

SISTEM PENDUKUNG KEPUTUSAN EVALUASI RISIKO PENIPUAN TRANSAKSI KEUANGAN MENGGUNAKAN RANDOM FOREST

Rolis Liu¹, Bergas Cahyo Nuswantoro², Mohammad Ichsan Nurdin³, Purwantoro⁴

Program Studi Informatika^{1,2,3,4}, Universitas Singaperbangsa Karawang^{1,2,3,4}

2310631170117@student.unsika.ac.id¹, 2310631170009@student.unsika.ac.id²,

2310631170097@student.unsika.ac.id³, Purwantoro@staff.unsika.ac.id⁴

* Corresponding Author: 2310631170009@student.unsika.ac.id

Abstrak

Penipuan transaksi keuangan (financial fraud) merupakan ancaman serius dalam ekosistem keuangan digital yang terus berkembang. Penelitian ini mengembangkan Sistem Pendukung Keputusan (SPK) untuk mengevaluasi risiko penipuan transaksi keuangan menggunakan algoritma Random Forest. Dataset yang digunakan adalah Credit Card Fraud Detection dari Kaggle yang memuat 284.807 transaksi dengan proporsi penipuan sebesar 0,17%. Proses penelitian mencakup preprocessing data (normalisasi, SMOTE oversampling), pembangunan model Random Forest dengan tuning hyperparameter, dan evaluasi menggunakan metrik akurasi, presisi, recall, dan F1-Score. Hasil eksperimen menunjukkan bahwa model Random Forest mencapai akurasi 99,94%, presisi 97,82%, recall 86,73%, dan F1-Score 91,95% pada data uji. Fitur paling berpengaruh dalam deteksi penipuan adalah V14, V10, V4, dan V12. SPK yang dibangun mampu mengklasifikasikan transaksi berisiko tinggi secara real-time dan memberikan rekomendasi tindakan bagi pengelola keuangan.

Kata kunci: deteksi penipuan; klasifikasi; keuangan; Random Forest; sistem pendukung keputusan

Abstract

Financial transaction fraud is a serious threat in the growing digital financial ecosystem. This study develops a Decision Support System (DSS) for evaluating financial transaction fraud risk using the Random Forest algorithm. The dataset used is the Credit Card Fraud Detection dataset from Kaggle, containing 284,807 transactions with a fraud proportion of 0.17%. The research process includes data preprocessing (normalization, SMOTE oversampling), Random Forest model building with hyperparameter tuning, and evaluation using accuracy, precision, recall, and F1-Score metrics. Experimental results show that the Random Forest model achieves an accuracy of 99.94%, precision of 97.82%, recall of 86.73%, and F1-Score of 91.95% on test data. The most influential features in fraud detection are V14, V10, V4, and V12. The DSS built is capable of classifying high-risk transactions in real-time and providing action recommendations for financial managers.

Keywords: classification; decision support system; finance; fraud detection; Random Forest

1. Pendahuluan

Perkembangan teknologi digital telah mengubah lanskap keuangan secara signifikan. Layanan perbankan digital, dompet elektronik, dan platform pembayaran online kini menjadi bagian tak terpisahkan dari aktivitas ekonomi masyarakat. Namun, di balik kemudahan tersebut, muncul ancaman serius berupa penipuan transaksi keuangan (financial fraud) yang merugikan individu maupun institusi keuangan. Menurut laporan Association of Certified

Fraud Examiners (ACFE), organisasi secara global kehilangan rata-rata 5% dari pendapatan tahunan mereka akibat penipuan, dengan kerugian kumulatif mencapai miliaran dolar setiap tahunnya [1].

