



Pengaruh Hyperparameter Tuning Gridsearch pada Random Forest untuk Klasifikasi Diabetes Melitus

Resa Astari Br Tarigan¹, I Gusti Ayu Agung Diatri Indradewi², Gede Surya Mahendra³
Program Studi Sistem Informasi, Fakultas Teknik dan Kejuruan, Universitas Pendidikan Ganesha
resa.astari@student.undiksha.ac.id

Abstrak

Diabetes Melitus merupakan salah satu penyakit kronis yang memerlukan deteksi dini dan klasifikasi yang akurat untuk mendukung pengambilan keputusan dalam bidang kesehatan. Salah satu algoritma yang banyak digunakan untuk tugas klasifikasi adalah Random Forest, namun performanya dipengaruhi oleh pemilihan hyperparameter yang tepat. Penelitian ini bertujuan untuk menganalisis pengaruh hyperparameter tuning menggunakan metode GridSearchCV terhadap kinerja algoritma Random Forest dalam klasifikasi Diabetes Melitus. Dataset yang digunakan diperoleh dari Mendeley Data dan terdiri atas 826 data setelah melalui tahap preprocessing. Penelitian dilakukan dengan membandingkan performa model Random Forest sebelum dan sesudah proses hyperparameter tuning menggunakan metrik accuracy, precision, recall, dan F1-score. Hasil pengujian menunjukkan bahwa model tanpa tuning menghasilkan accuracy sebesar 97,59%, precision 99,07%, recall 85,75%, dan F1-score 90,93%. Setelah dilakukan tuning menggunakan GridSearchCV, diperoleh parameter terbaik yaitu $\text{criterion} = \text{gini}$, $\text{max_depth} = 15$, $\text{min_samples_leaf} = 1$, $\text{max_features} = \text{sqrt}$, dan $\text{n_estimators} = 300$, yang menghasilkan accuracy sebesar 98,19%, precision 99,30%, recall 89,91%, dan F1-score 93,98%. Hasil penelitian menunjukkan bahwa penerapan GridSearchCV memberikan pengaruh positif terhadap performa Random Forest, terutama dalam meningkatkan kemampuan model mengenali kelas minoritas yang ditunjukkan oleh peningkatan nilai recall dan F1-score. Dengan demikian, hyperparameter tuning menggunakan GridSearchCV dapat digunakan untuk mengoptimalkan kinerja Random Forest dalam klasifikasi Diabetes Melitus gayanya.

Kata Kunci: Komponen; Diabetes Melitus, Random Forest, Hyperparameter Tuning, GridSearchCV, Klasifikasi

1. Latar Belakang

Perkembangan teknologi informasi telah mendorong penerapan metode berbasis data di berbagai bidang, termasuk dalam sektor kesehatan. Machine learning (ML) merupakan salah satu solusi yang paling banyak digunakan untuk mendukung proses diagnosis klinis dan pengambilan keputusan, mengingat kebutuhan akan sistem yang mampu menyediakan informasi dengan cepat dan akurat [1]. Machine learning merupakan pendekatan yang memungkinkan sistem komputer untuk mempelajari pola dari data secara sistematis guna membangun model yang merepresentasikan hubungan dalam data. Selain itu, machine learning mampu mengekstraksi informasi penting melalui analisis pola dan hubungan antar data dalam jumlah besar [2]. Salah satu penerapan machine learning dalam bidang kesehatan adalah klasifikasi, yaitu proses pengelompokan data ke dalam kategori tertentu berdasarkan karakteristik yang dimiliki. Teknik ini banyak digunakan dalam diagnosis dan analisis data medis untuk mendukung pengambilan keputusan klinis [3]. Dalam konteks klasifikasi penyakit, model dibangun menggunakan berbagai atribut medis seperti hasil pemeriksaan laboratorium dan data klinis pasien. Namun, keberhasilan model klasifikasi tidak hanya ditentukan oleh algoritma yang digunakan, tetapi juga dipengaruhi oleh kualitas data serta konfigurasi parameter yang digunakan dalam proses pelatihan model [4]. pemilihan metode klasifikasi berpengaruh signifikan terhadap accuracy, sensitivity, dan specificity [5]. Random Forest merupakan salah satu algoritma klasifikasi berbasis ensemble yang banyak digunakan karena kemampuannya dalam menangani data dengan hubungan non-linear serta mengurangi risiko overfitting melalui pendekatan bagging [6]. Meskipun demikian, performa Random Forest sangat dipengaruhi oleh pemilihan hyperparameter yang tepat, seperti n_estimators , max_depth , criterion , min_samples_leaf , max_features dan parameter lainnya. Pemilihan hyperparameter yang tidak optimal dapat menyebabkan model tidak bekerja secara maksimal, baik dalam bentuk underfitting maupun overfitting. Untuk mengatasi permasalahan tersebut, diperlukan teknik optimasi hyperparameter yang mampu menentukan kombinasi parameter terbaik. Hyperparameter dalam algoritma Machine Learning adalah parameter yang ditentukan sebelum proses pelatihan model dan tidak dipelajari secara langsung dari data. Fungsinya adalah untuk mengontrol proses pembangunan model, sehingga pemilihan

