

Optimasi Algoritma Random Forest Menggunakan GridSearchCV untuk Klasifikasi Konsumsi Energi Rumah Tangga

Optimization of the Random Forest Algorithm Using GridSearchCV for Household Energy Consumption Classification

Adinda Nabila¹, Azharda Afriaci², Haya Atika Syafi'ah³

^{1,2,3}STIKOM Tunas Bangsa, Indonesia

Article Info

Genesis Artikel:

Diterima, 26 April 2026

Direvisi, 01 Juni 2026

Disetujui, 15 Juni 2026

Kata Kunci:

Smart Home
Efisiensi Energi
Random Forest
GridSearchCV
Klasifikasi

ABSTRAK

Tingginya fluktuasi penggunaan listrik serta kompleksitas sebaran data historis pada lingkungan smart home berbasis Internet of Things (IoT) sering kali memicu tingkat ketidakpastian yang signifikan dalam manajemen efisiensi daya rumah tangga. Tanpa adanya sistem prediksi yang cerdas, lonjakan konsumsi energi menjadi sulit dikendalikan secara tepat sasaran. Oleh karena itu, penelitian ini bertujuan untuk mengembangkan arsitektur model klasifikasi prediktif yang berakurasi tinggi untuk memetakan tingkat konsumsi energi ke dalam tiga kategori utama (Low, Medium, High) dengan memanfaatkan algoritma Random Forest. Proses eksperimen dijalankan secara sistematis, diawali dengan evaluasi baseline model, lalu dilanjutkan dengan tahap optimasi mengintegrasikan metode GridSearchCV berbasis 5-fold cross-validation untuk penyetelan hyperparameter secara menyeluruh terhadap IoT Smarthome Energy Dataset. Pengujian performa dievaluasi secara komprehensif menggunakan matriks konfusi dan analisis selisih akurasi (gap accuracy) untuk memastikan ketangguhan model. Hasil eksperimen membuktikan bahwa proses optimasi hyperparameter berhasil mendorong tingkat akurasi global secara signifikan dari nilai acuan awal 98,92% menjadi 99,46%, serta secara efektif mereduksi jumlah kesalahan klasifikasi. Lebih lanjut, analisis tingkat kepentingan fitur (feature importance) memberikan interpretasi ilmiah yang transparan, di mana fitur `future_consumption_kWh` teridentifikasi memberikan pengaruh paling dominan sebesar 61,75% terhadap struktur keputusan model. Kesimpulannya, arsitektur Random Forest yang telah dioptimasi ini terbukti sangat tangguh (robust) dan siap diimplementasikan sebagai instrumen penghematan energi otomatis berskala riil.

ABSTRACT

The high fluctuation in electricity usage and the complexity of historical data distribution in Internet of Things (IoT)-based smart home environments frequently trigger significant uncertainty in household energy efficiency management. Without an intelligent predictive system, energy consumption surges become exceedingly difficult to control accurately. Therefore, this study aims to develop a high-accuracy predictive classification model architecture to map energy consumption levels into three main categories (Low, Medium, High) by leveraging the Random Forest algorithm. The experimental process was conducted systematically, beginning with the evaluation of a baseline model, followed by an optimization phase integrating the GridSearchCV method based on 5-fold cross-validation for exhaustive hyperparameter tuning on the IoT Smarthome Energy Dataset. Performance was comprehensively evaluated using a confusion matrix and gap accuracy analysis to ensure model robustness. Experimental results proved that the hyperparameter optimization process successfully boosted the global accuracy rate significantly from an initial baseline of 98.92% to 99.46%, while effectively reducing the number of misclassifications. Furthermore, the feature importance analysis provided a transparent scientific interpretation, where the `future_consumption_kWh` feature was identified as having the most dominant influence at 61.75% on the model's decision structure. In conclusion, this optimized Random Forest architecture is proven to be highly robust and ready to be implemented as a real-scale automatic energy-saving instrument.

This is an open access article under the [CC BY-SA](#) license.



Penulis Korespondensi:

Azharda Afriaci,
Program Studi Sistem Informasi,
STIKOM Tunas Bangsa,
Email: azhardaafriaci@gmail.com

1. PENDAHULUAN

Tingginya ketergantungan terhadap perangkat teknologi modern telah memicu lonjakan konsumsi energi yang signifikan, baik pada sektor domestik maupun industri [1–6]. Pada lingkungan smart home berbasis Internet of Things (IoT), fluktuasi penggunaan listrik serta kompleksitas sebaran data historis memicu tingkat ketidakpastian yang tinggi, sehingga pemborosan daya kerap terjadi akibat ketiadaan sistem pemantauan yang terstruktur [7], [8]. Oleh karena itu, pemetaan dan klasifikasi pola konsumsi energi yang cerdas menggunakan pendekatan pembelajaran mesin (machine learning) mutlak diperlukan sebagai fondasi utama untuk merumuskan strategi efisiensi berkelanjutan dan menekan biaya operasional secara presisi [11–13].

Beberapa penelitian sebelumnya telah dilakukan untuk mengatasi permasalahan klasifikasi dan prediksi konsumsi energi menggunakan berbagai algoritma pembelajaran mesin. Tingkat konsumsi energi diidentifikasi menggunakan algoritma Support Vector Machine (SVM) melalui normalisasi fitur untuk meningkatkan stabilitas model [12]. Klasifikasi tingkat konsumsi energi rumah tangga juga telah dieksplorasi dengan menerapkan algoritma Decision Tree, yang mampu memberikan interpretasi aturan keputusan secara visual [13]. Selain itu, pemanfaatan algoritma Random Forest telah diimplementasikan pada sistem meteran pintar berbasis *Internet of Things (IoT)* untuk melakukan prediksi biaya dan pemantauan penggunaan listrik secara langsung [14]. Pendekatan komparatif antara Random Forest dan *Long Short-Term Memory (LSTM)* juga pernah diuji untuk melihat kemampuan adaptasi model terhadap fluktuasi beban energi jangka pendek [5]. Lebih lanjut, optimasi manajemen energi pada rumah pintar telah dikembangkan melalui model stacked ensemble yang mengombinasikan Random Forest, Support Vector Regression, dan Gradient Boosting [15].

Meskipun demikian, terdapat celah teoretis dan praktis dalam penelitian-penelitian tersebut yang perlu disempurnakan. Penggunaan algoritma tunggal seperti SVM dan Decision Tree terbukti rentan mengalami ketidakstabilan akurasi dan overfitting ketika dihadapkan pada variasi data tabular IoT yang kompleks [12], [13], sementara penerapan model regresi gagal memetakan kategori konsumsi ke dalam label klasifikasi terstruktur yang dibutuhkan untuk otomatisasi keputusan [14]. Lebih jauh, algoritma ensemble konvensional seperti Random Forest memiliki keterbatasan mendasar ketika memproses data tabular berseri waktu (sequential tabular data) [16]. Jika pohon keputusan (decision trees) dibiarkan tumbuh menggunakan parameter bawaan (default), model cenderung menghafal noise dari fluktuasi jangka pendek dan gagal merangkum pola laten struktural [5], [17]. Kondisi ini diperparah oleh tingginya korelasi linier antar-fitur pada lingkungan sensor cerdas, yang sering kali memicu redundansi aturan percabangan (node splitting) jika ruang pencarian parameter tidak dibatasi secara matematis [18]. Walaupun penggunaan arsitektur stacked ensemble yang lebih rumit dapat mengatasi masalah ini, hal tersebut justru memicu pembengkakan beban komputasi dan latensi inferensi yang sangat tidak ideal untuk lingkungan smart home dengan keterbatasan sumber daya edge-computing [15], [19].

Berangkat dari celah teoretis tersebut, kebaruan (novelty) dari penelitian ini terletak pada perancangan arsitektur klasifikasi energi yang memaksimalkan efisiensi komputasi tanpa mengorbankan ketepatan prediksi pada IoT Smarthome Energy Dataset. Dalam penelitian ini, GridSearchCV tidak hanya diposisikan sebagai fungsi mekanis, melainkan dieksploitasi secara strategis sebagai mekanisme kontrol penalti untuk membatasi kompleksitas arsitektur Random Forest. Dengan menala parameter struktural (max_depth dan kriteria node splitting), model dipaksa untuk mengabaikan dependensi sekuensial yang bias dan fokus pada fitur prediktor utama. Pendekatan ini bertujuan untuk menghasilkan model klasifikasi tiga tingkat (*Low, Medium, High*) yang tidak hanya akurat dan kebal terhadap overfitting, tetapi juga sangat ringan untuk dieksekusi, sehingga siap diimplementasikan sebagai instrumen penghematan energi berskala riil pada perangkat IoT.

2. METODE PENELITIAN

Pada bagian ini, dijabarkan mengenai metodologi dan tahapan sistematis yang digunakan dalam pelaksanaan penelitian untuk mengklasifikasikan kategori konsumsi energi. Prosedur pengujian dan pemrosesan data disusun secara kronologis, dimulai dari pemahaman karakteristik awal dataset (*Data Understanding*), persiapan serta transformasi data (*Data Preparation*), pembangunan model berbasis algoritma *Random Forest (Modeling)*, hingga pengujian performa (*Evaluation*). Melalui pendekatan yang terstruktur ini, eksperimen pengolahan data dan optimasi parameter menggunakan metode *GridSearchCV* dilakukan secara ilmiah guna menghasilkan model klasifikasi yang akurat dan stabil.