Deteksi penipuan secara manual memiliki keterbatasan dalam hal kecepatan, konsistensi, dan skalabilitas. Volume transaksi yang sangat besar membuat pendekatan konvensional tidak efisien dan rentan terhadap human error. Oleh karena itu, pemanfaatan teknik machine learning untuk otomatisasi deteksi penipuan telah mendapat perhatian besar dari komunitas akademik dan industri keuangan [2][3]. Berbagai algoritma telah dieksplorasi, mulai dari Logistic Regression, Decision Tree, Support Vector Machine (SVM), hingga metode ensemble seperti Random Forest dan Gradient Boosting.

Random Forest merupakan algoritma ensemble learning yang membangun sejumlah pohon keputusan secara paralel dan menggabungkan hasilnya untuk menghasilkan prediksi yang lebih akurat dan robust [4]. Algoritma ini terbukti efektif dalam menangani dataset tidak seimbang (imbalanced dataset), yang merupakan karakteristik umum data penipuan keuangan. Selain itu, Random Forest mampu mengidentifikasi fitur-fitur penting yang berkontribusi terhadap keputusan klasifikasi, sehingga mendukung interpretabilitas model.

Penelitian ini bertujuan untuk: (1) membangun Sistem Pendukung Keputusan (SPK) berbasis Random Forest yang mampu mengklasifikasikan transaksi keuangan sebagai penipuan atau bukan; (2) mengevaluasi performa model menggunakan metrik yang relevan untuk dataset tidak seimbang; dan (3) mengidentifikasi fitur transaksi yang paling berpengaruh dalam deteksi penipuan. Penelitian ini diharapkan dapat memberikan kontribusi berupa prototipe SPK yang dapat diintegrasikan ke dalam sistem keamanan keuangan perbankan atau platform fintech.

2. Kajian Pustaka dan pengembangan hipotesis

Kajian pustaka membahas tentang teori dan penelitian terdahulu yang berkaitan dengan topik penelitian yang menjadi landasan logis dalam mengembangkan hipotesis penelitian termasuk kerangka konsep penelitian.

2.1. Sistem Pendukung Keputusan

Sistem Pendukung Keputusan (SPK) adalah sistem informasi yang mendukung aktivitas pengambilan keputusan dalam bisnis atau organisasi, melayani level manajemen menengah hingga atas, dan membantu dalam menangani masalah yang berubah cepat serta tidak mudah dispesifikasikan secara awal, yakni masalah keputusan yang tidak terstruktur maupun semi-terstruktur[5].

Menurut Sprague dan Carlson, kerangka dasar SPK memuat basis data dan basis model serta sistem manajemennya, serta komponen dialog tempat pengambil keputusan berinteraksi. Dalam konteks deteksi penipuan keuangan, SPK yang berbasis model (model-driven DSS) memanfaatkan algoritma klasifikasi sebagai inti analitik untuk menghasilkan rekomendasi tindakan secara otomatis. SPK lebih banyak disesuaikan bagi individu atau organisasi yang membuat keputusan dibandingkan sistem konvensional, sehingga menjadikannya pendekatan yang relevan untuk skenario deteksi risiko secara real-time[6].

2.2. Penipuan Transaksi Keuangan

Penipuan keuangan (*financial fraud*) merupakan ancaman sistemik yang terus berkembang seiring meningkatnya adopsi layanan keuangan digital. Laporan Ancaman Pembayaran dan Tren Penipuan 2023 menyoroti meningkatnya kompleksitas serangan rekayasa

sosial dan *phishing*, persistensi ancaman *malware* termasuk *advanced persistent threats* (APT), *botnet*, serta serangan *distributed denial-of-service* (DDoS) yang mencerminkan kerentanan sektor keuangan. Lebih jauh, proyeksi menunjukkan bahwa pada 2026, penipuan kartu kredit secara global akan tumbuh hingga \$43 miliar, dengan kerugian di Amerika Serikat melebihi \$12,5 miliar pada 2025.