parameter tersebut memiliki dampak yang signifikan terhadap kinerja prediktif model. Salah satu metode yang umum digunakan dalam optimasi hyperparameter adalah Grid Search, yaitu teknik yang mengevaluasi berbagai kombinasi parameter secara sistematis untuk menemukan konfigurasi optimal [7]. Penerapan Grid Search yang dikombinasikan dengan cross-validation memungkinkan proses evaluasi model dilakukan secara lebih menyeluruh, sehingga menghasilkan performa yang lebih stabil. Penelitian sebelumnya menunjukkan bahwa optimasi hyperparameter dapat meningkatkan kinerja model klasifikasi. Penelitian pada klasifikasi diabetes menggunakan algoritma XGBoost menunjukkan peningkatan accuracy yang signifikan setelah penerapan Grid Search, dari 75% menjadi 95% [8]. Hasil serupa juga ditemukan pada klasifikasi penyakit ginjal menggunakan Random Forest, di mana accuracy meningkat setelah dilakukan optimasi parameter [9]. Namun demikian, tidak semua penelitian menunjukkan hasil yang konsisten, karena pada kasus klasifikasi genre musik, penerapan Grid Search justru menyebabkan penurunan accuracy model [10]. Diabetes melitus merupakan salah satu penyakit kronis dengan tingkat prevalensi yang terus meningkat dan memerlukan deteksi dini untuk mencegah komplikasi yang lebih serius. Penyakit ini ditandai dengan gangguan metabolisme yang menyebabkan kadar gula darah tinggi dan dapat berdampak pada berbagai organ tubuh [11]. Peningkatan jumlah penderita diabetes secara global menunjukkan perlunya pendekatan berbasis teknologi yang mampu membantu proses klasifikasi penyakit secara lebih akurat [12]. Berdasarkan uraian tersebut penelitian ini bertujuan untuk menganalisis pengaruh hyperparameter tuning menggunakan metode Grid SearchCV terhadap kinerja algoritma Random Forest dalam klasifikasi diabetes melitus berbasis data klinis. Evaluasi kinerja model dilakukan menggunakan metrik accuracy, precision, recall, dan F1-score guna menilai sejauh mana optimasi parameter dapat meningkatkan performa model.

2. Metode Penelitian

Knowledge Discovery in Databases (KDD) adalah metode data mining yang bertujuan untuk menemukan informasi berharga dan pola tersembunyi dalam data. Berbagai algoritma digunakan dalam proses ini untuk menemukan hubungan antar data. KDD terdiri dari beberapa tahapan utama: pemilihan, pre-processing, transformasi, data mining, dan interpretasi dan evaluasi. Tahapan-tahapan ini dilakukan secara sistematis untuk menghasilkan pengetahuan yang bermanfaat [13] tahap penelitian dapat dilihat dari gambar 1.



Gambar 1. Metode Penelitian

2.1. Pengumpulan Data

Dataset yang digunakan dalam penelitian ini diperoleh dari Mendeley Data dengan judul Diabetes Dataset yang dipublikasikan pada tanggal 18 Juli 2020, oleh Ahlam Rashid. Dataset terdiri dari sejumlah fitur yang

merepresentasikan kondisi kesehatan pasien, dataset yang digunakan memiliki 1000 data dengan jumlah fitur 12 dengan variabel target Diabetes, Pre-Diabetes, dan Non-Diabetes.