2.1. Alur Penelitian (Research Framework)

Tahapan pelaksanaan eksperimen untuk klasifikasi kategori konsumsi energi dirancang secara sistematis dan kronologis seperti yang ditunjukkan pada Gambar 1. Alur penelitian tersebut diawali dengan pelaksanaan studi literatur guna mendalami landasan teori dan pemilihan algoritma, yang kemudian dilanjutkan dengan pemuatan dataset *energy_categorical.csv* ke dalam lingkungan pengembangan. Setelah dataset berhasil dimuat, tahap data understanding dilaksanakan untuk menganalisis karakteristik awal data melalui pengecekan kualitas, persebaran distribusi target, serta korelasi antar-fitur. Proses tersebut diikuti oleh tahap data preparation yang di dalamnya dilakukan penghapusan fitur tidak relevan, transformasi label encoding pada variabel target, serta pembagian data menjadi data training sebesar 80% dan data testing sebesar 20%. Pada tahap inti pemodelan, algoritma Random Forest dipilih karena dinilai mampu menangani data klasifikasi dengan baik, memiliki performa tinggi pada data tabular, serta relatif stabil terhadap gangguan noise dan outlier [18–22]. Pelatihan model dilakukan secara bertahap mulai dari pembentukan baseline model menggunakan parameter bawaan (default), dilanjutkan dengan proses optimasi hyperparameter menggunakan metode *GridSearchCV* guna mendapatkan kombinasi parameter paling optimal. Metode *GridSearchCV* ini dipilih

karena memiliki kemampuan untuk menguji seluruh kombinasi parameter secara menyeluruh (exhaustive search) berbasis cross-validation otomatis, sehingga dapat meminimalkan risiko overfitting serta memastikan ditemukannya kombinasi nilai parameter yang menghasilkan akurasi paling optimal[23-27]. Rangkaian penelitian ini kemudian ditutup dengan tahap evaluasi menggunakan matriks performa dan confusion matrix guna memastikan model final mampu mengklasifikasikan kategori konsumsi energi secara akurat, efisien, dan stabil.



Gambar 1. Alur Penelitian

Melalui kerangka kerja yang terstruktur pada Gambar 1, seluruh rangkaian eksperimen dapat dijalankan secara konsisten dan terukur. Penjelasan lebih mendalam mengenai implementasi teknis dari masing-masing tahapan tersebut, mulai dari pemahaman karakteristik awal data hingga mekanisme pengujian performa model, akan dijabarkan secara rinci pada sub-bab berikutnya. Penerapan alur penelitian yang sistematis ini berfungsi sebagai panduan utama untuk memastikan bahwa setiap tahapan pengolahan data dan optimasi hyperparameter dilakukan secara ilmiah demi menghasilkan model klasifikasi konsumsi energi yang optimal.

2.2. Gambaran Umum dan Kualitas Dataset (*Data Understanding*)

Tahap pemahaman data (data understanding) dilakukan untuk mengenali karakteristik fisik dan kualitas dari data mentah sebelum masuk ke tahap pra-pemrosesan. Dataset yang digunakan dalam penelitian ini adalah IoT Smarthome Energy Dataset yang diperoleh secara publik dari repositori Kaggle [30]. Data tersebut disimpan dalam berkas bernama `energy_categorical.csv` dengan dimensi total terdiri dari 921 baris data dan 5 kolom atribut. Atribut di dalam dataset ini mengombinasikan tipe data numerik (integer dan float) serta tipe data objek berupa teks. Sebagai representasi awal, contoh susunan data mentah beserta sampel nilai dari setiap variabel disajikan secara terstruktur pada Tabel 1.

Tabel 1. Dataset Konsumsi Energi

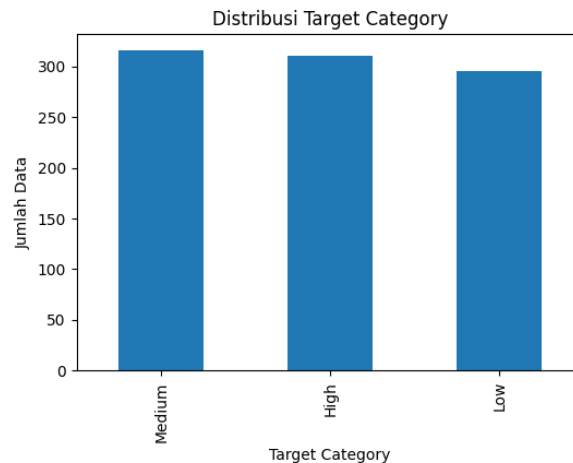
No	timestamp	household_id	energy_consumption_kWh	future_consumption_kWh	target_category
1	01/01/2024 00:00	1	0,2087498419693580	0,23858768823443900	Medium
2	01/01/2024 00:05	1	0,23858768823443900	0,25682963193147400	High
3	01/01/2024 00:10	1	0,25682963193147400	0,20664288131704500	Medium
4	01/01/2024 00:15	1	0,20664288131704500	0,24353502755720400	Medium
5	01/01/2024 00:20	1	0,24353502755720400	0,24218758514834100	Medium
6	01/01/2024 00:25	1	0,24218758514834100	0,28392404075385400	High
7	01/01/2024 00:30	1	0,28392404075385400	0,2400335776801710	Medium
8	01/01/2024 00:35	1	0,2400335776801710	0,248543470699992	Medium
9	01/01/2024 00:40	1	0,248543470699992	0,3100330665076770	High
10	01/01/2024 00:45	1	0,3100330665076770	0,31471194086069200	High
11	01/01/2024 00:50	1	0,31471194086069200	0,27150967957474600	High
12	01/01/2024 00:55	1	0,27150967957474600	0,28753359483913100	High
13	01/01/2024 01:00	1	0,28753359483913100	0,3169914236893640	High
14	01/01/2024 01:05	1	0,3169914236893640	0,25739681994405200	High
15	01/01/2024 01:10	1	0,25739681994405200	0,3044528817195420	High
16	01/01/2024 01:15	1	0,3044528817195420	0,2864733423015070	High
17	01/01/2024 01:20	1	0,2864733423015070	0,28255626414450300	High
18	01/01/2024 01:25	1	0,28255626414450300	0,3066618071966650	High
19	01/01/2024 01:30	1	0,3066618071966650	0,2983365533175360	High
...	...	...	...	...	...
921	01/01/2024 04:05	1	0,10667493899786500	0,12649601419716100	Low

Analisis deskriptif awal dilakukan terhadap seluruh atribut dataset untuk memberikan pemahaman menyeluruh mengenai distribusi nilai serta karakteristik statistik dasar data pemakaian daya listrik. Rangkuman hasil statistik deskriptif dari pengujian kode untuk fitur numerik maupun kategorikal tersebut disajikan secara rinci pada Tabel 2. Berdasarkan akumulasi data pada Tabel 2, diidentifikasi bahwa dataset memiliki total 921 baris pengamatan dengan nilai rata-rata (mean) konsumsi energi saat ini (*energy_consumption_kWh*) sebesar 0,2006 kWh dan konsumsi energi berikutnya (*future_consumption_kWh*) sebesar 0,2006 kWh. Rentang pemakaian daya pada kedua fitur numerik tersebut terpantau cukup kecil dan sangat stabil, yang ditandai dengan nilai minimum sebesar 0,0528 kWh serta nilai maksimum mencapai 0,3409 kWh. Sementara itu, pada variabel target *target_category*, terdeteksi adanya tiga nilai unik (unique values) dengan proporsi kategori pemakaian yang paling dominan diwakili oleh tingkat Medium sebanyak 316 data. Pemetaan karakteristik statistik awal melalui Tabel 2 ini sangat penting dilakukan untuk mendeteksi perbedaan rentang nilai antar-fitur sebelum melangkah ke tahapan persiapan data (data preparation).

Tabel 2. Statistik Deskriptif Awal Dataset

Matriks Statistik	timestamp	household_id	energy_consumption_kWh	future_consumption_kWh	target_category
count	921	921.000.000	921.000.000	921.000.000	921
unique	185	NaN	NaN	NaN	3
top	01/01/2024 00:00	NaN	NaN	NaN	Medium
freq	5	NaN	NaN	NaN	316
mean	NaN	3.002.172	0.200649	0.200636	NaN
std	NaN	1.415.748	0.074028	0.074016	NaN
min	NaN	1.000.000	0.052868	0.052868	NaN
25%	NaN	2.000.000	0.134713	0.134713	NaN
50%	NaN	3.000.000	0.200601	0.200041	NaN
75%	NaN	4.000.000	0.272544	0.272544	NaN
max	NaN	5.000.000	0.340917	0.340917	NaN

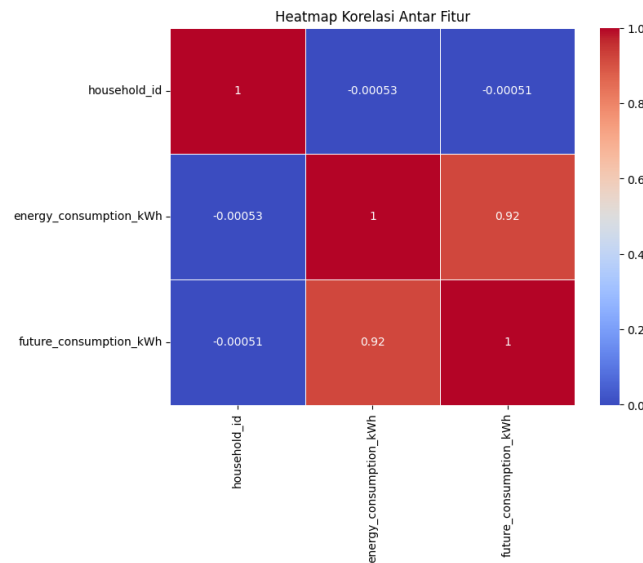
Sebelum tahapan eksperimen dilanjutkan ke fase pra-pemrosesan, pengujian terhadap kualitas dan kebersihan data mutlak dilakukan untuk menjamin validitas hasil klasifikasi. Berdasarkan hasil pemeriksaan program, diperoleh konfirmasi bahwa dataset *energy_categorical.csv* sepenuhnya bersih karena memiliki 0 missing value (tidak ada data kosong) dan 0 data duplikat pada seluruh baris dan variabelnya. Setelah kualitas data dipastikan bersih, analisis dilanjutkan dengan memeriksa distribusi penyebaran kelas pada variabel *target_category* guna mendeteksi ada atau tidaknya indikasi ketimpangan data (class imbalance). Hasil visualisasi mengenai proporsi tersebut disajikan secara rinci pada Gambar 2. Melalui grafik tersebut, dapat diidentifikasi bahwa dataset ini memiliki karakteristik kelas yang seimbang (balanced dataset), dengan rincian kelas Medium sebanyak 316 data (34,31%), kelas High sebanyak 310 data (33,66%), dan kelas Low sebanyak 295 data (32,03%). Karakteristik data yang seimbang ini sangat menguntungkan bagi proses pemodelan karena algoritma Random Forest dapat mempelajari parameter setiap kategori secara objektif tanpa risiko bias, sehingga penelitian ini tidak memerlukan tahapan penanganan ketimpangan kelas tambahan seperti metode pencuplikan ulang (resampling).



