

Perbandingan Optimasi Algoritma Klasifikasi *Decision Tree*, *Naive Bayes* dan *KNN* Menggunakan *Optimize Parameter Grid* Pada Tingkat Resiko Ibu Hamil

Aa Kurniawan

Program Studi Teknik Informatika, Universitas Pamulang

Article Info

Article history:

Received Feb 17, 2025

Revised Sept 09, 2025

Accepted Sept 15, 2025

Keywords:

classification

optimize parameter grid

decision tree

naïve bayes

knn

ABSTRAK

The high Maternal Mortality Rate (MMR) in Indonesia (2023) highlights the critical need for accurate maternal risk classification to enable timely healthcare interventions. This research compares the performance of four classification algorithms Decision Tree ID3, Decision Tree C4.5, Naïve Bayes, and KNN on the UCI "Maternal Health Risk" dataset. Each model's performance was evaluated before and after hyperparameter tuning via the Grid Search method. Prior to optimization, KNN yielded the highest accuracy (82.62%), while Naïve Bayes was the least effective (61.31%). Post-optimization, a significant shift occurred: both Decision Tree ID3 and C4.5 emerged as the top-performing models with an identical accuracy of 85.65%. Notably, Naïve Bayes showed no improvement, whereas Decision Tree C4.5 exhibited the most substantial accuracy gain of 10.82%. These findings suggest that optimized Decision Tree models provide the most robust approach for this maternal health risk classification task.

Corresponding Author:

Teknik Informatika Fakultas Ilmu Komputer

Universitas Pamulang

Jl. Surya Kencana No.1, Pamulang, Kota Tangerang Selatan

Email: dosen02361@unpam.ac.id

1. PENDAHULUAN

Menurut data yang dirilis oleh Kementerian Kesehatan Republik Indonesia pada Januari 2023, Angka Kematian Ibu (AKI) masih berada di angka 305 dari 100.000 kelahiran hidup. Angka ini masih jauh dari target yang ditetapkan, yaitu 183 dari 100.000 kelahiran hidup pada tahun 2024 [1]. AKI menjadi indikator utama dalam menilai keberhasilan pembangunan kesehatan ibu.

Upaya peningkatan kesehatan ibu masih menghadapi berbagai kendala, terutama dalam mengelola faktor risiko bagi ibu serta janin. Edukasi pemahaman mengenai gizi yang tepat, pola hidup sehat dan deteksi dini terhadap kondisi medis yang dapat mempengaruhi kehamilan masih menjadi tantangan di lapangan. Oleh sebab itu, diperlukan kerja sama antara tenaga kesehatan, masyarakat dan pemerintah dalam menyediakan akses yang lebih baik terhadap informasi serta layanan kesehatan bagi ibu hamil. Langkah ini menjadi bagian penting dalam strategi pencegahan guna meningkatkan kualitas kesehatan ibu dan janin.

Perawatan yang tepat dan dilakukan secara tepat waktu bagi ibu hamil sangat penting untuk mengurangi risiko komplikasi serta memastikan kelahiran yang sehat. Namun faktanya, klasifikasi dalam menentukan tingkat risiko kesehatan ibu hamil sering kali menjadi tantangan bagi tenaga medis. Sistem yang digunakan saat ini masih bersifat subjektif, sehingga hasil penilaian dapat bervariasi dan kurang konsisten [2].

Dalam meningkatkan pemahaman terhadap faktor risiko yang mempengaruhi kesehatan ibu hamil, perlu dilakukan identifikasi variabel data yang dapat digunakan untuk mengembangkan model klasifikasi tingkat risiko. Selanjutnya, pendekatan *machine learning* dapat diterapkan dalam analisis guna meningkatkan akurasi dan objektivitas klasifikasi risiko tersebut [3], [4], [5].

Dengan metode berbasis data yang lebih objektif, algoritma *machine learning* mampu mengidentifikasi faktor risiko dan pola yang berkontribusi terhadap kesehatan ibu dan janin dengan lebih

akurat dan konsisten. Penelitian ini bertujuan untuk mengembangkan metode klasifikasi risiko kesehatan ibu hamil yang lebih andal dan efisien, sehingga dapat berkontribusi dalam menurunkan angka kematian ibu dan bayi.

Di era industri 4.0, perkembangan teknologi informasi dan komunikasi dapat dimanfaatkan untuk memprediksi tingkat kesehatan ibu hamil. Salah satu metode yang dapat diterapkan adalah penggunaan algoritma klasifikasi teknik pemanfaatannya dengan menggunakan suatu algoritma klasifikasi dalam *data mining*. Perbandingan tingkat akurasi dari berbagai algoritma menjadi indikator penting dalam menentukan metode terbaik yang dapat digunakan untuk implementasi dan pengembangan selanjutnya.

Teknik klasifikasi adalah bidang pembahasan yang sering digunakan oleh beberapa peneliti sebelumnya dalam mengekstrak pengetahuan. Dari beberapa literatur terdapat beberapa penelitian terkait dengan klasifikasi prediksi tingkat resiko ibu hamil. Penelitian pertama, menggunakan metode *Decision Tree* dalam melakukan teknik klasifikasi dengan data uji dari UCI dataset yang diimplementasikan dalam aplikasi, menghasilkan tingkat akurasi sebesar 61.54% [6].