2.2. Preprocessing

1. Data Cleaning

Tahapan preprocessing data dilakukan untuk mempersiapkan data mentah menjadi lebih terstruktur dan siap digunakan dalam proses pemodelan, sehingga dapat meningkatkan kualitas hasil klasifikasi. Pada Tahap ini Data Cleaning dilakukan agar mengatasi terjadinya missing value, data duplikat.

2.3 Tranformasi Data

Tahapan ini bertujuan untuk mempersiapkan data agar dapat digunakan dalam proses pemodelan. Proses preprocessing melakukan Encoding variabel kategorikal, pada tahap ini, variabel-variabel kategorikal dalam dataset diubah ke dalam bentuk numerik agar dapat diproses oleh algoritma machine learning.

2.4 Data Mining

Pada tahap ini dilakukan proses pemodelan untuk membangun model klasifikasi menggunakan algoritma Random Forest. Proses ini diawali dengan pembentukan model awal (baseline model) tanpa penerapan hyperparameter tuning. Model baseline ini digunakan sebagai acuan untuk mengetahui performa awal algoritma sebelum dilakukan optimasi parameter.

Selanjutnya, untuk meningkatkan kinerja model, dilakukan proses hyperparameter tuning menggunakan metode GridSearchCV. Metode ini bekerja dengan menguji berbagai kombinasi parameter secara sistematis untuk menemukan konfigurasi terbaik yang dapat meningkatkan performa model. Proses pencarian dilakukan dengan mengombinasikan teknik grid search dan cross-validation (K-Fold), sehingga setiap kombinasi parameter diuji pada beberapa subset data latih untuk menghasilkan evaluasi yang lebih stabil dan objektif.

Parameter yang dioptimasi dalam penelitian ini meliputi `n_estimators` banyaknya pohon yang dibangun dalam hutan pohon, `max_depth` merupakan batas maksimum kedalaman setiap pohon, `criterion` yaitu batas maksimum kedalaman setiap pohon, `min_samples_leaf` merupakan jumlah minimum sampel yang diperlukan pada daun pohon, dan `max_features` merupakan jumlah fitur yang dipertimbangkan saat menentukan pemisahan terbaik [14]. Rentang nilai dari masing-masing parameter ditentukan dalam parameter grid yang disajikan pada Tabel 1.

Tabel 1. Nilai parameter GridSearchCV

Parameter	Grid Search Values
<code>n_estimators</code>	100,200,300
<code>max_depth</code>	5,7,9,15, none
<code>criterion</code>	Gini, Entropy
<code>min_samples_leaf</code>	1,2,3,4,5
<code>max_features</code>	<code>sqrt</code> , <code>log2</code>

Setiap kombinasi parameter akan dievaluasi berdasarkan nilai rata-rata performa yang dihasilkan selama proses validasi silang, kemudian dipilih kombinasi terbaik yang memberikan skor tertinggi. Model hasil optimasi kemudian dibandingkan dengan model baseline untuk melihat pengaruh penerapan hyperparameter tuning terhadap peningkatan kinerja model [15]. Dengan demikian, dapat dianalisis sejauh mana proses tuning memberikan kontribusi dalam meningkatkan akurasi dan kestabilan model Random Forest dalam klasifikasi data.

2.5 Evaluasi Model

Setelah final model diperoleh, dan telah diimplementasikan pada data uji langkah berikutnya adalah mengevaluasi model tersebut. Penilaian model ini adalah tahap yang dilakukan untuk menilai kinerja model yang sudah dibuat. Penilaian performa model dilakukan dengan memanfaatkan metrik evaluasi klasifikasi yang meliputi accuracy, precision, recall, dan F1-score.