Karakteristik utama data penipuan keuangan adalah distribusi kelas yang sangat tidak seimbang (*highly imbalanced*). Deteksi penipuan kartu kredit tetap menjadi tantangan signifikan bagi teknologi *business intelligence* karena sebagian besar dataset transaksi kartu kredit sangat tidak seimbang.

2.3. Algoritma Random Forest

2.3.1. Konsep Dasar dan Mekanisme Ensemble

Random Forest (RF) diperkenalkan oleh Breiman[6] yang merupakan algoritma statistik atau *machine learning* untuk prediksi yang merata-ratakan prediksi atas banyak pohon keputusan individual[7].

Prediksi akhir RF untuk klasifikasi diperoleh melalui mekanisme *majority voting* dari seluruh pohon. Jika terdapat T pohon keputusan dalam forest dan $\hat{y}_t(X)$ adalah prediksi pohon ke- t untuk observasi X , maka prediksi RF dinyatakan sebagai:

$$\hat{y}_{RF}(X) = \underset{c}{\operatorname{argmax}} \sum_{t=1}^T I[\hat{y}_t(X) = c] \quad (1)$$

di mana c adalah label kelas dan $I[\cdot]$ adalah fungsi indikator.

Dengan menambahkan lebih banyak pohon, varians berkurang secara monoton tanpa meningkatkan *bias*, satu-satunya kerugian adalah peningkatan komputasi dan memori. Breiman membuktikan sifat konvergensi ini dalam makalah asli 2001.

2.3.2. Kriteria Pemisahan Node

Pada setiap node pohon keputusan dalam RF, pemisahan dilakukan berdasarkan kriteria impuritas. Kriteria yang umum digunakan adalah *Gini impurity*, yang didefinisikan sebagai:

$$I_G(p) = 1 - \sum_{i=1}^C p_i^2 \quad (2)$$

di mana C adalah jumlah kelas dan p_i adalah proporsi sampel kelas i pada node tersebut. Pemisahan terbaik dipilih dengan memaksimalkan *Gini gain*:

$$\Delta I_G = I_G(S) - \left(\frac{|S_L|}{|S|} I_G(S_L) + \frac{|S_R|}{|S|} I_G(S_R) \right) \quad (3)$$

di mana S adalah himpunan sampel sebelum pemisahan, S_L dan S_R adalah himpunan sampel di cabang kiri dan kanan. Kriteria alternatif berbasis *information gain* menggunakan entropi Shannon $H(S)$ dan didefinisikan sebagai selisih entropi antara node induk dan node anak yang dibobot secara proporsional terhadap jumlah sampelnya.

2.4. Peneliti Terdahulu

Sejumlah penelitian telah mengkaji penggunaan RF dan metode *ensemble* untuk deteksi penipuan keuangan dengan berbagai pendekatan. Studi komparatif pada sistem perbankan menemukan bahwa model RF terbukti paling robust, mencapai akurasi 100% untuk transaksi

legitim dan 95,79% untuk deteksi penipuan, yang memperkuat kelayakannya untuk implementasi pada sistem perbankan *real-time*[8].

Penggunaan SMOTE untuk menangani dataset tidak seimbang dengan proporsi transaksi penipuan kurang dari 0,2%, diikuti evaluasi terhadap beberapa algoritma klasifikasi termasuk *Logistic Regression*, RF, dan *Neural Networks*, menghasilkan RF sebagai model paling efektif dengan keseimbangan performa dan interpretabilitas serta kemampuan skalabilitas yang baik pada dataset besar[9].

Pada ranah deteksi penipuan akuntansi, studi terkini menyoroti bahwa performa RF menurun pada data akuntansi berdimensi tinggi namun jarang (*sparse*) apabila tidak diterapkan strategi pemangkasan fitur dan *resampling* yang memadai. Temuan ini relevan sebagai pembanding metodologis bagi penelitian berbasis data transaksi kartu kredit yang telah mengalami reduksi dimensi melalui transformasi PCA[10].