Studi lain, membandingkan algoritma *machine learning* dalam klasifikasi ibu hamil dilakukan di *Google Colaboratory* dengan Bahasa pemrograman *Phyton*. Algoritma yang digunakan meliputi *Naïve Bayes*, *Decision Tree*, dan *K-Nearest Neighbors (KNN)* dengan data dari *Uci dataset*. Hasil penelitian menunjukkan bahwa *Decision Tree* memiliki tingkat akurasi tertinggi sebesar 90% mengungguli *KNN* (86%) dan *Naïve Bayes* (65%) [7].

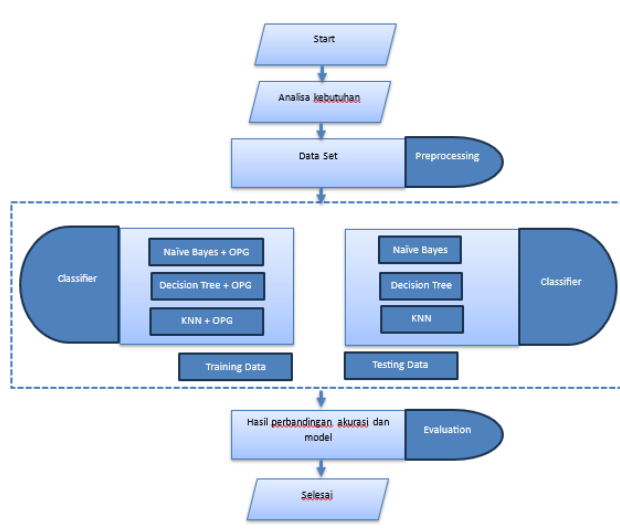
Penelitian lainnya membandingkan beberapa algoritma klasifikasi menggunakan *Hyperparameter Tuning* yaitu *Random Forest*, *Extra Trees*, *Extreme Gradient Boosting*, *Decision Tree* dan *Light Gradient Boosting Machine*. Hasilnya cukup mengejutkan bahwa sebelum optimalisasi, algoritma *Random Forest* memiliki akurasi tertinggi sebesar 82,15%. Namun, setelah dilakukan optimalisasi, *Extreme Gradient Boosting* menunjukkan akurasi terbaik sebesar 79,03% [8].

Secara umum, algoritma *machine learning* tidak akan mencapai hasil optimal jika parameter yang digunakan tidak sesuai. Oleh karena itu, dalam membangun model klasifikasi dengan tingkat akurasi tinggi, penting untuk memilih algoritma yang tepat serta melakukan optimalisasi parameter. Optimisasi parameter secara manual membutuhkan waktu yang lama, terutama jika algoritma yang digunakan memiliki banyak parameter [9]. Oleh karena itu, dalam penelitian ini, digunakan metode *Optimize Parameter Grid* untuk meningkatkan akurasi model

Dengan memanfaatkan dataset yang informatif dari *UCI dataset*, penelitian ini berupaya menganalisis penerapan model klasifikasi guna mengidentifikasi pola risiko kesehatan ibu dan anak, diharapkan penelitian ini dapat memberikan wawasan baru dalam edukasi terkait risiko kesehatan ibu hamil serta meningkatkan kesadaran akan pentingnya perawatan dan manajemen kehamilan yang sehat, baik bagi ibu maupun calon bayi [10], [11].

2. METODE

Penelitian ini menggunakan metode komparatif yang terdiri dari empat tahapan utama, yaitu analisis kebutuhan, preprocessing data, komparasi klasifikasi data, dan evaluasi. Alur penelitian ditunjukkan pada Gambar 1.



Gambar 1. Tahapan-tahapan penelitian

2.1 Analisis Kebutuhan

Data mining merupakan teknik pengolahan data data skala besar yang digunakan untuk menganalisis hubungan antar *dataset*, yang kemudian diinterpretasikan dalam bentuk yang lebih mudah difahami dan bermanfaat bagi pengguna [12]. *Data mining* mencakup berbagai disiplin ilmu, seperti basis data, *machine learning*, teknik visualisasi, pengenalan pola, serta statistik, untuk memperoleh informasi dari data yang besar [13].

Dalam penelitian ini, data latih dan data uji diperoleh dari *Uci repository* dengan tema *maternal health risk* yang terdiri dari 1014 data. Dataset ini memiliki enam atribut independen dan satu atribut dependen. Rincian atribut yang digunakan dapat dilihat pada tabel.1 berikut.

Tabel 1. Keterangan Atribut

No.	Atribut	Keterangan
1.	Age	Usia Calon Penderita
2.	Systolic BP	Angka Tekanan darah maksimum
3.	Diastolic BP	Angka Tekanan Darah Minimum
4.	BS	Konsentrasi Glukosa Darah
5.	BodyTemp	Suhu Tubuh Calon Penderita
6.	Heart Rate	Denyut Jantung
7.	Riskleve	Prediksi Tingkat Risiko Ibu Hamil