3. Hasil dan Pembahasan

3.1. Data Selection

Dataset yang digunakan dalam penelitian ini merupakan dataset penyakit Diabetes Melitus yang diperoleh dari platform Mendeley Data. Dataset ini digunakan untuk melakukan proses klasifikasi guna memprediksi kondisi pasien berdasarkan atribut yang tersedia. Dataset terdiri dari sejumlah fitur yang merepresentasikan kondisi kesehatan pasien, dataset yang digunakan memiliki 1000 data dengan jumlah fitur 14 yaitu ID, no pation, usia pasien, jenis kelamin, kadar gula darah, indeks massa tubuh (BMI), rasio kreatinin (Cr), urea, kolesterol total, profil lipid puasa (LDL, VLDL, HDL, trigliserida), serta HbA1c, dengan kelas target berupa Diabetes, Non-Diabetes, dan Prediksi Diabetes dengan variabel target Diabetes, Pre-Diabetes, dan Non-Diabetes. Distribusi kelas awal dataset ditunjukkan pada Tabel 2

Tabel 2. Distribusi kelas awal

Kelas	Keterangan	Jumlah data
Y	Diabetes	844
N	Non-Diabetes	103
P	Pre-Diabetes	53

Dari 14 variabel tersebut, terdapat dua atribut yang tidak diikutsertakan dalam proses klasifikasi, yaitu *ID* dan *No_Pation*, karena kedua atribut tersebut hanya berfungsi sebagai pengenal unik rekaman dan tidak merepresentasikan informasi medis yang relevan terhadap prediksi diabetes. Sehingga, jumlah variabel yang digunakan dalam proses klasifikasi adalah 11 variabel, dan 1 variabel label kelas

3.2 Pra-Pemrosesan Data

3.2.1 Analisis Korelasi

Pada tahap ini Analisis korelasi dilakukan untuk mengidentifikasi hubungan linear antar variabel numerik dalam *dataset*. Hasil perhitungan korelasi antar variabel numerik ditampilkan dalam bentuk *Heatmap* Korelasi Antar Variabel matriks korelasi pada Gambar 2 sebagian besar hubungan variabel menunjukkan nilai korelasi yang rendah. Hal ini mengindikasikan bahwa hubungan linear antar variabel relatif lemah, tidak ditemukan pasangan variabel yang memiliki nilai koefisien korelasi di rentang nilai ($>|0,80|$), hal ini menunjukkan bahwa tidak terdapat masalah *multikolinieritas* atau *redundansi* data antar fitur pendukung. Oleh karena itu, seluruh fitur dipertahankan untuk digunakan dalam proses klasifikasi.



Gambar 2 Analisis Korelasi

3.2.2 Data Cleaning

Tahap pra-pemrosesan dilakukan untuk mempersiapkan data sebelum proses pelatihan model. Tahapan tersebut meliputi.

DOI: <https://doi.org/10.31004/riggs.v5i2.11967>

Lisensi: Creative Commons Attribution 4.0 International (CC BY 4.0)

1. Missing Value

Pemeriksaan missing value dilakukan dengan mengevaluasi jumlah data pada setiap variabel dalam dataset. Berdasarkan hasil pemeriksaan struktur dataset, diperoleh bahwa data penelitian terdiri dari 1.000 entry dengan total 14 variabel. Seluruh variabel memiliki jumlah data yang lengkap tanpa ditemukan nilai kosong. Dengan demikian, dapat disimpulkan bahwa dataset tidak mengandung missing value dan siap untuk digunakan pada tahap analisis selanjutnya tanpa memerlukan proses penanganan tambahan.

2. Data Duplikat

Pada tahap preprocessing, dilakukan pengecekan terhadap keberadaan data duplikat dalam dataset. Berdasarkan hasil pengecekan ditemukan sebanyak 174 data duplikat dalam dataset awal. Oleh karena itu, dilakukan proses penghapusan data duplikat. Setelah proses penghapusan dilakukan, jumlah data berkurang menjadi 826 data, dan tidak ditemukan lagi data duplikat dalam dataset.