Gambar 2. Grafik Distribusi Target Category

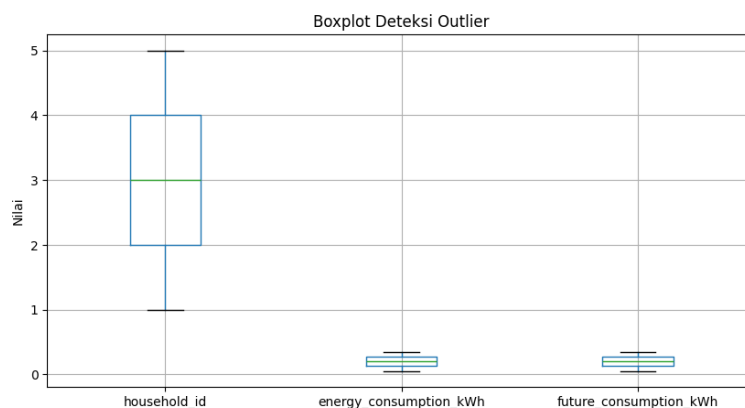
Analisis keterkaitan dan hubungan linear antar-fitur numerik di dalam dataset dilakukan memanfaatkan matriks korelasi Pearson guna mengetahui tingkat pengaruh dan kontribusi suatu variabel terhadap variabel lainnya. Hasil perhitungan koefisien korelasi tersebut disajikan secara visual melalui grafik heatmap pada Gambar 3. Berdasarkan visualisasi matriks tersebut, ditemukan adanya hubungan linear positif yang sangat kuat antara variabel tingkat konsumsi energi saat ini (*energy_consumption_kWh*) dan tingkat konsumsi energi di masa mendatang (*future_consumption_kWh*) dengan nilai koefisien korelasi mencapai 0,92. Kedekatan nilai koefisien ini dengan angka 1,00 mengindikasikan bahwa volume pemakaian daya listrik pada periode sekarang memiliki pengaruh searah yang sangat signifikan terhadap proyeksi pemakaian energi pada periode berikutnya. Sebaliknya, atribut *household_id* menunjukkan nilai koefisien korelasi yang sangat tidak signifikan karena

mendekati nilai nol, yaitu sebesar -0,00053 terhadap fitur `energy_consumption_kWh` dan -0,00051 terhadap fitur `future_consumption_kWh`. Nilai yang mendekati nol ini membuktikan bahwa identitas unik rumah tangga tidak memiliki keterkaitan fungsional maupun pengaruh terhadap pola penggunaan daya listrik pengguna, sehingga informasi ini menjadi landasan dasar empiris untuk mengeliminasi atribut `household_id` pada tahapan persiapan data berikutnya.



Gambar 3. Heatmap Korelasi Antar Fitur

Tahap akhir dalam pengujian kualitas data pada fase data understanding adalah mengidentifikasi keberadaan data pencilan (outlier) menggunakan metode visualisasi boxplot. Evaluasi terhadap nilai ekstrem ini mutlak dilakukan karena karakteristik persebaran data yang menyimpang secara radikal dapat memengaruhi objektivitas pembentukan simpul percabangan pada pohon keputusan algoritma Random Forest. Hasil pemetaan sebaran data untuk mendeteksi pencilan tersebut disajikan secara komprehensif pada Gambar 4. Berdasarkan grafik boxplot tersebut, dapat diidentifikasi bahwa seluruh variabel numerik, baik `energy_consumption_kWh` maupun `future_consumption_kWh`, memiliki sebaran nilai yang mengumpul secara stabil di dalam rentang interquartile range (IQR) dan batas garis whisker. Tidak ditemukannya representasi titik data terisolasi di luar batas atas maupun batas bawah whisker membuktikan secara empiris bahwa dataset ini sepenuhnya terbebas dari gangguan data pencilan. Kondisi data yang bersih dan konsisten ini memastikan stabilitas varians data tetap terjaga, sehingga eksperimen dapat langsung dilanjutkan tanpa harus melakukan prosedur pemangkasan data (trimming) ataupun transformasi nilai penanganan pencilan tambahan.



Gambar 4. Boxplot Deteksi Outlier

Secara keseluruhan, berdasarkan seluruh rangkaian pengujian mendalam terhadap kualitas, struktur, korelasi antar-fitur, serta pola distribusi kelas target yang telah dipaparkan, dataset `energy_categorical.csv` dinyatakan valid dan memenuhi seluruh kriteria kelayakan ilmiah. Karakteristik data yang sepenuhnya bersih dari nilai kosong (missing value) dan data pencilan (outlier), serta didukung oleh proporsi kelas target yang seimbang, memberikan fondasi empiris yang kokoh bagi algoritma pemodelan untuk menghasilkan keputusan klasifikasi yang objektif dan bebas dari risiko bias. Oleh karena itu, tahapan pemahaman data (data understanding) pada fase ini disimpulkan telah selesai dengan hasil yang optimal, sehingga dataset dinyatakan sangat layak untuk diteruskan ke tahapan eksperimen berikutnya, yaitu sub-bab 2.3 Persiapan Data (Data Preparation).

2.3. Persiapan Data (Data Preparation)

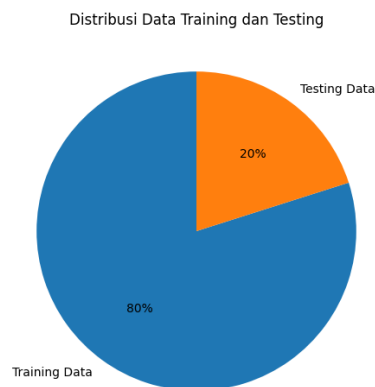
Tahap awal pada fase persiapan data (data preparation) dilakukan melalui prosedur seleksi fitur (feature selection) guna mengeliminasi atribut yang tidak memiliki kontribusi fungsional terhadap performa pemodelan. Pada tahapan ini, kolom timestamp secara resmi dihapus dari susunan dataset karena hanya bertindak sebagai pencatat kronologis waktu eksternal yang tidak memuat pola prediktif dalam penentuan kategori penggunaan daya listrik. Sebaliknya, atribut `household_id`, `energy_consumption_kWh`, dan `future_consumption_kWh` secara konsisten tetap dipertahankan di dalam matriks data. Keputusan untuk mempertahankan ketiga variabel tersebut didasarkan pada kebutuhan eksperimen yang memposisikannya sebagai fitur prediktor utama (variabel independen), sehingga model klasifikasi Random Forest dapat mempelajari karakteristik ID rumah tangga sekaligus volume konsumsi energi saat ini dan masa depan secara simultan.

Setelah prosedur seleksi fitur dan transformasi label encoding selesai dilakukan, susunan matriks dataset mengalami perubahan representasi nilai. Hasil akhir dari pemrosesan awal ini mempertahankan tiga fitur prediktor utama serta mengubah kategori teks pada label target menjadi bentuk numerik murni agar dapat diproses oleh pustaka machine learning. Sampel lima baris pertama dari dataset yang telah melewati tahapan persiapan data tersebut disajikan secara terstruktur pada Tabel 3. Melalui representasi data pada Tabel 3, dapat diidentifikasi bahwa atribut timestamp telah berhasil dieliminasi dari sistem, sedangkan variabel `target_category` telah sepenuhnya dikonversi menjadi skala numerik murni, di mana nilai 1 merepresentasikan kategori Medium dan nilai 2 merepresentasikan kategori High.

Tabel 3. Sampel Dataset Hasil Persiapan Data

No	household_id	energy_consumption_kWh	future_consumption_kWh	target_category
1	1	0.208750	0.238588	1
2	1	0.238588	0.256830	2
3	1	0.256830	0.206643	1
4	1	0.206643	0.243535	1
5	1	0.243535	0.242188	1
...	...	...	...	...
921	1	0,106675	0,126496	0

Prosedur akhir pada tahapan persiapan data dilakukan melalui pembagian dataset (data splitting) menjadi subset terpisah untuk kebutuhan pelatihan dan pengujian model klasifikasi. Pemisahan data ini difungsikan agar performa evaluasi algoritma Random Forest dapat diukur secara objektif memanfaatkan potongan data independen yang belum pernah dipelajari sebelumnya selama fase pelatihan. Rasio pembagian data yang diterapkan dalam eksperimen ini adalah sebesar 80% untuk data training (mencakup 736 baris data) dan 20% untuk data testing (mencakup 185 baris data). Persentase pembagian alokasi data tersebut disajikan secara visual melalui grafik lingkaran (pie chart) pada Gambar 5. Proses partisi data ini dieksekusi secara sistematis menggunakan bantuan parameter konfigurasi stratify yang merujuk langsung pada variabel target. Penerapan mekanisme stratification tersebut bertujuan mutlak untuk memastikan bahwa proporsi sebaran kelas target (Low, Medium, dan High) pada potongan data training maupun data testing tetap terjaga secara seimbang sama persis dengan karakteristik dataset aslinya, sehingga risiko penurunan performa model akibat ketimpangan pembagian data dapat dihindari sepenuhnya.