Sumber : Uci repository

Berdasarkan sumber data yang ada, tidak ditemukan nilai kosong (*missing values*), sehingga data dapat langsung dianalisis lebih lanjut. Sebelum memasuki tahapan preprocessing ada beberapa hal yang perlu untuk difahami terkait karakteristik data, berikut penjelasannya:

a. Analisis Statistik Deskriptif

- Usia ibu (*Age*) berkisar antara 10 hingga 70 tahun dengan rata-rata 29,8 tahun. Mayoritas usia berada pada rentang usia produktif (20–35 tahun), namun terdapat nilai ekstrem di atas 45 tahun .
- Tekanan darah sistolik (*SystolicBP*) memiliki nilai rata-rata 113 mmHg, dengan rentang 70–160 mmHg. Nilai di atas 140 mmHg dapat mengindikasikan hipertensi, yang merupakan faktor risiko tinggi pada kehamilan.
- Tekanan darah diastolik (*DiastolicBP*) rata-rata 76 mmHg dengan rentang 49–100 mmHg. Beberapa nilai rendah mendekati 49 mmHg menunjukkan kemungkinan adanya hipotensi.
- Kadar gula darah (*BS*) menunjukkan nilai rata-rata 8,7 mmol/L dengan rentang 6–19 mmol/L. Nilai di atas 11 mmol/L dapat dikategorikan sebagai hiperglikemia.
- Suhu tubuh (*BodyTemp*) relatif stabil, berkisar antara 98°F hingga 103°F, dengan rata-rata 98,6°F.
- Denyut jantung (*HeartRate*) memiliki rata-rata 74 denyut/menit dengan variasi antara 7 hingga 90 denyut/menit.

b. Distribusi Kelas

Variabel target *RiskLevel* terdiri dari tiga kategori: *low risk*, *mid risk*, dan *high risk*. Hasil eksplorasi awal menunjukkan adanya indikasi ketidakseimbangan data, di mana kelas *low risk* lebih dominan dibandingkan *high risk*. Ketidakseimbangan ini berpotensi menimbulkan bias dalam klasifikasi.

c. Korelasi Antar Variabel

Analisis korelasi menunjukkan adanya keterkaitan yang cukup kuat antara tekanan darah (sistolik maupun diastolik) dengan tingkat risiko kehamilan. Selain itu, kadar gula darah (*BS*) juga memperlihatkan peran signifikan dalam membedakan kategori risiko. Variabel usia, meskipun kontribusinya tidak sekuat tekanan darah, tetap memiliki pengaruh terhadap risiko maternal.

d. Deteksi Outlier

Beberapa variabel memperlihatkan potensi outlier, misalnya usia ibu ≥ 60 tahun, serta nilai *HeartRate* dan gula darah yang sangat tinggi.

Hasil EDA menunjukkan bahwa dataset *Maternal Health Risk* relatif bersih tanpa nilai kosong, namun mengandung ketidakseimbangan kelas dan outlier yang signifikan. Variabel utama yang paling berhubungan dengan risiko maternal adalah tekanan darah (sistolik dan diastolik) serta kadar gula darah.

2.2 Preprocessing

Sebelum melakukan klasifikasi, tahap *preprocessing* dilakukan untuk memastikan data memiliki format dan skala yang sesuai, sehingga model *machine learning* dapat bekerja secara optimal [14].

Proses *preprocessing* dalam penelitian mencakup:

- a. Pengecekan *dataset* menggunakan operator *filter example* untuk mengatasi *missing value*, yang dapat mempengaruhi validitas hasil penelitian.
- b. *Split dataset* menjadikan *dataset* yang ada dibagi menjadi 2 bagian, yaitu data latih dan data uji. Sebelumnya peneliti sudah melakukan beberapa variasi perbandingan misalkan 60:40, 70:30, 80:20 dan 90:10 yang menghasilkan tingkat akurasi berbeda. Namun pada perbandingan 70:30 memiliki performa yang hampir baik di semua metode dan memiliki tingkat akurasi yang stabil.
- c. Penggunaan *Cross Validation* untuk validasi data latih dan data uji, dengan tujuan memastikan kesamaan jumlah data dalam *dataset* yang digunakan [15]. Dimana nilai k yang digunakan adalah 10.
- d. Optimalisasi parameter menggunakan *Optimize Parameter Grid* untuk meningkatkan kinerja model dengan menguji berbagai kombinasi parameter guna menemukan konfigurasi terbaik.

2.3 Komparasi Klasifikasi

Klasifikasi adalah metode pembelajaran yang digunakan untuk mengelompokkan data berdasarkan atribut tertentu. Dalam penelitian ini, dilakukan perbandingan tiga algoritma klasifikasi: *Decision Tree*, *Naive Bayes*, dan *KNN*, serta optimalisasi menggunakan *Optimize Parameter Grid*.

Pada tahapan komparasi algoritma, dilakukan klasifikasi beberapa algoritma yang berbeda. Penelitian ini membandingkan 3 algoritma yang terdiri dari *Decision Tree*, *Naive Bayes* dan *KNN* beserta optimasinya yang digunakan adalah *Optimize Parameter Grid*.

e. *Decision Tree*

Algoritma *Decision Tree* dikenal efektif dalam klasifikasi dan *clustering*. Metode ini mengubah data menjadi struktur pohon keputusan yang dapat dipresentasikan dalam bentuk aturan yang mudah difahami [16], [17].