3. Ketidakkonsistenan Data

Ditemukan beberapa ketidakkonsistenan data seperti perbedaan huruf besar dan kecil dan tambahan spasi. Karena itu dilakukan standarisasi data agar format atribut menjadi konsisten sebelum proses pemodelan. Tahap ini bertujuan untuk memeriksa konsistensi data pada atribut kategorikal, yaitu Gender dan CLASS. Hasil pemeriksaan menunjukkan adanya ketidakkonsistenan penulisan, seperti perbedaan huruf kapital dan huruf kecil pada atribut Gender (F, M, f) serta adanya spasi tambahan pada atribut CLASS (N, N , Y, Y). Oleh karena itu dilakukan proses standarisasi untuk menghapus spasi dan menyeragamkan seluruh nilai sehingga data menjadi konsisten seperti pada Tabel 3

Tabel 3. Nilai ketidakkonsistenan data

Variabel	Nilai Tidak Konsisten	Nilai Setelah Standarisasi
Gender	'M': 565, 'F': 434, 'f':	'M': 565, 'F': 435
CLASS	'Y': 840, 'N': 102, 'P': 53, 'Y ': 4, 'N ': 1	'Y': 844, 'N': 103, 'P': 53

3.3 Transformasi Data

Untuk membuat variabel kategorikal dapat diproses oleh algoritma pembelajaran mesin, dilakukan transformasi data. Metode yang digunakan dalam penelitian ini adalah label encoding. Kumpulan data tersebut berisi dua variabel kategorikal Gender dan CLASS. Variabel Gender diubah menjadi nilai numerik, dengan laki-laki menjadi angka 0 dan perempuan oleh angka 1. Variabel target, CLASS, diubah menjadi tiga kelas: Non-diabetes (0), Prediabetes (1), dan Diabetes (2). Setelah proses transformasi selesai, semua variabel dalam dataset telah diubah menjadi nilai numerik dan siap digunakan.

3.4 Data Mining

Pada tahap data mining, dilakukan pembangunan model klasifikasi menggunakan algoritma Random Forest tanpa penerapan hyperparameter tuning. Model dibangun menggunakan data yang telah melalui tahap preprocessing dan pembagian data menjadi data latih dan data uji. Pada tahap ini, Random Forest menggunakan parameter bawaan tanpa proses optimasi parameter.

Untuk memperoleh hasil evaluasi yang lebih stabil dan mengurangi pengaruh pembagian data secara acak, proses pelatihan model dilakukan menggunakan metode 5-Fold Cross Validation. Proses ini dilakukan sebanyak lima kali hingga seluruh data memperoleh kesempatan menjadi data validasi. Nilai performa yang diperoleh merupakan rata-rata hasil dari kelima proses validasi tersebut sehingga dapat mengurangi bias yang disebabkan oleh pemilihan data secara acak.

Algoritma Random Forest bekerja dengan membangun sejumlah pohon keputusan berdasarkan subset data dan subset fitur yang dipilih secara acak. Setiap pohon menghasilkan prediksi kelas, kemudian hasil akhir ditentukan

melalui mekanisme voting mayoritas dari seluruh pohon yang terbentuk. Model yang dihasilkan pada tahap ini digunakan sebagai baseline untuk mengetahui performa awal algoritma sebelum dilakukan proses optimasi menggunakan GridSearchCV.

Selanjutnya, hasil performa model Random Forest tanpa hyperparameter tuning dibandingkan dengan model Random Forest yang telah melalui proses hyperparameter tuning menggunakan GridSearchCV untuk mengetahui pengaruh optimasi parameter terhadap kinerja klasifikasi Diabetes Melitus

3.5 Performa Model Tanpa Hyperparameter Tuning

Evaluasi awal terhadap model klasifikasi menggunakan algoritma Random Forest tanpa penerapan hyperparameter tuning diperoleh nilai Accuracy sebesar 97.59%, dengan Precision sebesar 99.07%, Recall sebesar 85.75%, dan F1-score sebesar 90.93%.

```
***** EVALUASI RF TANPA TUNING *****
Accuracy : 97.59%
Precision : 99.07%
Recall : 85.75%
F1-Score : 90.93%

Classification Report:

              precision    recall  f1-score   support

   Normal (N)         1.00      0.95      0.97         19
  Pre-diabetes (P)         1.00      0.62      0.77          8
    Diabetes (Y)         0.97      1.00      0.99        139

   accuracy              0.98         166
  macro avg              0.99      0.86      0.91         166
 weighted avg              0.98      0.98      0.97         166

Confusion Matrix:
[[ 18  0  1]
 [ 0  5  3]
 [ 0  0 139]]
```