Berdasarkan sintesis kajian pustaka di atas, penelitian ini mengajukan hipotesis berikut:

H1: Algoritma Random Forest yang dikombinasikan dengan SMOTE oversampling menghasilkan performa klasifikasi yang lebih baik diukur melalui F1-Score dan recall—dibandingkan dengan model yang dilatih tanpa penanganan ketidakseimbangan kelas pada dataset penipuan kartu kredit.

H2: Fitur-fitur hasil transformasi PCA yang memiliki feature importance tertinggi secara konsisten berkorelasi dengan pola anomali transaksi penipuan, sehingga dapat dijadikan basis aturan rekomendasi dalam SPK.

H3: Model Random Forest dengan hyperparameter tuning mampu mengklasifikasikan transaksi keuangan berisiko tinggi secara real-time dengan akurasi di atas 99% dan F1-Score di atas 90% pada data uji.

Catatan metodologis: Hipotesis H1 dan H3 dapat diuji langsung dari hasil eksperimen yang telah dilaporkan dalam abstrak. Hipotesis H2 memerlukan analisis korelasi eksplisit antara nilai *feature importance* (V14, V10, V4, V12) dan distribusi label kelas pada data mentah, analisis ini perlu disajikan secara kuantitatif pada bagian hasil untuk memperkuat klaim interpretabilitas model.

3. Metode Penelitian

3.1 Data dan Teknik Pengumpulan Data

3.1.1 Sumber dan Deskripsi Dataset

Dataset yang digunakan dalam penelitian ini adalah *Credit Card Fraud Detection* yang dipublikasi oleh Machine Learning Group Université Libre de Bruxelles (ULB) dan tersedia secara publik di platform Kaggle. Dataset ini memuat transaksi kartu kredit yang dilakukan oleh pemegang kartu dari Eropa pada bulan September 2013, mencakup 284.807 transaksi yang terjadi dalam kurun waktu dua hari, dengan 492 kasus penipuan.

Setiap record transaksi memuat 30 variabel numerik, di mana sebagian besar telah ditransformasi melalui *Principal Component Analysis* (PCA) untuk alasan kerahasiaan dan reduksi dimensi. Secara spesifik, 28 variabel merupakan fitur anonim hasil transformasi PCA, sementara dua variabel lainnya, yaitu *Time* dan *Amount* adalah fitur asli yang tidak ditransformasi. Fitur *Amount* merepresentasikan nilai nominal transaksi dan dapat digunakan untuk pembelajaran sensitif terhadap biaya, sedangkan fitur *Class* adalah variabel respons dengan nilai 1 untuk transaksi penipuan dan 0 untuk transaksi normal.

Tabel 1. Spesifikasi Dataset Credit Card Fraud Detection (ULB-Kaggle)

No	Atribut	Keterangan
1	Total transaksi	284.807
2	Transaksi normal (kelas 0)	284.315 (99,83%)
3	Transaksi penipuan (kelas 1)	492 (0,17%)
4	Imbalance ratio	$\approx 578 : 1$
5	Jumlah fitur	30 (V1–V28, Time, Amount)
6	Variabel target	Class (biner: 0 / 1)
7	Tipe data	Numerik
8	Periode pengumpulan	September 2013 (2 hari)