Dalam penelitian ini, digunakan dua varian utama dari *Decision Tree*:

i. *Decision Tree ID3*

Algoritma ini menggunakan strategi pembelajaran induktif secara rekursif dari atas ke bawah. Pemilihan atribut dilakukan menggunakan *Information Gain* untuk mengukur efektivitas atribut dalam memisahkan data latih berdasarkan kelasnya.

ii. *Decision Tree C4.5*

Merupakan pengembangan dari ID3, dengan perbedaan utama pada penggunaan *Gain Ratio* sebagai metode pemilihan atribut.

f. *Naive Bayes*

Naive Bayes adalah algoritma klasifikasi berbasis probabilitas yang mengasumsikan bahwa semua atribut berkontribusi terhadap hasil klasifikasi secara independen (*class conditional independence*) [18]. Perhitungan dalam metode ini menggunakan Teorema Bayes.

g. *KNN (K-Nearest Neighbor)*

KNN adalah metode pembelajaran berbasis pengawasan yang menentukan kelas suatu data berdasarkan tetangga terdekat. Metode ini menggunakan jarak Euclidean dalam menentukan tetangga terdekatnya [19].

h. *Optimize Parameter Grid*

Teknik ini digunakan untuk mengoptimalkan parameter dalam model *machine learning* untuk meningkatkan kinerja model. Proses ini melibatkan evaluasi berbagai kombinasi parameter untuk menemukan konfigurasi yang menghasilkan akurasi terbaik [20]. *Optimize Parameter Grid* memiliki dua metode utama, yaitu:

2.4 Evaluasi

Evaluasi dilakukan untuk menilai kinerja algoritma yang digunakan dalam klasifikasi. Salah satu indikator utama adalah tingkat akurasi.

Akurasi dapat didefinisikan sebagai perbandingan nilai presentase dari *dataset* dimana sampel target dan non-target yang diprediksi dengan baik serta mampu menjadi gambaran kemampuan klasifikasi data uji berdasarkan data yang ada [21]. Nilai akurasi banyak dipengaruhi oleh dua hal yaitu dari segi jumlah data dan

juga berdasarkan *imbalancing* data yang kita gunakan sebagai data latih dan data uji. Tingkat nilai dapat dihitung menggunakan dengan rumus berikut:

$$Akurasi = \frac{TP+TN}{TP+FP+TN+FN} \times 100$$

Dimana :

True Positive (TP) = jumlah sampel positif yang diprediksi benar

False Positive (FP) = jumlah sampel positif yang diprediksi salah

True Negative (TN) = jumlah sampel negatif yang diprediksi benar

False Negative (FN) = jumlah sampel negatif yang diprediksi salah

Precision menunjukkan seberapa tepat model dalam memprediksi kelas positif. Semakin tinggi nilai *precision* maka semakin jarang model tersebut salah dalam melakukan klasifikasi. Nilai ini dapat dihitung menggunakan rumus:

$$Precision = \frac{TP}{TP + FP}$$

Dimana:

True Positive (TP) = jumlah data positif yang diprediksi benar

False Positive(FP) = jumlah data negatif yang salah diprediksi positif

Recall memperlihatkan seberapa lengkap model dalam menemukan semua data yang positif. Nilai *recall* yang tinggi menunjukkan bahwa model tersebut mampu menangkap hampir semua data positif. Berikut rumusnya:

$$Recall = \frac{TP}{TP + FN}$$

Dimana:

True Positive (TP) = jumlah data positif yang diprediksi benar

False Negative (FN) = jumlah data positif yang salah prediksi negatif

Untuk mencari nilai rata-rata harmonis antara *precision* dan *recall* menggunakan *F1-Score*. Dipakai untuk menyeimbangkan keduanya, terutama jika datanya tidak seimbang.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

Melalui evaluasi ini, tingkat akurasi, *precision*, *recall* dan *F1 Score* dari setiap algoritma akan dibandingkan untuk menentukan metode terbaik, baik sebelum maupun setelah dilakukan optimasi. . Hal ini dilakukan untuk memberikan gambaran yang lebih komprehensif mengenai kinerja model terutama dalam mengklasifikasikan kelas *High Risk* yang sangat penting di bidang kesehatan.

3. HASIL DAN PEMBAHASAN

Penelitian ini memanfaatkan dataset dari *UCI repository* yang kemudian diklasifikasikan menggunakan beberapa algoritma, yaitu *Decision Tree*, *Naive Bayes*, dan *KNN*. Selanjutnya, metode optimasi *Optimize Parameter Grid* diterapkan untuk mengevaluasi apakah optimalisasi ini dapat meningkatkan akurasi algoritma yang digunakan. Setelah dilakukan pengujian, diperoleh nilai akurasi yang kemudian dianalisis lebih lanjut.

3.1. Hasil Penelitian

Pengujian dilakukan menggunakan metode *cross validation*, yang membagi dataset menjadi beberapa bagian dengan ukuran yang sama. Data kemudian diproses menggunakan teknik *split validation*, dengan 70% data digunakan untuk pelatihan (*training dataset*) dan 30% untuk pengujian (*testing dataset*). Hasil pengujian klasifikasi “Risiko Ibu Hamil” ditampilkan dalam tabel 2 berikut.