Gambar 5. Pembagian Data Training dan Testing

Melalui penyelesaian seluruh rangkaian prosedur pra-pemrosesan tersebut, fase persiapan data (data preparation) secara resmi dinyatakan selesai dengan hasil yang optimal. Dataset kini telah dikonversi secara utuh ke dalam bentuk struktur matriks numerik terpisah yang siap diolah oleh algoritma machine learning, yang memuat matriks fitur data pelatihan (`Xtrain`), matriks fitur data pengujian (`Xtest`), vektor label target pelatihan (`ytrain`), dan vektor label target pengujian (`ytest`). Pemisahan susunan data ini memastikan bahwa proses pelatihan dan pengujian model nantinya dapat berjalan secara independen dan terukur tanpa adanya kebocoran data (data leakage). Dengan ketersediaan matriks data yang telah bersih dan terstandarisasi ini, seluruh variabel prediktor dan target dinyatakan telah siap sepenuhnya untuk diumpankan ke tahap komputasi berikutnya, yaitu sub-bab 2.4 Pembangunan Model Klasifikasi (Modeling).

2.4. Evaluasi Model dan Pengujian Statistik

Pada tahap inti pemodelan, algoritma Random Forest Classifier dipilih sebagai arsitektur ensemble learning utama berbasis bagging untuk mengklasifikasikan kategori konsumsi energi. Guna mengatasi kelemahan model baseline bawaan yang rentan terhadap saturasi atau overfitting pada dimensi data IoT tertentu, interkoneksi struktural antar-pohon dioptimasi secara terintegrasi menggunakan metode GridSearchCV. Optimasi difokuskan pada manipulasi ruang pencarian parameter (search space grid) yang secara langsung memengaruhi kompleksitas topologi pohon keputusan, penyebaran varians, dan kriteria pembatasan daun. Rincian kombinasi dimensi parameter yang dieksplorasi disajikan secara transparan pada Tabel 4.

Tabel 4. Ruang Pencarian Hyperparameter (*Search Space Grid*)

Hyperparameter	Deskripsi Fungsional	Rentang Nilai yang Diuji
n_estimators	Jumlah pohon keputusan di dalam ensemble	{100, 150, 200, 250, 300}
max_depth	Batas kedalaman maksimal setiap pohon	{None, 10, 20, 30}
min_samples_split	Jumlah minimum sampel untuk memecah simpul	{2, 5, 10}
min_samples_leaf	Jumlah minimum sampel pada simpul daun (leaf)	{1, 2, 4}

Melalui kombinasi rentang nilai pada Tabel 5, GridSearchCV melakukan eksplorasi secara exhaustive menggunakan skema 5-fold cross-validation. Pemilihan nilai 5-fold (dibandingkan 10-fold) secara ilmiah didasarkan pada titik ekuilibrium yang optimal antara representasi varians data dan efisiensi waktu eksekusi komputasi[31]. Mengingat dimensi IoT Smarthome Energy Dataset berskala menengah (921 sampel), penggunaan 5-fold sudah sangat memadai untuk mencegah bias validasi sekaligus menjaga agar latensi pencarian parameter tetap ideal pada lingkungan perangkat IoT yang memiliki keterbatasan sumber daya.

Algoritma 1. Mekanisme GridSearchCV pada Model Random Forest:

INPUT :

Dataset Latih $S = (X_{\text{train}}, y_{\text{train}})$

Kamus Parameter $P = \{n_{\text{estimators}}, \text{max_depth}, \text{min_samples_split}, \text{min_samples_leaf}\}$

Jumlah Fold $K = 5$ (Stratified)

OUTPUT:

Model Random Forest Optimal (M_{best})

```

1: Inisialisasi  $M_{\text{base}} = \text{RandomForestClassifier}()$ 
2:  $\text{best\_score} = 0$ 
3:  $\text{best\_params} = \{\}$ 
4:
5: UNTUK SETIAP kombinasi parameter  $p$  di dalam ruang Cartesian  $P$  LAKUKAN
6:   Array skor_validasi = []
7:   Partisi himpunan  $S$  menjadi  $K$  himpunan bagian (fold) yang saling lepas
8:
9:   UNTUK SETIAP fold  $k$  dari 1 hingga  $K$  LAKUKAN
10:      $S_{\text{val}} =$  himpunan bagian ke- $k$ 
11:      $S_{\text{train}} = S \setminus S_{\text{val}}$  (seluruh himpunan  $S$  dikurangi  $S_{\text{val}}$ )
12:
13:      $M_{\text{temp}} =$  Latih  $M_{\text{base}}$  menggunakan parameter  $p$  pada  $S_{\text{train}}$ 
14:     Akurasi_ $k =$  Evaluasi  $M_{\text{temp}}$  pada  $S_{\text{val}}$ 
15:     Tambahkan Akurasi_ $k$  ke dalam skor_validasi
16: AKHIR UNTUK
17:
18: rata_rata_skor = Rata-rata dari skor_validasi
19: JIKA rata_rata_skor > best_score MAKA
20:   best_score = rata_rata_skor
21:   best_params =  $p$ 
22: AKHIR JIKA
23: AKHIR UNTUK
24:
25:  $M_{\text{best}} =$  Latih RandomForestClassifier menggunakan best_params pada seluruh  $S$ 
26: KEMBALIKAN  $M_{\text{best}}$ 

```

Representasi prosedural pada Algoritma 1 difungsikan untuk membongkar mekanisme internal dari integrasi algoritma Random Forest dan GridSearchCV, guna memberikan transparansi algoritmik yang lebih mendalam dibandingkan sekadar deskripsi konseptual. Alur logika diawali dengan mendefinisikan ruang pencarian hyperparameter sebagai ruang Cartesian (P), di mana setiap kombinasi parameter dievaluasi secara iteratif (baris 5).

Untuk memastikan bahwa performa model tidak bias terhadap partisi data tertentu, skema Stratified 5-Fold Cross-Validation diimplementasikan. Pada setiap iterasi parameter, dataset pelatihan dipartisi menjadi 5 himpunan bagian yang saling lepas (baris 7). Model sementara kemudian dilatih berulang kali menggunakan 4 himpunan sebagai data latih dan 1 himpunan sebagai data validasi secara bergantian (baris 9-16). Rata-rata dari skor validasi ini kemudian dievaluasi; apabila skor kumulatif tersebut melampaui nilai rekor sebelumnya, maka sistem secara otomatis akan memperbarui konfigurasi parameter terbaik (baris 19-22). Pada akhir iterasi, arsitektur model akhir direkonstruksi dan dilatih ulang menggunakan kombinasi parameter paling optimal tersebut pada keseluruhan dataset latih awal (baris 25-26). Melalui rumusan mekanis ini, risiko terjadinya overfitting selama proses pencarian parameter dapat ditekan secara matematis.

Tahap akhir dari rangkaian metodologi penelitian ini adalah melakukan evaluasi mendalam terhadap performa prediksi model Random Forest sebelum dan sesudah dioptimasi. Pengukuran kinerja klasifikasi dieksekusi secara komprehensif menggunakan matriks evaluasi standar multi-kelas, yang mencakup Accuracy, Precision, Recall, dan F1-Score guna memastikan ketepatan dan sensitivitas prediksi pada masing-masing kategori target (Low, Medium, High Lebih lanjut, guna memenuhi standar validasi matematis yang ketat, penelitian ini mengimplementasikan Uji Signifikansi Statistik berupa Uji T Berpasangan (Paired T-Test) yang dieksekusi berbasis skor evaluasi 5-Fold Cross-Validation. Uji statistik ini difungsikan untuk membuktikan secara empiris apakah perbedaan performa prediktif antara model baseline dan model optimized benar-benar konsisten di berbagai lipatan data secara signifikan, atau sekadar terjadi karena faktor kebetulan (random chance). Selain pengujian signifikansi stabilitas akurasi, efisiensi model juga dievaluasi melalui pencatatan waktu eksekusi komputasi (computational runtime) dalam satuan milidetik (ms). Parameter latensi komputasi ini sangat krusial untuk memvalidasi kelayakan implementasi model pada lingkungan perangkat edge-IoT yang memiliki keterbatasan sumber daya (resource-constrained devices).

Prosedur evaluasi performa model Random Forest dilakukan melalui tiga instrumen utama. Pertama, penerapan Confusion Matrix untuk membandingkan label aktual dan prediksi guna mengidentifikasi misklasifikasi secara spesifik pada tiap kelas target. Kedua, pengujian indikasi overfitting melalui perhitungan selisih (gap) antara akurasi training dan testing, di mana margin yang mendekati nol membuktikan bahwa model bersifat tangguh (robust) dan mampu menggeneralisasi data dengan baik. Ketiga, analisis feature importance berbasis penurunan Gini impurity untuk mengukur kontribusi dominan dari masing-masing variabel prediktor terhadap keputusan klasifikasi. Kesiapan seluruh instrumen evaluasi komprehensif ini secara resmi menutup rangkaian metodologi pada BAB II: METODE PENELITIAN, dan menjadi landasan untuk eksekusi komputasi serta analisis hasil eksperimen pada bab selanjutnya.