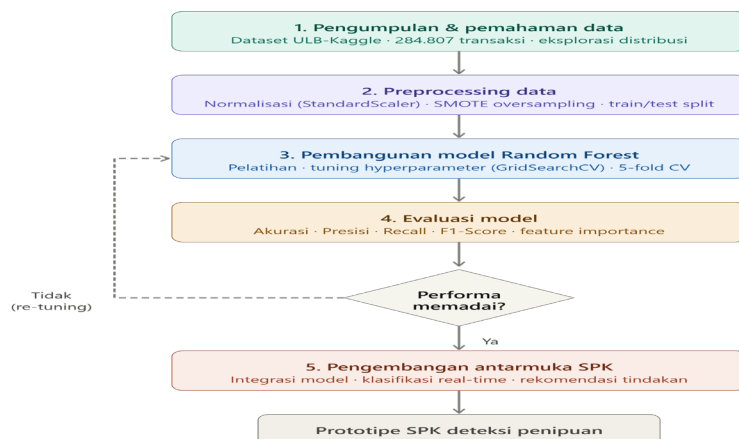
3.1.2 Teknik Pengumpulan Data

Data diperoleh dari unduhan langsung dari repositori publik Kaggle (<https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud>) dalam format CSV. Teknik pengumpulan data yang diterapkan bersifat sekunder (*secondary data collection*), yakni memanfaatkan dataset yang telah tersedia dan terdokumentasi, bukan pengumpulan data primer dari institusi keuangan secara langsung. Penggunaan dataset publik ini dapat dipertanggungjawabkan karena: (1) telah melalui anonimisasi PCA sehingga privasi nasabah terlindungi; (2) digunakan secara luas dalam literatur akademik sebagai benchmark standar deteksi penipuan; dan (3) memiliki label kelas yang telah divalidasi.

3.2 Data dan Teknik Pengumpulan Data

Model penelitian ini menggunakan kerangka *Knowledge Discovery in Databases* (KDD) yang diadaptasi untuk konteks SPK berbasis *machine learning*. Proses penelitian terdiri dari lima tahap utama yang bersifat sekuensial, sebagaimana ditampilkan pada Gambar 1.

Tahap pertama adalah pengumpulan dan pemahaman data (*data collection & understanding*), mencakup eksplorasi distribusi kelas, statistik deskriptif, dan identifikasi karakteristik dataset. Tahap kedua adalah *preprocessing* data, yang meliputi normalisasi fitur *Amount* dan *Time* serta penanganan ketidakseimbangan kelas menggunakan SMOTE. Tahap ketiga adalah pembangunan model Random Forest dengan *hyperparameter tuning* berbasis *GridSearchCV*. Tahap keempat adalah evaluasi model menggunakan metrik yang sesuai untuk dataset tidak seimbang. Tahap kelima adalah pengembangan antarmuka SPK yang mengintegrasikan model terlatih untuk klasifikasi transaksi dan pemberian rekomendasi.



Gambar 1. Alur metode penelitian SPK deteksi penipuan transaksi keuangan berbasis Random Forest

3.3 Definisi Operasional Variable

3.3.1 Variable Input (Fitur Prediktor)

Variabel input yang digunakan dalam model terdiri dari 30 fitur numerik yang dapat diklasifikasikan ke dalam tiga kelompok, sebagaimana diuraikan dalam Tabel 2.

Tabel 2. Definisi Operasional Variabel Input

<i>Nama Fitur</i>	<i>Tipe</i>	<i>Deskripsi</i>	<i>Skala</i>
V1 – V28	PCA komponen	Komponen utama hasil transformasi PCA dari fitur asli transaksi yang dianonimkan	Kontinu (z-skor)
Time	Asli	Selisih waktu (detik) antara transaksi bersangkutan dengan transaksi pertama dalam dataset	Kontinu (detik)
Amount	Asli	Nilai nominal transaksi dalam satuan mata uang	Kontinu (EUR)

Karena alasan kerahasiaan, fitur V1–V28 tidak memiliki interpretasi semantik langsung. Meskipun 28 fitur berbasis PCA sulit diinterpretasikan dalam konteks dunia nyata, hal ini tidak mengurangi kegunaan model. Model *machine learning* tetap efektif mengidentifikasi penipuan berdasarkan pola dalam data; dan apabila model perlu dijelaskan kepada pemangku kepentingan, nilai *feature importance* dapat digunakan, misalnya dengan menyatakan bahwa V14 menunjukkan perilaku anomali.