Gambar 3. Classification report dan confusion matriks tanpa haperparameter tuning

Berdasarkan hasil confusion matrix pada Gambar 3, masih terdapat beberapa data yang belum berhasil diklasifikasikan dengan benar oleh model. Kesalahan klasifikasi terutama terjadi pada kelas minoritas. Pada kelas Non-diabetes, terdapat 1 data yang salah diprediksi sebagai Diabetes. Pada kelas Pre-diabetes, sebanyak 3 data salah diklasifikasikan ke kelas Diabetes. Sementara itu, seluruh data pada kelas Diabetes berhasil diklasifikasikan dengan benar tanpa adanya kesalahan prediksi.

3.6 Performa Model Dengan Hyperparameter Tuning

Pada tahap ini dilakukan optimasi model menggunakan metode GridSearchCV untuk memperoleh kombinasi parameter terbaik pada algoritma Random Forest. Proses tuning dilakukan dengan menguji beberapa kombinasi parameter secara sistematis menggunakan teknik cross-validation.

Berdasarkan hasil Grid Search, diperoleh parameter terbaik Random Forest yaitu criterion = gini, max_depth = 15, min_samples_leaf = 1, max_features = sqrt, dan n_estimators = 300. Penggunaan criterion = gini memungkinkan pemisahan data berdasarkan tingkat kemurnian kelas yang optimal. Nilai max_depth = 15 memungkinkan model mempelajari pola yang cukup kompleks dengan pertumbuhan pohon yang tetap terkontrol. Parameter min_samples_leaf = 1 memungkinkan pembentukan node daun dengan minimal satu sampel, sedangkan max_features = sqrt membantu meningkatkan keragaman antar pohon dengan memilih sebagian fitur secara acak pada setiap proses pemisahan. Selain itu, nilai n_estimators = 300 dimana penggunaan 300 pohon keputusan meningkatkan stabilitas dan kemampuan generalisasi model dalam melakukan klasifikasi.

Hasil evaluasi menunjukkan bahwa model Random Forest dengan hyperparameter tuning memiliki performa yang sangat baik, dengan nilai Accuracy sebesar 98.19%, Precision sebesar 99.30%, Recall sebesar 89.91% dan F1-score sebesar 93.98%. Selain itu, terjadi peningkatan kemampuan model dalam mengenali kelas minoritas, yang ditunjukkan oleh meningkatnya nilai *recall* dan *F1-score* pada kelas tersebut

```

===== EVALUASI RF DENGAN TUNING =====
Accuracy : 98.19%
Precision : 99.30%
Recall : 89.91%
F1-Score : 93.98%

Classification Report:

              precision    recall  f1-score   support

Normal (N)       1.00      0.95      0.97        19
Pre-diabetes (P)  1.00      0.75      0.86         8
Diabetes (Y)     0.98      1.00      0.99       139

   accuracy          0.98          166
  macro avg       0.99      0.90      0.94       166
 weighted avg       0.98      0.98      0.98       166

Confusion Matrix:
[[ 18  0  1]
 [ 0  6  2]
 [ 0  0 139]]
    
```

Gambar 4. Classification report dan confusion matriks dengan haperparameter tuning

Berdasarkan hasil confusion matrix pada gambar 4, masih terdapat beberapa kesalahan klasifikasi pada model yang dihasilkan. Pada kelas Non-diabetes yang terdiri atas 19 data, terdapat 1 data yang salah diklasifikasikan sebagai Diabetes. Selanjutnya, pada kelas Pre-diabetes yang berjumlah 8 data, terdapat 2 data yang salah diklasifikasikan sebagai Diabetes. Sementara itu, seluruh data pada kelas Diabetes yang berjumlah 139 berhasil diklasifikasikan dengan benar oleh model. Hasil ini menunjukkan bahwa model memiliki kemampuan yang sangat baik dalam mengenali kelas Diabetes, namun masih mengalami kesalahan klasifikasi pada sebagian data kelas Non-diabetes dan Pre-diabetes.