Tabel 2. Hasil Akurasi Pengujian

No.	Algoritma	Sebelum Optimasi	Setelah Optimasi
1.	Decision Tree ID3	75.08 %	85.25 %
2.	Decision Tree C4.5	74.43 %	85.25 %
3.	Naive Bayes	61.31 %	61.31 %
4.	KNN	82.62 %	83.72 %

Tabel 3. Hasil *Recall* Pengujian

No.	Algoritma	Sebelum Optimasi	Setelah Optimasi
1.	Decision Tree ID3	74 %	83.65 %
2.	Decision Tree C4.5	69.48 %	84.90 %
3.	Naïve Bayes	58.66 %	81.47 %
4.	KNN	83.47 %	83.95 %

Tabel 4. Hasil *Precision* Pengujian

No.	Algoritma	Sebelum Optimasi	Setelah Optimasi
1.	Decision Tree ID3	76.52 %	84.30 %
2.	Decision Tree C4.5	80.72 %	85.60 %
3.	Naïve Bayes	66.05 %	82.15 %
4.	KNN	82.91 %	84.72 %

Tabel 5. Hasil *F1 Score* Pengujian

No.	Algoritma	Sebelum Optimasi	Setelah Optimasi
1.	Decision Tree ID3	75.15 %	83.97 %
2.	Decision Tree C4.5	74.68 %	85.25 %
3.	Naïve Bayes	62.16 %	81.80 %
4.	KNN	83.14 %	84.30 %

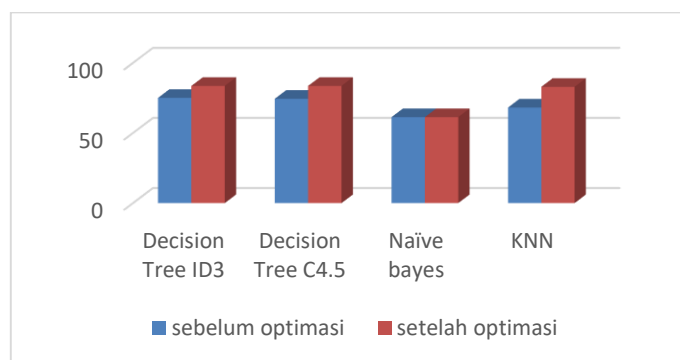
Sebelum dilakukan optimasi, hasil akurasi menunjukkan bahwa *KNN* memiliki performa terbaik dengan akurasi 82,62%, sementara *Naïve Bayes* mencatat nilai terendah sebesar 61,31%. Setelah dilakukan optimasi, terjadi peningkatan signifikan pada hampir semua algoritma, kecuali *Naïve Bayes*. *Decision Tree ID3* dan *C4.5* sama-sama mencapai akurasi terbaik sebesar 85,25%, disusul oleh *KNN* dengan 83,72%. Hal ini menegaskan bahwa penerapan *parameter tuning* mampu meningkatkan akurasi algoritma secara signifikan, meskipun dampaknya tidak seragam pada setiap metode.

Selain akurasi, evaluasi juga dilakukan pada metrik *precision*, *recall*, dan *F1-score*. Hasil menunjukkan pola yang konsisten: *Decision Tree C4.5* cenderung memiliki nilai *precision* dan *F1-score* yang lebih stabil dibandingkan *ID3*, sementara *KNN* unggul dalam *recall*. Sebaliknya, *Naïve Bayes* menunjukkan kinerja terendah di hampir semua metrik, bahkan setelah optimasi.

3.2. Pembahasan Penelitian

Optimasi parameter berperan penting dalam meningkatkan akurasi prediksi suatu algoritma. Hasil penelitian menunjukkan bahwa semua algoritma yang diuji mengalami peningkatan akurasi setelah dioptimasi, kecuali *Naive Bayes*. Setelah penerapan *Optimize Paramater Grid*, algoritma *Decision Tree ID3* dan *Decision Tree C4.5* menghasilkan nilai akurasi yang cukup menarik, karena mencapai akurasi yang sama yaitu 85,25 % menjadikannya metode klasifikasi dengan performa terbaik, padahal sebelumnya keduanya memiliki tingkat akurasi yang berbeda. Sementara itu, algoritma *Naive Bayes* tetap memiliki akurasi terendah sebesar 61,31% meskipun telah dilakukan optimasi.

Hasil ini memperlihatkan bahwa optimasi parameter berperan penting dalam meningkatkan kualitas prediksi, terutama pada algoritma yang memiliki banyak variabel parameter seperti *Decision Tree* dan *KNN*. Peningkatan akurasi sebesar 10,82 % pada *Decision Tree C4.5* dan 10,17% pada *ID3* menjadi bukti nyata bahwa penyesuaian parameter dapat memperbaiki kelemahan bawaan dari algoritma. Sebaliknya, *Naïve Bayes* yang berbasis probabilistik cenderung stabil meskipun tanpa optimasi, sehingga tidak menunjukkan perubahan kinerja setelah *tuning*. Grafik perbandingan akurasi hasil pengujian dapat dilihat pada Gambar 1.