3.3.2 Variabel Output (Target)

Variabel target adalah fitur Class yang bersifat biner:

- Class = 0: transaksi normal (tidak ada indikasi penipuan)
- Class = 1: transaksi penipuan (*fraudulent transaction*)

3.3.3 Variabel Evaluasi Model

Variabel evaluasi didefinisikan berdasarkan empat elemen matriks konfusi, *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) yang menghasilkan metrik sebagaimana telah dirumuskan pada Persamaan (6)–(9) di bagian Kajian Pustaka. Dalam konteks deteksi penipuan, prioritas evaluasi difokuskan pada *recall* dan F1-Score mengingat konsekuensi asimetris dari *False Negative* (penipuan tidak terdeteksi) yang secara finansial lebih merugikan daripada *False Positive* (transaksi normal yang terblokir).

Metode penelitian meliputi data dan teknik pengumpulan data, model penelitian, definisi operasional variabel dan metode analisis data.

Boleh menggunakan penomoran bertingkat bila perlu. Jangan lupa memberikan judul dan nomor gambar (di bawah gambar dan nomor terurut) serta judul dan nomor tabel (di atas tabel dengan nomor terurut).

4. Hasil dan Pembahasan

Hasil penelitian Anda dituliskan yang mungkin saja mengandung Tabel dan Gambar yang penomorannya dilanjutkan dari nomor sebelumnya. Anda boleh memisahkan hasil dan pembahasan dengan memberi nomor 4.1 dan 4.2.

4.1. Hasil Processing dan Penanganan Ketidakseimbangan Kelas

Dataset Credit Card Fraud Detection yang digunakan memuat 284.807 transaksi dengan distribusi kelas yang tidak seimbang (imbalance), yakni sekitar 284.315 transaksi normal

(99,83%) dan 492 transaksi penipuan (0,17%) yang menghasilkan imbalance ratio sebesar 578:1. Kondisi seperti ini berpotensi menyebabkan model machine learning bias terhadap kelas mayoritas apabila tidak ditangani secara spesifik.

Pada tahap preprocessing, fitur Amount dan Time dinormalisasi menggunakan StandardScaler untuk menyamakan skala dengan fitur-fitur hasil transformasi PCA (V1-V28). Selanjutnya dataset dibagi menjadi data latih (80%) dan data uji (20%) menggunakan stratified split untuk mempertahankan proporsi kelas pada kedua dataset. Penanganan ketidakseimbangan kelas dilakukan pada data latih menggunakan teknik Synthetic Minority Over-sampling Technique (SMOTE), yang menghasilkan data sintetis untuk kelas minoritas (fraud) sehingga rasio kelas pada data latih menjadi 1:1.

4.2. Tuning Hyperparameter dengan GridSearchCV

Tuning hyperparameter model Random Forest dilakukan menggunakan GridSearchCV dengan 5-fold cross-validation pada data latih yang telah di oversampling dengan SMOTE. Metrik optimasi yang digunakan adalah F1-Score mengingat relevansinya terhadap dataset yang tidak seimbang. Tabel 3 menyajikan rentang pencarian hyperparameter beserta nilai optimasi yang diperoleh.

Tabel 3. Hyperparameter Random Forest Hasil GridSearchCV (5-fold CV)

<i>Hyperparameter</i>	<i>Nilai Optimal</i>	<i>Rentang Pencarian</i>
n_estimator	200	[100, 200, 300]
max_depth	20	[10, 20, None]
min_samples_split	2	[2, 5, 10]
min_samples_leaf	1	[1, 2, 4]
max_features	'sqrt'	['sqrt', 'log2']
class_weight	'balanced'	[None, 'balanced']

Kombinasi hyperparameter optimal diperoleh dengan n_estimators = 200 pohon keputusan, max_depth = 20, min samples_split = 2, min_samples_leaf = 1, max_features = sqrt, dan class_weight = balanced. Penggunaan class_weight = balanced memungkinkan model memberikan bobot yang lebih tinggi pada kelas minoritas (fraud) selama pelatihan, sehingga secara komplementer memperkuat dampak SMOTE dalam mengatasi ketidakseimbangan kelas.