3.7 Pembahasan

Hasil pengujian menunjukkan bahwa model klasifikasi menggunakan algoritma Random Forest tanpa penerapan *hyperparameter tuning* telah menghasilkan performa yang baik dengan nilai accuracy sebesar 97,59%, *precision* sebesar 99,07%, *recall* sebesar 85,75%, dan *F1-score* sebesar 90,93%. Nilai *accuracy* dan *precision* yang tinggi menunjukkan bahwa model mampu melakukan klasifikasi dengan tingkat ketepatan yang baik pada sebagian besar data.

Namun, berdasarkan hasil *confusion matrix* dan nilai *recall* yang diperoleh, masih terdapat beberapa data yang belum berhasil diklasifikasikan dengan benar, terutama pada kelas minoritas. Kesalahan klasifikasi terjadi pada kelas *Non-diabetes* dan *Pre-diabetes* yang sebagian diprediksi sebagai *Diabetes*. Kondisi ini menunjukkan bahwa model cenderung lebih mudah mengenali pola pada kelas yang memiliki jumlah data lebih banyak dibandingkan kelas yang jumlah datanya lebih sedikit.

Setelah dilakukan *hyperparameter GridSearchCV*, diperoleh kombinasi *parameter* terbaik penerapan kombinasi parameter tersebut menghasilkan peningkatan pada seluruh metrik evaluasi. Nilai *accuracy* meningkat dari 97,59% menjadi 98,19%, *precision* meningkat dari 99,07% menjadi 99,30%, *recall* meningkat dari 85,75% menjadi 89,91%, dan *F1-score* meningkat dari 90,93% menjadi *tuning* menggunakan metode 93,98%.

Meskipun peningkatan *accuracy* dan *precision* relatif kecil, peningkatan yang lebih terlihat terjadi pada nilai *recall* dan *F1-score*. Hal ini menunjukkan bahwa proses *tuning* memberikan kontribusi terhadap peningkatan kemampuan model dalam mengenali data pada kelas minoritas. Dengan kata lain, *tuning* tidak hanya meningkatkan performa model secara keseluruhan, tetapi juga membantu menghasilkan klasifikasi yang lebih seimbang pada setiap kelas.

Berdasarkan hasil *confusion matrix* setelah *tuning*, jumlah kesalahan klasifikasi pada kelas *Pre-diabetes* mengalami penurunan. Hasil ini menunjukkan bahwa kombinasi *hyperparameter* yang diperoleh melalui *GridSearchCV* mampu meningkatkan kemampuan model dalam membedakan karakteristik antar kelas, meskipun masih terdapat beberapa data yang sulit dipisahkan.

Secara keseluruhan, hasil penelitian menunjukkan bahwa algoritma *Random Forest* efektif digunakan untuk klasifikasi Diabetes Melitus. Penerapan *hyperparameter tuning* menggunakan *GridSearchCV* memberikan peningkatan performa, terutama pada kemampuan model dalam mengenali kelas minoritas yang ditunjukkan oleh peningkatan nilai *recall* dan *F1-score*. Temuan ini mengindikasikan bahwa efektivitas *hyperparameter tuning* tidak hanya diukur dari peningkatan *accuracy*, tetapi juga dari kemampuannya dalam menghasilkan model yang lebih seimbang dan memiliki kemampuan generalisasi yang lebih baik pada seluruh kelas.

4. Kesimpulan dan Saran

Berdasarkan hasil penelitian, penerapan *hyperparameter tuning* menggunakan *GridSearchCV* memberikan pengaruh positif terhadap kinerja algoritma *Random Forest* dalam klasifikasi Diabetes Melitus. Model *Random Forest* tanpa *tuning* menghasilkan nilai *accuracy* 97,59%, *precision* 99,07%, *recall* 85,75%, dan *F1-score* 90,93%, sedangkan setelah *tuning* meningkat menjadi *accuracy* 98,19%, *precision* 99,30%, *recall* 89,91%, dan *F1-score* 93,98%. Meskipun peningkatan *accuracy* dan *precision* relatif kecil, peningkatan *recall* dan *F1-score* menunjukkan bahwa *GridSearchCV* mampu meningkatkan kemampuan model dalam mengidentifikasi kelas minoritas. Dengan demikian, *hyperparameter tuning* menggunakan *GridSearchCV* terbukti dapat mengoptimalkan performa *Random Forest* dan menghasilkan model klasifikasi yang lebih seimbang untuk klasifikasi Diabetes Melitus.