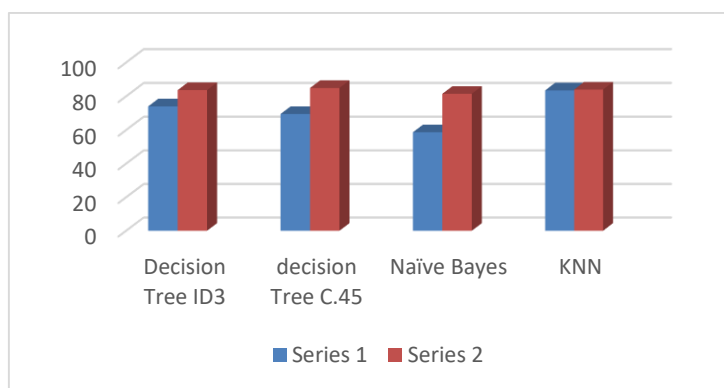


Gambar 2. Perbandingan algoritma sebelum dan sesudah optimasi

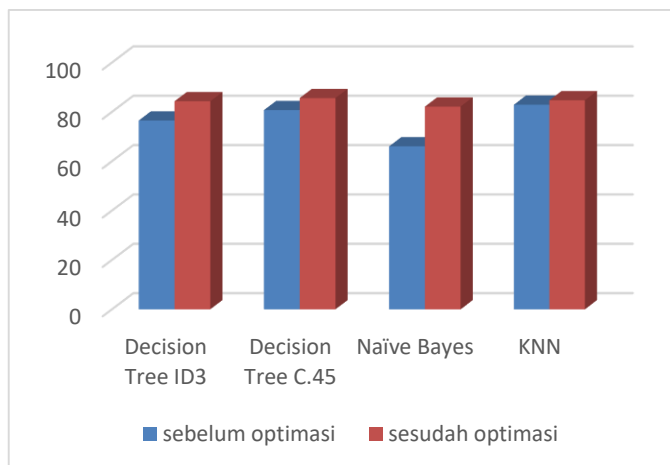
Algoritma *Decision Tree* sangat bergantung pada pemilihan atribut dan parameter. Sehingga ketika sebelum dioptimasi, pohon bisa terlalu sederhana (*underfitting*) atau terlalu kompleks (*overfitting*), sehingga akurasinya tidak maksimal. Setelah dilakukan *Optimize Parameter Grid*, parameter-parameter tersebut dapat disesuaikan secara sistematis sehingga menghasilkan pohon keputusan yang lebih seimbang, tidak terlalu dalam dan tidak terlalu dangkal. Inilah sebabnya akurasi *Decision Tree ID3* dan *C4.5* mengalami peningkatan performa. Fakta bahwa *ID3* dan *C4.5* menghasilkan akurasi yang sama setelah dioptimasi menunjukkan bahwa parameter *tunning* lebih berpengaruh daripada jenis algoritma *tree* itu sendiri pada data set *maternal health risk*.

Performa *KNN* sangat dipengaruhi oleh nilai parameter *k* dan metrik jarak yang digunakan. Jika nilai *k* terlalu kecil, model mudah *overfitting*, jika terlalu besar, prediksi menjadi terlalu *general* dan kehilangan detail. Optimasi parameter yang sesuai dengan karakteristik dataset tetapi tidak ada perubahan nilai parameter *k*, memberikan peningkatan yang signifikan terhadap *KNN*.

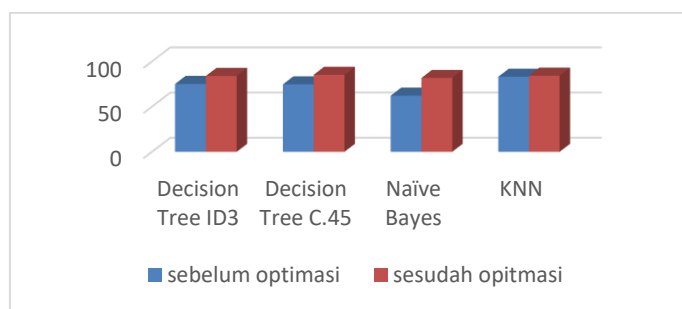
Naïve Bayes bekerja dengan asumsi bahwa setiap fitur independen satu sama lain (*conditional independence*). Namun, pada dataset *maternal health risk*, terdapat keterkaitan kuat antar atribut, misalnya tekanan darah sistolik dan diastolik yang saling berkorelasi. Karena asumsi independensi tidak terpenuhi, performa *Naïve Bayes* menjadi terbatas. Dari sini dapat kita lihat bahwa *Naïve Bayes* tidak memiliki banyak parameter untuk dioptimasi. Sehingga, meskipun dilakukan *Optimize Parameter Grid*, hasil akurasinya tidak berubah.



Gambar 3. Perbandingan recall sebelum dan sesudah optimasi



Gambar 4. Perbandingan precision sebelum dan sesudah optimasi



Gambar 5. Perbandingan F1 score sebelum dan sesudah optimasi

Berdasarkan gambar 2,3 dan 4 hasil pengujian menunjukkan bahwa setiap algoritma mengalami peningkatan kinerja setelah dilakukan optimasi parameter. Peningkatan ini terlihat dari nilai *recall*, *precision*, dan *F1-score* yang secara konsisten lebih tinggi dibandingkan sebelum optimasi.