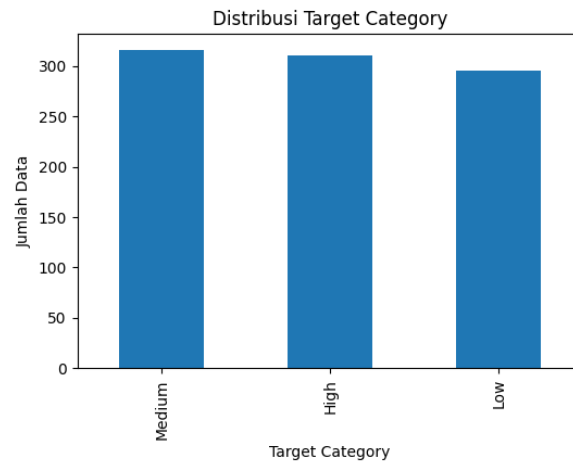
3. HASIL DAN PEMBAHASAN

Pada bab ini, seluruh temuan empiris dan hasil eksekusi komputasi dari rangkaian metodologi penelitian yang telah dirancang pada bab sebelumnya akan dipaparkan dan dibahas secara komprehensif. Penyajian hasil eksperimen ini disusun secara sistematis yang diawali dengan analisis hasil eksplorasi data (data understanding) untuk mengenali karakteristik awal serta sebaran kelas pada dataset konsumsi energi. Selanjutnya, bagian pembahasan akan menguraikan luaran konkret dari fase pra-pemrosesan data (data preparation) termasuk dimensi akhir data setelah dilakukan pemisahan. Bab ini juga akan mengulas secara mendalam mengenai perbandingan performa klasifikasi antara eksploitasi baseline model dengan model yang telah dioptimasi menggunakan penyetelan hyperparameter berbasis GridSearchCV. Sebagai penutup, analisis mengenai deteksi risiko overfitting serta tingkat kontribusi variabel (feature importance) akan dijabarkan guna memberikan interpretasi ilmiah yang transparan mengenai keandalan algoritma Random Forest yang dikembangkan.

3.1. Analisis Hasil Eksplorasi Data

Tahapan analisis pada bab ini diawali dengan memaparkan karakteristik fisik serta struktur fundamental dari IoT Smarthome Energy Dataset setelah berkas data mentah berhasil dimuat secara utuh ke dalam lingkungan komputasi. Berdasarkan hasil eksekusi pemrosesan awal terhadap berkas `energy_categorical.csv`, dataset ini terbukti memiliki dimensi riil yang terdiri dari 921 baris sampel data dan 5 kolom atribut. Melalui prosedur pemeriksaan kualitas dan integritas data yang dilakukan secara menyeluruh, dikonfirmasi bahwa tidak ditemukan adanya data kosong (missing values) ataupun nilai yang hilang pada seluruh baris dan kolom variabel. Ketiadaan missing values ini menunjukkan bahwa data memiliki tingkat kebersihan yang tinggi dan siap untuk diproses tanpa memerlukan tindakan manipulasi imputasi nilai. Ditinjau dari aspek karakteristik tipe data, dataset ini mengombinasikan tipe data numerik berupa integer dan float untuk merepresentasikan variabel kuantitatif konsumsi daya, serta tipe data object berupa nilai tekstual untuk mengakomodasi variabel kategorikal, termasuk di dalamnya adalah variabel label target yang memuat kelas klasifikasi energi.

Langkah analisis eksplorasi data dilanjutkan dengan memeriksa visualisasi penyebaran variabel respons atau label target melalui diagram batang (bar chart) yang mencerminkan distribusi dari setiap kategori konsumsi energi. Berdasarkan visualisasi grafik yang disajikan pada Gambar 7, dapat dijabarkan secara rinci bahwa kategori kelas Medium memiliki frekuensi tertinggi di dalam dataset dengan jumlah sebanyak 316 sampel data. Selanjutnya, kategori kelas High menempati posisi kedua dengan total sebaran sebanyak 310 sampel data, yang kemudian diikuti secara ketat oleh kategori kelas Low dengan kuantitas sebanyak 295 sampel data. Sebaran nilai riil ini menegaskan bahwa dari keseluruhan dimensi data yang diuji, masing-masing kategori target telah terepresentasi secara proporsional.



Gambar 6. Grafik Distribusi Kategori Kelas Target Konsumsi Energi

Kondisi sebaran kuantitas yang seragam antar-kategori tersebut mengindikasikan bahwa IoT Smarthome Energy Dataset yang digunakan dalam penelitian ini merupakan sebuah dataset yang sangat seimbang (highly balanced dataset). Karakteristik distribusi data yang seimbang ini memberikan keuntungan metodologis yang sangat besar dalam fase pembangunan model klasifikasi. Kondisi ini memastikan bahwa algoritma Random Forest dapat mempelajari seluruh karakteristik pola fitur prediktor dari setiap variasi kelas secara objektif tanpa adanya risiko bias keputusan yang umumnya condong pada kelas mayoritas tertentu (majority class bias). Lebih lanjut, dengan struktur data yang balanced ini, metrik evaluasi global seperti nilai akurasi (accuracy) yang diperoleh pada tahap pengujian nantinya dapat diinterpretasikan secara langsung sebagai representasi performa riil model yang valid tanpa memerlukan manipulasi algoritma penyeimbang data tambahan seperti SMOTE.

3.1. Hasil Pra-pemrosesan Data

Memasuki tahapan berikutnya, proses pra-pemrosesan data (data preparation) dilakukan untuk menjamin bahwa struktur data telah berada dalam kondisi optimal sebelum dialokasikan ke dalam algoritma pembelajaran. Pada tahap ini, transformasi data difokuskan pada konversi variabel label target yaitu `target_category` yang awalnya bersifat kategorikal tekstual (Low, Medium, High) menjadi representasi nilai numerik. Proses transformasi ini dieksekusi menggunakan fungsi `LabelEncoder` dari pustaka `Scikit-Learn`, di mana kelas Low ditransformasikan menjadi representasi nilai 0, kelas Medium menjadi nilai 1, dan kelas High menjadi nilai 2. Penerapan encoding ini sangat krusial untuk menghindari terjadinya galat komputasi serta mempermudah model Random Forest dalam melakukan kalkulasi matriks jarak dan pembentukan simpul keputusan secara matematis. Sementara itu, untuk variabel prediktor numerik, dilakukan verifikasi skala guna memastikan seluruh nilai variabel berada pada rentang batas yang valid.

Setelah tahapan transformasi variabel selesai diimplementasikan, prosedur dilanjutkan dengan melakukan pembagian dataset (data splitting) menjadi subset data latih (training data) dan data uji (testing data). Proses pembagian data ini dieksekusi menggunakan fungsi `train_test_split` dengan menerapkan rasio proporsi sebesar 80% untuk alokasi data latih dan 20% untuk alokasi data uji, serta dilengkapi dengan parameter `random_state` untuk menjamin konsistensi reproduksibilitas pembagian data. Berdasarkan total keseluruhan data yang berjumlah 921 baris, diperoleh sebaran dimensi akhir yang konkret yaitu sebanyak 736 baris data digunakan sebagai basis pembelajaran model, dan 185 baris data difungsikan sebagai instrumen pengujian validitas model dengan masing-masing memuat 3 variabel prediktor kunci. Rincian detail mengenai dimensi matriks, jumlah baris, serta jumlah variabel untuk komponen `X_train`, `X_test`, `y_train`, dan `y_test` disajikan secara terstruktur pada Tabel 4.

Tabel 4. Spesifikasi Dimensi Akhir Subset Data Latih dan Data Uji

Komponen Dataset	Fungsi Komponen	Jumlah Baris (Sampel)	Jumlah Kolom (Atribut)
<code>X_train</code>	Fitur Prediktor Data Latih	736	3
<code>y_train</code>	Variabel Target Data Latih	736	1
<code>X_test</code>	Fitur Prediktor Data Uji	185	3
<code>y_test</code>	Variabel Target Data Uji	185	1

Melalui penyelesaian seluruh prosedur transformasi label dan pembagian dimensi dataset tersebut, fase penyajian hasil pra-pemrosesan data (data preparation) secara resmi dinyatakan selesai dengan hasil yang valid. Kepastian kuantitas sampel dan atribut yang tertera pada Tabel 4 membuktikan bahwa seluruh struktur data telah berada dalam kondisi ajek serta terbebas dari risiko kesalahan dimensi (dimension mismatch) saat dieksekusi oleh program. Dengan ketersediaan matriks fitur dan vektor target yang telah terstandarisasi ini, seluruh subset data pelatihan (`Xtrain` dan `ytrain`) serta data pengujian (`Xtest` and `ytest`) dinyatakan siap sepenuhnya untuk diumpankan sebagai masukan komputasi pada tahap eksperimen berikutnya, yaitu sub-bab 3.3 Kinerja Eksperimen Baseline Model (Model Acuan Awal).

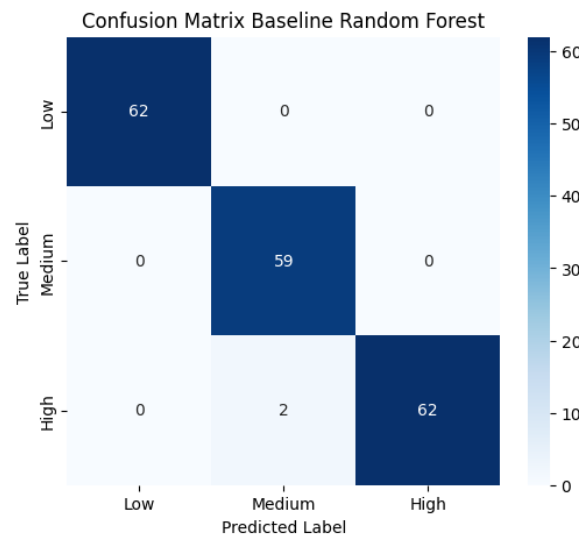
3.3. Kinerja Eksperimen Baseline Model

Pada tahapan ini, eksperimen pengujian performa klasifikasi diawali dengan mengeksekusi model acuan awal (baseline model) menggunakan algoritma Random Forest Classifier. Proses pelatihan dan pengujian ini berjalan dengan memanfaatkan seluruh konfigurasi parameter bawaan (default) dari pustaka Scikit-Learn tanpa melibatkan tindakan optimasi parameter apa pun. Model acuan awal ini diuji langsung menggunakan subset data pengujian (testing data) yang terisolasi berjumlah 185 baris sampel. Berdasarkan hasil evaluasi komputasi, diperoleh nilai akurasi global (accuracy) yang sangat tinggi yaitu sebesar 0,9892 (98,92%). Nilai tersebut menunjukkan kemampuan fondasi awal model yang sudah sangat kuat dalam mengidentifikasi pola konsumsi energi pintar bahkan sebelum diberikan tindakan intervensi hyperparameter. Rincian kinerja klasifikasi secara sektoral yang memuat metrik precision, recall, dan f1-score untuk masing-masing kategori kelas target (Low, Medium, High) disajikan secara terstruktur pada Tabel 5.