Referensi

- [1] M. Banurea, D. B. Hutagao, and O. Sihombing, "Klasifikasi Penyakit Stunting Dengan Menggunakan Algoritma Support Vector Machine Dan Random Forest," *J. TEKINKOM*, vol. 6, no. 2, pp. 540–549, 2023, doi: 10.37600/tekinkom.v6i2.927.
- [2] P. Sidik, I. M. G. Sunarya, and I. G. A. Gunadi, "Comparison of Random Forest and Support Vector Machine Methods in Sentiment Analysis of Student Satisfaction Questionnaire Comments at ITB STIKOM Bali," vol. 9, no. 3, pp. 794–802, 2025.
- [3] J. Z. and J. J. K. Qi An, Saifur Rahman, "A Comprehensive Review on Machine Learning in Healthcare Industry: Classification, Restrictions, Opportunities and Challenges," *Sensors*, vol. 23, no. 9, 2023.
- [4] A. Jarmakovica, "Machine learning-based strategies for improving healthcare data quality : an evaluation of accuracy , completeness , and reusability," no. July, pp. 1–14, 2025, doi: 10.3389/frai.2025.1621514.
- [5] C. Series, "Comparative study of classification method on diagnosis of plasmodium phase Comparative study of classification method on diagnosis of plasmodium phase", doi: 10.1088/1742-6596/1516/1/012021.
- [6] L. Breiman, "No Title," pp. 1–33, 2001.
- [7] C. Arnold, "The role of hyperparameters in machine learning models and how to tune them," pp. 841–848, 2024, doi: 10.1017/psrm.2023.61.
- [8] G. Abdurrahman, H. Oktavianto, and M. Sintawati, "Optimasi Algoritma XGBoost Classifier Menggunakan Hyperparameter Gridsearch dan Random Search Pada Klasifikasi Penyakit Diabetes," *INFORMAL Informatics J.*, vol. 7, no. 3, p. 193, 2022, doi: 10.19184/isj.v7i3.35441.
- [9] S. K. Parinduri, P. Alkhairi, and H. Qurniawan, "Classification Model Optimization using Grid Search and Random Search in Machine Learning Algorithms," vol. 4, no. 2, pp. 71–78, 2025, doi: 10.61944/bids.v4i2.136.
- [10] A. A. Maitsa and N. A. S. Winarsih, "Comparison of Hyperparameter Tuning in Decision Tree and Random Forest Algorithms for Song Genre Classification," vol. 9, no. 4, pp. 1825–1834, 2025.
- [11] D. A. N. Pengobatan, "Telaah komprehensif diabetes melitus: klasifikasi, gejala, diagnosis, pencegahan, dan pengobatan," vol. 7, no. August 2020, pp. 304–317, 2021.
- [12] A. Zamani *et al.*, "Investigating the Prevalence of Undiagnosed Diabetes and Its Associated Factors Among Healthcare Workers : A Cross-sectional Study in South of Iran," pp. 413–420, 2024, doi: 10.4103/jod.jod.
- [13] I. P. Dedy, W. Darmawan, G. A. Pradnyana, I. Bagus, and N. Pascima, "Optimasi Parameter Support Vector Machine Dengan Algoritma Genetika Untuk Analisis Sentimen Pada Media Sosial Instagram," vol. 6, no. 1, pp. 58–67, 2023.
- [14] L. E. O. Breiman, "Random Forests," pp. 5–32, 2001.
- [15] A. Fatoni, D. Putra, M. N. Azmi, S. Utama, I. G. Putu, and W. Wedashwara, "Optimizing Rain Prediction Model Using Random Forest and Grid Search Cross-Validation for Agriculture Sector," vol. 23, no. 3, pp. 519–530, 2024, doi: 10.30812/matrik.v23i3.3891.