Pada algoritma *Decision Tree ID3*, nilai *recall* meningkat dari 74% menjadi 83,65%, *precision* dari 76,52% menjadi 84,30%, serta *F1-score* dari 75,15% menjadi 83,97%. Hal ini menunjukkan bahwa setelah optimasi, model lebih mampu mengenali kelas positif sekaligus mengurangi kesalahan prediksi. Peningkatan ini dapat dijelaskan karena optimasi parameter pada *ID3* mampu menghasilkan struktur pohon keputusan yang lebih representatif, sehingga tingkat generalisasi model terhadap data uji semakin baik. Algoritma *Decision Tree C4.5* mengalami peningkatan yang signifikan, dengan *recall* dari 69,48% menjadi 84,90%, *precision* dari 80,72% menjadi 85,60%, dan *F1-score* dari 74,68% menjadi 85,25%. Kinerja *C4.5* yang semakin baik setelah optimasi disebabkan oleh proses pemilihan atribut melalui *gain ratio* yang lebih tepat serta penerapan *pruning*, sehingga model tidak hanya kompleks, tetapi juga lebih efisien dalam mengklasifikasikan data.

Sementara itu, algoritma *Naïve Bayes* menunjukkan peningkatan yang paling drastis. *Recall* naik dari 58,66% menjadi 81,47%, *precision* dari 66,05% menjadi 82,15%, serta *F1-score* dari 62,16% menjadi 81,80%. Lonjakan performa ini terjadi karena optimasi mampu menyesuaikan parameter distribusi dan mengurangi dampak asumsi independensi antar fitur yang sering menjadi keterbatasan utama *Naïve Bayes*. Dengan demikian, model menjadi lebih sesuai dengan karakteristik data yang digunakan.

Berbeda dengan algoritma lain, *K-Nearest Neighbor (KNN)* hanya mengalami peningkatan kecil. *Recall* meningkat dari 83,47% menjadi 83,95%, *precision* dari 82,91% menjadi 84,72%, serta *F1-score* dari 83,14% menjadi 84,30%. Hal ini menunjukkan bahwa *KNN* sejak awal sudah memiliki kinerja yang cukup baik, sehingga optimasi hanya memberikan perbaikan pada stabilitas model, terutama melalui pemilihan nilai *k* dan metrik jarak yang lebih sesuai. Dari perspektif implementasi, hasil ini memiliki implikasi penting bagi tenaga medis dalam mendukung proses pengambilan keputusan. Algoritma dengan akurasi lebih tinggi, khususnya *Decision Tree C4.5*, dapat digunakan untuk membangun sistem pendukung keputusan dalam mengidentifikasi tingkat risiko ibu hamil. Hal ini krusial karena prediksi yang lebih akurat memungkinkan tenaga medis melakukan intervensi lebih dini, sehingga dapat menekan angka komplikasi bahkan kematian ibu.

Dalam penelitian ini menyoroti pentingnya memperhatikan ketidakseimbangan kelas (*class imbalance*) pada dataset *maternal health risk*. Kelas *low risk* yang dominan berpotensi menyebabkan bias model terhadap prediksi kelas mayoritas. Oleh karena itu, penelitian selanjutnya dapat mengintegrasikan metode penyeimbangan data untuk menguji apakah akurasi dan *F1-score* pada kelas *high risk* dapat lebih ditingkatkan.

4. PENUTUP

Penelitian ini menunjukkan bahwa penerapan metode *Optimize Parameter Grid* mampu meningkatkan kinerja sebagian besar algoritma klasifikasi dalam mengidentifikasi tingkat risiko ibu hamil. Sebelum optimasi, algoritma *KNN* memberikan akurasi terbaik sebesar 82,62%, sementara *Naïve Bayes* mencatat hasil terendah sebesar 61,31%. Setelah dilakukan optimasi, *Decision Tree ID3* dan *C4.5* mencapai performa terbaik dengan akurasi 85,25%, diikuti oleh *KNN* dengan 83,72%. Peningkatan ini juga konsisten pada metrik evaluasi lain seperti *precision*, *recall*, dan *F1-score*.

Secara umum, optimasi memberikan dampak yang signifikan pada algoritma berbasis pohon keputusan dan *KNN*, sedangkan *Naïve Bayes* relatif tidak terpengaruh karena keterbatasan asumsi independensi antar fitur. Hal ini membuktikan bahwa pemilihan serta penyesuaian parameter menjadi faktor penting dalam membangun model klasifikasi yang lebih akurat dan andal. Temuan ini memiliki implikasi praktis dalam bidang kesehatan, khususnya dalam mendukung tenaga medis untuk mendeteksi risiko ibu hamil secara lebih objektif dan konsisten. Dengan akurasi yang lebih baik, sistem berbasis *machine learning* dapat

digunakan sebagai alat bantu pengambilan keputusan guna mencegah komplikasi dan menekan angka kematian ibu. Untuk penelitian selanjutnya, disarankan penggunaan metode penyeimbangan data (*class balancing*) serta pengujian pada algoritma lain yang lebih kompleks seperti *ensemble learning* atau *deep learning*. Hal ini diharapkan dapat meningkatkan kemampuan model dalam mengklasifikasikan kategori *high risk*, yang memiliki peran krusial dalam upaya pencegahan risiko kehamilan.