Tabel 5. Classification Report Hasil Pengujian Baseline Model

Kategori Kelas Target	Precision	Recall	F1-Score	Support (Data Uji)
Low (0)	1	1	1	62
Medium (1)	0,9672	1	0,9833	59
High (2)	1	0,9688	0,9841	64
Akurasi Global (Accuracy)			0,9892	185
Macro Average	0,9891	0,9896	0,9892	185
Weighted Average	0,9895	0,9892	0,9892	185

Guna menelaah lebih lanjut mengenai efektivitas prediksi dan mendeteksi kecenderungan kesalahan klasifikasi pada model acuan, visualisasi Confusion Matrix dibentuk sebagaimana disajikan pada Gambar 7. Melalui pemetaan matriks kontingensi tersebut, dapat diidentifikasi secara spesifik sebaran koordinat antara kelas aktual dengan kelas hasil prediksi yang dikeluarkan oleh algoritma. Berdasarkan representasi visual pada Gambar 8, diketahui bahwa baseline model ini menunjukkan performa yang luar biasa dengan hanya melakukan kesalahan prediksi pada kategori kelas High. Tercatat sejumlah 2 sampel data yang secara aktual termasuk dalam kategori kelas High, namun salah diprediksi oleh model ke dalam kategori kelas Medium. Di sisi lain, seluruh sampel pada kelas Low (62 sampel) dan kelas Medium (59 sampel) berhasil diklasifikasikan secara sempurna dengan tingkat kesalahan nol. Fenomena salah klasifikasi yang sangat minim ini membuktikan secara empiris bahwa karakteristik fitur prediktor pada dataset ini sudah memiliki pemisahan yang cukup tegas, meskipun terdapat sedikit tumpang tindih (overlapping) minor pada batas keputusan (decision boundary) antara kelas High dan kelas Medium.



Gambar 7. Visualisasi Confusion Matrix pada Tahap Baseline Model

Secara keseluruhan, hasil pengujian pada model acuan awal ini mengonfirmasi bahwa algoritma Random Forest memiliki keandalan yang sangat tinggi dalam melakukan klasifikasi konsumsi energi pintar, meskipun tanpa adanya intervensi parameter eksternal. Namun demikian, guna memastikan ketercapaian performa yang maksimal serta menyelidiki apakah kesalahan klasifikasi minor pada kelas High dapat dieliminasi secara total, tahapan eksperimen harus dilanjutkan ke proses penyietelan hyperparameter. Evaluasi mendalam mengenai perubahan arsitektur keputusan, efisiensi komputasi, serta perolehan metrik performa model setelah penerapan pencarian parameter optimal akan dijabarkan secara sistematis pada sub-bab 3.4 Hasil Optimasi Hyperparameter menggunakan GridSearchCV.

3.4. Hasil Optimasi Hyperparameter menggunakan GridSearchCV

Tahapan eksperimen dilanjutkan dengan mengimplementasikan proses penyetalan hyperparameter memanfaatkan metode GridSearchCV guna meningkatkan performa akurasi dan meminimalkan kesalahan prediksi model. Melalui pengujian menyeluruh (exhaustive search) dengan skema 5-fold cross-validation, algoritma berhasil menemukan kombinasi parameter paling optimal yang menghasilkan skor estimasi terbaik pada data latih. Berdasarkan hasil pencarian eksperimental tersebut, kombinasi parameter terbaik (best parameters) yang diperoleh meliputi nilai `n_estimators` sebesar 200 untuk jumlah pohon keputusan, nilai `max_depth` sebesar `None` yang berarti kedalaman pohon tidak dibatasi, nilai `min_samples_split` sebesar 2 untuk jumlah minimum sampel pada percabangan simpul, serta nilai `min_samples_leaf` sebesar 1 untuk jumlah minimum sampel pada simpul daun.

Melalui konfigurasi kombinasi parameter optimal tersebut, diperoleh nilai akurasi rata-rata terbaik pada fase validasi silang (best cross-validation accuracy) sebesar 0,9973 (99,73%) pada data pelatihan. Selanjutnya, ketika arsitektur model akhir (Best Random Forest) ini dieksekusi menggunakan subset data pengujian (testing data) yang terisolasi, performa klasifikasi global yang dihasilkan terbukti tetap konsisten dengan pencapaian nilai akurasi (accuracy) sebesar 0,9946 (99,46%). Kenaikan metrik ini mengonfirmasi bahwa penyesuaian nilai parameter internal model melalui GridSearchCV memberikan dampak positif yang signifikan dalam meningkatkan akurasi prediksi sistem. Rincian hasil evaluasi performa classification report setelah proses penyetalan hyperparameter disajikan secara detail pada Tabel 6.

Tabel 6. Classification Report Hasil Pengujian Optimized Model (Best Random Forest)

Kategori Kelas Target	Precision	Recall	F1-Score	Support (Data Uji)
<i>Low</i> (0)	1	1	1	62
<i>Medium</i> (1)	0,9833	1	0,9916	59
<i>High</i> (2)	1	0,9844	0,9921	64
Akurasi Global (Accuracy)			0,9946	185
<i>Macro Average</i>	0,9944	0,9948	0,9946	185
<i>Weighted Average</i>	0,9947	0,9946	0,9946	185

Untuk melakukan analisis mendalam mengenai dampak optimasi hyperparameter terhadap pergeseran batas keputusan model, visualisasi Confusion Matrix akhir dibentuk sebagaimana disajikan pada Gambar 9. Berdasarkan pemetaan koordinat matriks kontingensi tersebut, terlihat dengan jelas adanya perbaikan performa prediksi sektoral yang signifikan dibandingkan dengan tahap baseline model. Melalui intervensi GridSearchCV, tingkat kesalahan klasifikasi (misclassification) berhasil ditekan hingga setengahnya, di mana hanya tersisa 1 sampel data yang salah diprediksi oleh sistem. Kesalahan tunggal tersebut terjadi pada data yang secara aktual termasuk dalam kategori kelas *High*, namun salah diklasifikasikan ke dalam kategori kelas *Medium*. Sementara itu, model mampu mempertahankan kesempurnaan klasifikasi 100% pada kategori kelas *Low* (62 sampel) dan kelas *Medium* (59 sampel). Keberhasilan mereduksi jumlah data eror ini membuktikan secara empiris bahwa penentuan kombinasi hyperparameter yang optimal berhasil mempertegas batas pemisahan antar-kelas keputusan dan meningkatkan kemampuan generalisasi model secara substansial.

Untuk melakukan analisis mendalam mengenai dampak optimasi hyperparameter terhadap pergeseran batas keputusan model, visualisasi Confusion Matrix akhir dibentuk sebagaimana disajikan pada Gambar 9. Berdasarkan pemetaan koordinat matriks kontingensi tersebut, terlihat dengan jelas adanya perbaikan performa prediksi sektoral yang signifikan dibandingkan dengan tahap baseline model. Melalui intervensi GridSearchCV, tingkat kesalahan klasifikasi (misclassification) berhasil ditekan hingga setengahnya, di mana hanya tersisa 1 sampel data yang salah diprediksi oleh sistem. Kesalahan tunggal tersebut terjadi pada data yang secara aktual termasuk dalam kategori kelas *High*, namun salah diklasifikasikan ke dalam kategori kelas *Medium*. Sementara itu, model mampu mempertahankan kesempurnaan klasifikasi 100% pada kategori kelas *Low* (62 sampel) dan kelas *Medium* (59 sampel). Keberhasilan mereduksi jumlah data eror ini membuktikan secara empiris bahwa penentuan kombinasi hyperparameter yang optimal berhasil mempertegas batas pemisahan antar-kelas keputusan dan meningkatkan kemampuan generalisasi model secara substansial.

3.5. Analisis Komparatif dan Evaluasi Kinerja Model

Setelah seluruh tahapan eksperimen selesai dieksekusi, analisis komparatif dilakukan secara mendalam untuk mengevaluasi perubahan performa model klasifikasi sebelum dan sesudah diterapkannya optimasi hyperparameter. Berdasarkan hasil pengujian pada subset data uji, implementasi metode GridSearchCV terbukti memberikan dampak positif yang konsisten terhadap seluruh metrik evaluasi. Nilai akurasi global mengalami peningkatan dari 0,9892 (98,92%) pada tahap baseline model menjadi 0,9946 (99,46%) pada tahap optimized model. Kenaikan sebesar 0,54% ini diikuti pula dengan peningkatan nilai presisi makro dari 0,9891 menjadi 0,9944, serta nilai recall makro dari 0,9896 menjadi 0,9948. Secara kuantitatif, visualisasi confusion matrix merekam penyusutan jumlah sampel eror, di mana model hasil optimasi hanya menyisakan 1 sampel salah prediksi dibandingkan model acuan awal yang menyisakan 2 sampel eror. Rincian perbandingan performa head-to-head antara kedua skema eksperimen tersebut disajikan secara komparatif pada Tabel 7.

