengan demikian, ketidaksesuaian yang terjadi dapat dipandang sebagai fenomena yang wajar secara klinis dan tidak mengurangi validitas sistem, melainkan menegaskan bahwa algoritma telah mampu mendukung proses penilaian risiko kesehatan secara objektif dan sistematis.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa model *ensemble (stacking)* yang menggabungkan Decision Tree sebagai *base learner* dan Logistic Regression sebagai *meta-learner* berhasil dikembangkan untuk mengklasifikasikan risiko penyakit metabolik berdasarkan data Posbindu Universitas Gadjah Mada ke dalam tiga kategori, yaitu Normal, Berisiko, dan Risiko Tinggi. Hasil evaluasi menunjukkan bahwa model memiliki performa yang baik dengan akurasi sebesar 94,87% dan F1-score sebesar 0,9187, serta mampu menghasilkan prediksi yang lebih stabil dibandingkan model tunggal, terutama pada data dengan distribusi kelas yang tidak seimbang. Analisis *confusion matrix*, ROC-AUC, dan *feature importance* juga menunjukkan bahwa model mampu meminimalkan kesalahan klasifikasi pada kasus-kasus kritis, sementara Indeks Massa Tubuh (IMT), kolesterol, tekanan darah, dan gula darah merupakan faktor yang paling berpengaruh dalam menentukan risiko penyakit metabolik. Secara keseluruhan, penerapan kerangka kerja Knowledge Discovery in Databases (KDD) dan metode *stacking ensemble* terbukti efektif dalam menghasilkan model klasifikasi yang akurat, stabil, dan relevan untuk mendukung deteksi dini penyakit metabolik di lingkungan Universitas Gadjah Mada. Berdasarkan hasil tersebut, penelitian selanjutnya disarankan menggunakan dataset yang lebih besar dan lebih beragam, serta menambahkan variabel seperti aktivitas fisik, pola makan, kebiasaan merokok, dan riwayat penyakit keluarga agar kemampuan prediksi model semakin baik. Selain itu, penerapan metode Explainable Artificial Intelligence (XAI), seperti SHAP atau LIME, dapat dipertimbangkan untuk meningkatkan transparansi hasil klasifikasi sehingga lebih mudah dipahami oleh tenaga kesehatan. Model yang telah dikembangkan juga berpotensi diimplementasikan ke dalam sistem informasi Posbindu dalam bentuk dashboard interaktif atau modul prediksi otomatis guna mendukung proses deteksi dini, penentuan prioritas intervensi, dan pengambilan keputusan berbasis data, sehingga dapat mendukung pelaksanaan program Health Promoting University (HPU) secara lebih efektif dan berkelanjutan.

UCAPAN TERIMAKASIH

Terima kasih disampaikan kepada pihak-pihak yang telah mendukung terlaksananya penelitian ini.

REFERENCES

- [1] Y. Zhang et al., "Construction of Xinjiang metabolic syndrome risk prediction model based on interpretable models," BMC Public Health, vol. 22, no. 1, pp. 1–11, 2022, doi: 10.1186/s12889-022-12617-y.
- [2] M. J. Zhou, M. S. Chen, T. S. Lee, C. T. Yang, Y. L. Chiu, and C. J. Lu, "A hybrid risk factor evaluation scheme for metabolic syndrome and stage 3 chronic kidney disease based on multiple machine learning techniques," Healthcare, vol. 10, no. 12, 2022, doi: 10.3390/healthcare10122496.
- [3] Universitas Gadjah Mada, "Usaha Cegah Penyakit Tak Menular Lewat Posbindu-PTM," 2024. [Online]. Available: <https://lib.ugm.ac.id/usaha-cegah-penyakit-tak-menular-lewat-posbindu-ptm/>
- [4] Kementerian Kesehatan Republik Indonesia, "Peraturan Menteri Kesehatan Republik Indonesia Nomor 71 Tahun 2015 tentang Penanggulangan Penyakit Tidak Menular," 2015. [Online]. Available: <https://p2ptm.kemkes.go.id/dokumen-ptm/peraturan-menteri-kesehatan-republik-indonesia-nomor-71-tahun-2015-tentang-penanggulangan-penyakit-tidak-menular>
- [5] E. O. Paul, "Hybrid decision tree-based machine learning models for diabetes prediction," SCIREA Journal of Information Science and Systems Science, Jan. 2024, doi: 10.54647/iss120327.
- [6] A. Arista, "Comparison Decision Tree and Logistic Regression machine learning classification algorithms to determine Covid-19," Sinkron, vol. 7, no. 1, pp. 59–65, 2022, doi: 10.33395/sinkron.v7i1.11243.