DAFTAR PUSTAKA

- [1] S. N. Tarmizi, "Sehat Negeriku," Oct. 29, 2023, *Kepala Biro dan Komunikasi Publik Kemkes*. [Online]. Available: <https://sehatnegeriku.kemkes.go.id/baca/rilis-media/20230115/4842206/turunkan-angka-kematian-ibu-melalui-deteksi-dini-dengan-pemenuhan-usg-di-puskesmas/>
- [2] H. Amalia, R. Rahmadanti, A. Syaiin, and S. Salsabila, "Prediksi Risiko Kesehatan Ibu Hamil Dengan Menggunakan Metode," *J. SWABUMI*, vol. 11, p. 1, 2023.
- [3] O. W. Yuda, D. Tuti, L. S. Yee, and Susanti, "Penerapan Data Mining Untuk Klasifikasi Kelulusan Mahasiswa Tepat Waktu Menggunakan Metode Random Forest," *SATIN - Sains dan Teknol. Inf.*, vol. 8, no. 2, pp. 122–131, 2022.
- [4] S. Chowdhury, S. Mishra, A. Miranda, and P. Mallick, "Energy Consumption Prediction Using Light Gradient Boosting Machine Model," in *Lecture Notes in Electrical Engineering*, vol. 690, 2021, pp. 413–422.
- [5] L. N. et al., "Streamflow forecasting using extreme gradient boosting model coupled with Gaussian mixture model," *J. Hydrol.*, vol. 586, p. 124901, 2020.
- [6] H. Amalia, D. Rahmadanti, A. Syaiin, S. Salsabila, Yunita, and Sriyadi, "Prediksi Resiko Kesehatan Ibu Hamil Dengan Menggunakan Metode Decision Tree," *J. SWABUMI*, vol. 11, pp. 48–53, 2023.
- [7] R. A. Fitra, W. A. Harahap, and W. K. Rahman, "Analisis Perbandingan Algoritma Machine Learning untuk Klasifikasi Tingkat Risiko Ibu Hamil," *Student Res. J.*, vol. 1, pp. 246–261, 2023.
- [8] J. J. Purnama, N. K. Hikmawati, and S. Rahayu, "Analisis Algoritma Klasifikasi Untuk Mengidentifikasi Potensi Risiko Kesehatan Ibu Hamil," *J. Appl. Comput. Sci. Technol.*, vol. 5, pp. 120–127, 2024.
- [9] F. Friedrichs and C. Igel, "Evolutionary tuning of multiple SVM parameters," *Neurocomputing*, vol. 64, pp. 107–117, 2005.
- [10] O. Sagi and L. Rokach, "Approximating XGBoost with an interpretable decision tree," *Inf. Sci. (Nij.)*, vol. 572, pp. 522–542, 2021.
- [11] T. Rahman, "Perbandingan Produkt. Kerja Karyawan Sebelum," 2018.
- [12] A. Fauzi, "Analisis Data Bank Direct Marketing dengan Perbandingan Klasifikasi Data Mining Berbasis Optimize Selection (Evolutionary)," *J. Inform. Univ. Pamulang*, vol. 6, p. 102, 2021.
- [13] A. H. Yunial, "Analisis Optimasi Algoritma Klasifikasi Support Vector Machine, Decision Trees, dan Neural Network Menggunakan Adaboost dan Bagging," *J. Inform. Univ. Pamulang*, vol. 5, p. 247, 2020.
- [14] F. Admojo and S. Jabir, "Analisis Performa Metode Naive Bayes Clasifier pada Electronic Nose Dalam Identifikasi Formalin Pada Tahu," *Indones. J. Data Sci.*, vol. 4, pp. 1–16, 2023.
- [15] R. R. Sudriyanto and M. A. R. Hariri, "Implementasi Algoritma Decision Tree (C4.5) dengan Optimize Wheights (PSO) Untuk Memprediksi Kelulusan Mahasiswa Tepat Waktu," *J. Inform. Univ. Pamulang*, vol. 6, pp. 252–257, 2021.
- [16] N. Romadloni, I. Santoso, and S. Bidolaksono, "Perbandingan Metode Naive Bayes, Knn dan Decision Tree Terhadap Analisis Sentimen Transportasi Krl Commuter Line," *J. IKRA-ITH Inform. J. Komput. dan Inform.*, vol. 3, pp. 1–9, 2019.
- [17] C. Cahyaningtyas, Y. Nataliani, and I. R. Widiyari, "Analisis Sentimen Pada Rating Aplikasi Shopee Menggunakan Metode Decision Tree Berbasis SMOTE," *AITI*, vol. 18, pp. 173–184, 2021.
- [18] Q. Hasanah, A. Andrianti, and M. Hidayat, "Sistem Informasi Posyandu Ibu Hamil dengan Penerapan Klasifikasi Risiko Kehamilan Menggunakan Metode Naive Bayes".
- [19] Yahya and W. P. Hidayanti, "Penerapan Algoritma K-Nearest Neighbor Untuk Klasifikasi Efektivitas Penjualan Vape (Rokok Elektrik) pada 'Lombok Vape On,'" *Infotek J. Inform. dan Teknol.*, vol. 3, pp. 104–114, 2020.
- [20] "Optimize Parameters (Grid)," Jan. 24, 2024. [Online]. Available: https://docs.rapidminer.com/9.10/studio/operators/modeling/optimization/parameters/optimize_parameters_grid.html
- [21] J. Lin, H. Chen, S. Li, Y. Liu, X. Li, and B. Yu, "Accurate Prediction of Potential Druggable Proteins Based on Genetic Algorithm and Bagging-SVM Ensemble Classifier," *Artif. Intell. Med.*, vol. 98, pp. 35–47, 2019.