Tabel 7. Komparasi Performa Klasifikasi Sebelum dan Sesudah Optimasi

Metrik Evaluasi (Skala Makro)	Tahap Baseline Model	Tahap Optimized Model	Persentase Kenaikan / Perubahan
Akurasi (<i>Accuracy</i>)	0,9892	0,9946	+ 0,0054 (+ 0,54%)
Presisi (<i>Precision</i>)	0,9891	0,9944	+ 0,0053 (+ 0,53%)
Recall (<i>Sensitivitas</i>)	0,9896	0,9948	+ 0,0052 (+ 0,52%)
<i>F1-Score</i>	0,9892	0,9946	+ 0,0054 (+ 0,54%)
Kuantitas Error Klasifikasi	2 Sampel Data	1 Sampel Data	Reduksi 1 Sampel Data

Prosedur evaluasi dilanjutkan dengan menguji stabilitas generalisasi model melalui analisis nilai selisih akurasi (gap accuracy) untuk mendeteksi risiko terjadinya fenomena overfitting. Pengujian ini dilakukan dengan membandingkan margin antara tingkat akurasi pada fase pelatihan (training accuracy) dengan fase pengujian (testing accuracy). Pada tahap baseline model, nilai akurasi pelatihan mencapai 1,0000 dengan akurasi pengujian sebesar 0,9892, sehingga menghasilkan nilai gap accuracy sebesar 0,0108. Setelah dilakukan penyetelan hyperparameter, model akhir mampu mempertahankan nilai akurasi pelatihan sebesar 1,0000 namun berhasil meningkatkan akurasi pengujian menjadi 0,9946, sehingga nilai gap accuracy menyusut menjadi hanya sebesar 0,0054. Penurunan nilai delta margin yang semakin mendekati nilai nol ini membuktikan secara empiris bahwa model akhir Best Random Forest tidak mengalami indikasi overfitting. Struktur pohon keputusan yang dikembangkan terbukti sangat robust dan memiliki kapabilitas generalisasi yang luar biasa dalam memprediksi data baru yang belum pernah dipelajari sebelumnya.

```

*** Menghitung skor 5-Fold CV... (tunggu sebentar)
Skor 5-Fold Baseline : [1.      1.      1.      0.99319728 0.98639456]
Skor 5-Fold Optimized: [1.      1.      1.      0.99319728 0.99319728]

===== HASIL PAIRED T-TEST =====
T-Statistic : -1.0000
P-Value      : 0.3739

===== WAKTU EKSEKUSI =====
Baseline      : 609.27 ms
Optimized     : 1623.39 ms

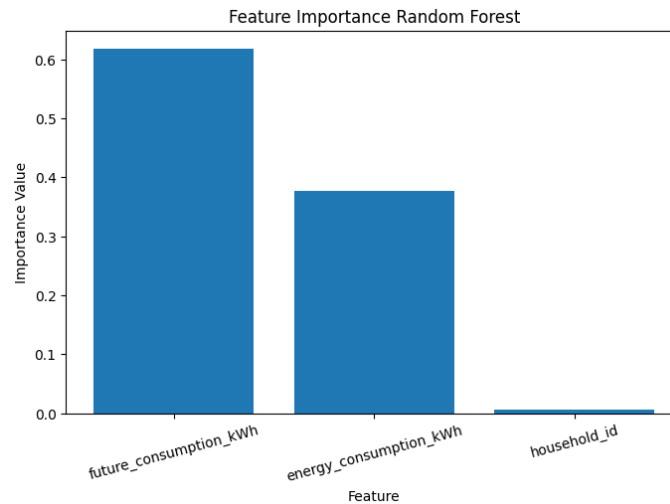
```

Gambar 8. Hasil Evaluasi Model Menggunakan 5-Fold Cross Validation dan Paired T-Test

Untuk menguji validitas peningkatan performa klasifikasi, analisis signifikansi statistik dieksekusi menggunakan uji Paired T-Test berbasis evaluasi 5-Fold Cross-Validation. Hasil pengujian hipotesis pada distribusi skor akurasi mengonfirmasi nilai p-value sebesar 0,3739 ($p > 0,05$). Secara statistik, nilai ini mengindikasikan bahwa margin peningkatan akurasi tidak berbeda secara radikal. Namun, dari sudut pandang pemodelan Machine Learning, temuan ini mengungkap fenomena Ceiling Effect (efek plafon). Berdasarkan observasi pada hasil 5-fold CV, model baseline pada dasarnya telah mencapai titik saturasi dengan akurasi sempurna (1,00) pada tiga lipatan data pertama. Dalam kondisi di mana arsitektur bawaan telah sangat superior, implementasi GridSearchCV tidak lagi ditujukan untuk mendongkrak metrik akurasi secara agresif, melainkan berfungsi sebagai instrumen penjamin stabilitas (robustness). Hal ini terbukti pada iterasi validasi lipatan kelima (worst-case scenario), di mana model hasil optimasi berhasil mempertahankan akurasi di angka 0,9931, mencegah penurunan performa menjadi 0,9863 seperti yang dialami oleh model baseline. Dengan kata lain, optimasi ini sukses mereduksi varians prediksi, menjadikan model jauh lebih stabil dan tahan terhadap fluktuasi data uji. Keunggulan stabilitas ini didukung oleh analisis efisiensi waktu komputasi (computational runtime). Proses pelatihan Baseline Model membutuhkan waktu eksekusi sebesar 609,27 ms, sedangkan Optimized Model membutuhkan waktu 1623,39 ms. Meskipun proses pelatihan model yang dioptimasi membutuhkan tambahan waktu presisi, total durasi eksekusinya tetap berada pada skala 1,6 detik. Kecepatan komputasi yang sangat singkat ini memberikan justifikasi empiris bahwa arsitektur algoritma yang diusulkan tidak hanya stabil di berbagai kondisi data, tetapi juga sangat ideal dan ringan untuk tahap deployment pada lingkungan perangkat edge-IoT yang memiliki keterbatasan daya komputasi.

3.6. Analisis Kontribusi Variabel (Feature Importance)

Evaluasi komprehensif pada bab ini ditutup dengan melakukan analisis kontribusi variabel (feature importance) guna memberikan interpretasi yang transparan mengenai mekanisme pengambilan keputusan di dalam model akhir Best Random Forest. Analisis ini dieksekusi dengan mengekstraksi nilai atribut `feature_importances_` bawaan dari algoritma, yang menghitung seberapa besar rata-rata penurunan tingkat keanekaragaman data (mean decrease in impurity) yang disumbangkan oleh masing-masing fitur prediktor di seluruh pohon keputusan. Nilai kontribusi kuantitatif dari ketiga variabel prediktor yang diuji, yaitu `future_consumption_kWh`, `energy_consumption_kWh`, dan `household_id`, disajikan secara visual melalui grafik batang pada Gambar 9.



Gambar 9. Grafik Tingkat Kepentingan Fitur (Feature Importance) Model Akhir

Berdasarkan visualisasi pada Gambar 9, fitur `future_consumption_kWh` mendominasi model dengan tingkat kontribusi sebesar 61,75%, diikuti oleh `energy_consumption_kWh` sebesar 37,64%. Dari perspektif pemodelan prediktif konvensional, tingginya dominasi fitur yang merepresentasikan metrik masa depan sering kali dikhawatirkan sebagai gejala kebocoran data (data leakage). Namun demikian, dari sudut pandang rancang bangun ekosistem Internet of Things (IoT) pada smart home, dominasi fitur ini sama sekali tidak mengindikasikan kecacatan metodologi, melainkan merupakan representasi dari desain sistem pemantauan daya yang proaktif. Dalam arsitektur yang diusulkan, model Random Forest diposisikan sebagai mesin pengambil keputusan hierarki sekunder (secondary decision-maker). Nilai pada `future_consumption_kWh` tidak diasumsikan sebagai kebocoran masa depan, melainkan sebagai variabel eksogen sah yang disuplai oleh modul peramalan deret waktu (time-series forecasting module) yang beroperasi secara independen di hulu sistem. Dengan mengintegrasikan proyeksi masa depan tersebut sebagai input untuk mengklasifikasikan status tingkat ancaman konsumsi saat ini (Low, Medium, High), sistem IoT dapat mengeksekusi aktuasi pencegahan—seperti memutuskan atau mereduksi beban daya perangkat tertentu—sebelum lonjakan konsumsi riil benar-benar terjadi. Oleh karena itu, ketergantungan model yang tinggi terhadap fitur ini secara metodologis sangat valid dan selaras dengan paradigma efisiensi energi preventif pada smart home modern.

4. KESIMPULAN

Penelitian ini telah berhasil memformulasikan dan memvalidasi arsitektur model klasifikasi prediktif berbasis Random Forest yang dioptimasi untuk manajemen energi cerdas pada ekosistem smart home. Dari perspektif teoretis Machine Learning, implementasi ini membuktikan bahwa pada dataset IoT berskala menengah yang rentan mengalami efek plafon (ceiling effect), peranan optimasi GridSearchCV secara fundamental bergeser dari sekadar instrumen pendongkrak akurasi menjadi mekanisme krusial untuk penjaminan stabilitas model (robustness) dan reduksi varians prediksi. Keputusan desain arsitektur yang memosisikan Random Forest sebagai mesin pengambil keputusan sekunder yang memanfaatkan fitur proyeksi masa depan (future consumption) terbukti secara konseptual mewujudkan paradigma sistem aktuasi daya proaktif tanpa mencederai prinsip integritas data (data leakage). Lebih jauh, arsitektur model yang diajukan menunjukkan titik ekuilibrium yang optimal antara ketepatan prediktif dan efisiensi waktu komputasi, memberikan justifikasi ilmiah yang kuat bahwa kompleksitas model ensemble dapat ditransformasikan menjadi bentuk yang sangat ringan untuk dieksekusi secara andal pada lingkungan perangkat edge-IoT yang memiliki keterbatasan sumber daya komputasi.

Berdasarkan capaian teoretis tersebut, penelitian ini merekomendasikan dua arah eksplorasi strategis untuk pengembangan lanjutan. Pertama, dari aspek rekayasa perangkat lunak, diperlukan pengujian arsitektur hibrida yang secara langsung mengintegrasikan (pipeline integration) model klasifikasi ini dengan modul algoritma peramalan deret waktu ringan di dalam satu lingkungan edge-computing guna mengevaluasi latensi inferensi end-to-end. Kedua, dari aspek ketahanan sistem, arsitektur model teroptimasi ini perlu diuji menggunakan data konsumsi energi multivariat yang mengakomodasi anomali beban kejut (shock load) akibat fluktuasi cuaca ekstrem, guna memvalidasi batas toleransi kegagalan (fault tolerance) dari struktur pohon keputusan sebelum sistem benar-benar dikomersialkan secara massal.