- [7] S. J. Lin, C. C. Liu, D. M. T. Tsai, Y. H. Shih, C. L. Lin, and Y. C. Hsu, "Prediction models using Decision Tree and Logistic Regression method for predicting hospital revisits in peritoneal dialysis patients," *Diagnostics*, vol. 14, no. 6, pp. 1–13, 2024, doi: 10.3390/diagnostics14060620.
- [8] Biro Pelayanan Kesehatan Terpadu, Universitas Gadjah Mada, "Health Promoting University di Universitas Gadjah Mada," 2024. [Online]. Available: <https://hpu.ugm.ac.id/>
- [9] Z. Xiaoshuai, T. Chuanping, and W. Shuohuan, "A stacking ensemble model for predicting the occurrence of carotid atherosclerosis," *Frontiers in Endocrinology*, vol. 15, pp. 1–8, 2024, doi: 10.3389/fendo.2024.1390352.
- [10] B. J. Bhatt et al., "Validation of a popular consumer-grade cuffless blood pressure device for continuous 24 h monitoring," *European Heart Journal - Digital Health*, pp. 1–9, 2025, doi: 10.1093/ehjdh/ztaf044.
- [11] M. Muniyandi et al., "Diagnostic accuracy of mercurial versus digital blood pressure measurement devices: a systematic review and meta-analysis," *Scientific Reports*, vol. 12, no. 1, pp. 1–8, 2022, doi: 10.1038/s41598-022-07315-z.
- [12] S. Stabouli et al., "Comparison of validation protocols for blood pressure measuring devices in children and adolescents," *Frontiers in Cardiovascular Medicine*, vol. 9, 2022, doi: 10.3389/fcvm.2022.1001878.
- [13] W. Widayat, P. Assiroj, Sohirin, I. A. Prabadhi, and P. A. Kautsar, "Data mining implementation: a survey," *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 36, no. 3, pp. 1960–1968, 2024, doi: 10.11591/ijeecs.v36.i3.pp1960-1968.
- [14] A. H. Nasrullah, "Implementasi Algoritma Decision Tree untuk klasifikasi produk laris," *Jurnal Ilmiah Ilmu Komputer*, vol. 7, no. 2, pp. 45–51, 2021, doi: 10.35329/jiik.v7i2.203.
- [15] L. R. Safitri, N. Chamidah, T. Saifudin, M. Firmansyah, and G. T. Alpandi, "Comparison of Logistic Regression and Support Vector Machine in predicting stroke risk," *Inferensi*, vol. 7, no. 2, p. 123, 2024, doi: 10.12962/j27213862.v7i2.20420.
- [16] T. Siswati, H. S. Kasjono, and Y. Olfah, "Posbindu PTM: The key of early detection and decreasing prevalence of non-communicable diseases in Indonesia," *Iranian Journal of Public Health*, vol. 51, no. 7, pp. 1683–1684, 2022, doi: 10.18502/ijph.v51i7.10105.
- [17] V. Widyarningsih et al., "Missed opportunities in hypertension risk factors screening in Indonesia: A mixed-methods evaluation of integrated health post (POSBINDU) implementation," *BMJ Open*, vol. 12, no. 2, pp. 1–11, 2022, doi: 10.1136/bmjopen-2021-051315.