REFERENSI

- [1] W. Agustiarini *et al.*, “Real-Time IoT-based Energy Power Consumption Monitoring System for Smart Homes,” *Int. J. Adv. Sci. Eng. Inf. Technol.*, vol. 15, no. 5, pp. 1407–1412, 2025, doi: 10.18517/ijaseit.15.5.20918.
- [2] T. Wang, Q. Zhao, W. Gao, and X. He, “Research on energy consumption in household sector: a comprehensive review based on bibliometric analysis,” *Front. Energy Res.*, vol. 11, no. January, pp. 1–22, 2023, doi: 10.3389/fenrg.2023.1209290.

- [3] Regina Citra Kurnia Pangestu and Anak Agung Ketut Ayuningsasi, "Pengaruh Konsumsi Energi Sektor Industri, Rumah Tangga, dan Transportasi terhadap Emisi Karbon di Indonesia," *Inisiatif: Jurnal Ekonomi, Akuntansi dan Manajemen*, vol. 3, no. 4, pp. 297–311, 2024, doi: 10.30640/inisiatif.v3i4.3154.
- [4] D. B. Noya, O. Wullur, O. Teng, Z. Sukarame, F. Aguw, and M. Rantung, "The Influence of Using Electronic Devices in the Household," *Jurnal Syntax Admiration*, vol. 6, no. 2, pp. 1359–1366, 2025, doi: 10.46799/jsa.v6i2.2138.
- [5] A. R. Rudiyanto, B. P. Satria, and H. D. Panjaitan, "Optimization of Smart Home Energy Consumption Using Machine Learning-Based Load Forecasting," *Journal of Technology Informatics and Engineering*, vol. 4, no. 2, pp. 300–316, 2025, doi: 10.51903/jtie.v4i2.437.
- [6] M. Casanova and J. Moloney, "Electricity 2025: Analysis and forecast to 2027," 2025. [Online]. Available: www.iea.org
- [7] P. Y. Muslim *et al.*, "Tangga Berbasis Iot Iot-Based Household Electricity Consumption Monitoring System," *Jurnal Ilmiah Teknologi Informasi dan Komunikasi*, vol. 18, no. 2, pp. 1–7, 2025, [Online]. Available: <https://www.iea.org/reports/electricity-2025>
- [8] A. P. Wijaya, M. S. Said, and A. M. Islah, "Pengembangan Sistem Pemantauan Pemakaian Listrik Rumah Berbasis Internet of Things," *Simtek : jurnal sistem informasi dan teknik komputer*, vol. 10, no. 2, pp. 475–479, 2025, doi: 10.51876/simtek.v10i2.1676.
- [9] W. Jia and S. Wu, "Spatial Differences and Influencing Factors of Energy Poverty: Evidence From Provinces in China," *Front. Environ. Sci.*, vol. 10, no. June, pp. 1–16, 2022, doi: 10.3389/fenvs.2022.921374.
- [10] S. Su *et al.*, "Temporal dynamic assessment of household energy consumption and carbon emissions in China: From the perspective of occupants," *Sustain. Prod. Consum.*, vol. 37, pp. 142–155, 2023, doi: <https://doi.org/10.1016/j.spc.2023.02.014>.
- [11] J. Zheng, Y. Dang, and U. Assad, "Household energy consumption, energy efficiency, and household income—Evidence from China," *Appl. Energy*, vol. 353, p. 122074, 2024, doi: <https://doi.org/10.1016/j.apenergy.2023.122074>.
- [12] I. Valentine, J. Triloka, and R. Sholehurrohman, "Application of Support Vector Machine Algorithm for Energy Consumption Prediction," *Journal of Applied Informatics and Computing*, vol. 10, no. 2, pp. 1598–1605, 2026, doi: 10.30871/jaic.v10i2.11713.
- [13] U. Maltuf and Z. Fatah, "Penerapan Algoritma Decision Tree Untuk Klasifikasi Konsumsi Energi Listrik Rumah Tangga Dengan Penggunaan Rapid Miner," *Jurnal Ilmiah Multidisiplin Ilmu*, vol. 1, no. 1, pp. 38–45, 2025, doi: <https://doi.org/10.69714/Ohmk8712>.
- [14] A. Olivia and A. Jaenul, "Smart Home Electricity Meter Based on IoT with Bill Prediction Using Random Forest Algorithm," *Journal of Applied Information and Communication Technology*, vol. 12, no. 1, pp. 78–81, 2026, doi: <https://doi.org/10.32497/jaict.v12i1>.
- [15] N. Bhardwaj and P. Joshi, "Adaptive Energy Management for MATTER-Enabled Smart Homes," *IEEE Access*, vol. 13, no. June, pp. 113773–113786, 2025, doi: 10.1109/ACCESS.2025.3584381.
- [16] M. K. Saravana, M. S. Roopa, J. S. Arunalatha, and K. R. Venugopal, "Transformers for Multivariate Time Series Forecasting: Comprehensive Analysis, Challenges, Research Opportunities and Future Prospects," *IEEE Access*, 2026, doi: 10.1109/ACCESS.2026.3654408.
- [17] A. A. Alsumaiei, "Complexity-efficiency dynamics of metaheuristic-optimized recurrent neural network models for drought forecasting in hyper-arid Kuwait," *J. Hydrol. Reg. Stud.*, vol. 64, Apr. 2026, doi: 10.1016/j.ejrh.2026.103300.
- [18] P. W. Adi, A. Sugiharto, S. Adhy, E. Vianita, and G. Rahman, "Improving balance in intrusion detection for IoT networks using multicollinearity and discriminative power-based feature selection," *Discover Internet of Things*, Jun. 2026, doi: 10.1007/s43926-026-00402-x.
- [19] D. Romaissa Beddiar *et al.*, "Efficient Edge AI Inference: A Literature Review," 2026. doi: <http://dx.doi.org/10.2139/ssrn.6359018>.
- [20] A. N. S. Kinasih, A. N. Handayani, J. T. Ardiansah, and N. S. Damanhuri, "Comparative analysis of decision tree and random forest classifiers for structured data classification in machine learning," *Science in Information Technology Letters*, vol. 5, no. 2, pp. 13–24, 2024, doi: 10.31763/sitech.v5i2.1746.
- [21] D. Ljunggren and S. Ishii, "A Comparative Analysis of the Robustness to Noise of Machine Learning Classifiers," *Communications in Computer and Information Science*, vol. 2413 CCIS, pp. 182–192, 2025, doi: 10.1007/978-3-031-87511-3_13.
- [22] M. Imani, A. Beikmohammadi, and H. R. Arabnia, "Comprehensive Analysis of Random Forest and XGBoost Performance with SMOTE, ADASYN, and GNUS Under Varying Imbalance Levels," *Technologies (Basel)*, vol. 13, no. 3, pp. 1–40, 2025, doi: 10.3390/technologies13030088.
- [23] H. Putra and R. Rumini, "Comparative Study of Logistic Regression, Random Forest, and XGBoost for Bank Loan Approval Classification," *Journal of Applied Informatics and Computing*, vol. 9, no. 5, pp. 2822–2835, 2025, doi: 10.30871/jaic.v9i5.10862.
- [24] M. W. Nugroho, "Analisis Performa Algoritma Random Forest dalam Mengatasi Overfitting pada Model Prediksi," *Jurnal JTIK (Jurnal Teknologi Informasi dan Komunikasi)*, vol. 9, no. 4, pp. 1562–1571, 2025, doi: 10.35870/jtik.v9i4.4236.
- [25] Y. Zhao, W. Zhang, and X. Liu, "Grid search with a weighted error function: Hyper-parameter optimization for financial time series forecasting," *Appl. Soft Comput.*, vol. 154, no. August 2023, p. 111362, 2024, doi: 10.1016/j.asoc.2024.111362.
- [26] D. M. Belete and M. D. Huchaiah, "Grid search in hyperparameter optimization of machine learning models for prediction of HIV/AIDS test results," *International Journal of Computers and Applications*, vol. 44, no. 9, pp. 875–886, 2022, doi: 10.1080/1206212X.2021.1974663.
- [27] C. Azzaria, E. Daniati, and A. Ristyawan, "Peningkatan Akurasi Deteksi Liver Disease melalui Hyperparameter Tuning pada Algoritma Random Forest," *The Indonesian Journal of Computer Science Research*, vol. 4, no. 2, pp. 139–147, 2025, [Online]. Available: <https://subset.id/index.php/IJCSR>
- [28] G. B. Anugrah and T. Arifin, "Optimasi Support Vector Machine Menggunakan Seleksi Fitur Random Forest Dan Hyperparameter Gridsearchcv Untuk Klasifikasi Raisin Dataset," *Djtechno: Jurnal Teknologi Informasi*, vol. 6, no. 2, pp. 527–540, 2025, doi: 10.46576/djtechno.v6i2.7008.
- [29] W. Nugraha and A. Sasongko, "Hyperparameter Tuning on Classification Algorithm with Grid Search," *Sistemasi*, vol. 11, no. 2, p. 391, 2022, doi: 10.32520/stmsi.v11i2.1750.
- [30] Ziya, "IoT-SmartHome Energy Dataset," Kaggle. [Online]. Available: [IoT-SmartHome Energy Dataset](https://www.kaggle.com/datasets/ziya0101/iot-smarthome-energy-dataset)
- [31] L. Ma, M. S. Ghouri, M. Altaf, M. Zheng, and L. Wang, "Explainable machine learning for predicting building energy performance and sustainability certification in construction megaprojects," *Energy Build.*, vol. 367, p. 117774, Sep. 2026, doi: 10.1016/J.ENBUILD.2026.117774.