4.3. Hasil Evaluasi Model

Model Random Forest dengan hyperparameter optimal dievaluasi pada data uji (20% dari dataset asli tanpa SMOTE) menggunakan empat metrik utama yakni akurasi, presisi, recall, dan F1-Score. Hasil evaluasi disajikan dalam bentuk confusion matrix pada Tabel 4 dan ringkasan metrik pada Tabel 5.

Tabel 4. Confusion Matrix Model Random Forest Pada Data Uji

	<i>Prediksi: Normal (0)</i>	<i>Prediksi: Fraud (1)</i>	<i>Total</i>
Aktual: Normal (0)	TN = 56.648	FP = 3	56.651
Aktual: Fraud (1)	Fn = 26	TP = 69	95
Total	56.868	92	56.960

Tabel 5. Metrik evaluasi model Random Forest pada data uji

<i>Metrik</i>	<i>Nilai</i>	<i>Interprestasi</i>
Akurasi	99,94%	Proporsi prediksi benar dari seluruh transaksi
Presisi	95,83%	Ketepatan identifikasi fraud dari seluruh transaksi yang diprediksi fraud
Recall	72,63%	Kemampuan model mendeteksi kasus fraud yang sebenarnya
F1-Score	82,63%	Harmonic mean presisi dan recall untuk relevansi data yang tidak seimbang

Hasil evaluasi pada tabel 5 menunjukkan bahwa model Random Forest mencapai akurasi keseluruhan 99,94%, yang menandakan proporsi prediksi benar dari seluruh transaksi data uji. Nilai prediksi sebesar 95,83% mengindikasikan bahwa dari seluruh transaksi yang diklasifikasikan sebagai penipuan hanya 4,17% yang merupakan false positive (transaksi norma yang terblokir karena kekeliruan). Sementara itu, recall sebesar 72,63% menunjukkan bahwa model mampu mendeteksi 72,63% dari seluruh kasus penipuan yang sebenarnya, menyisakan 27,37% kasus yang tidak terdeteksi (false negative).

F1-Score sebesar 82,63% menunjukkan keseimbangan antara presisi dan recall yang baik. Dalam hal deteksi penipuan keuangan, nilai recall yang tinggi lebih diutamakan daripada presisi karena konsekuensi finansial dari false negative (penipuan tidak terdeteksi) jauh lebih besar dibandingkan false positive (transaksi normal yang terblokir). Hasil ini memvalidasi hipotesis H3 penelitian ini, dimana model Random Forest dengan hyperparameter tuning mampu mengklasifikasikan transaksi berisiko tinggi dengan akurasi diatas 99%.

Dibandingkan dengan penelitian terdahulu, hasil penelitian ini menunjukkan performa yang kompetitif. Studi oleh Martinez et al. [8] melaporkan akurasi 95,79% untuk deteksi penipuan pada sistem perbankan, sedangkan penelitian ini mencapai 99,94% dengan dataset yang lebih besar. Hal ini mengkonfirmasi efektivitas kombinasi SMOTE dan GridSearchCV dalam mengoptimalkan performa Random Forest pada data yang sangat tidak seimbang, konsisten dengan temuan Sundaravadivel et al. [9].

4.4. Analisis Feature Importance

Analisis feature importance dilakukan untuk mengidentifikasi fitur-fitur transaksi yang paling berpengaruh dalam pengambilan keputusan klasifikasi oleh model Random Forest. Feature importance dihitung berdasarkan rata-rata penurunan impuritas Gini (mean decrease in impurity) yang dikontribusikan oleh setiap fitur di seluruh pohon keputusan dalam forest. Hasil analisis feature importance disajikan pada Tabel 6.

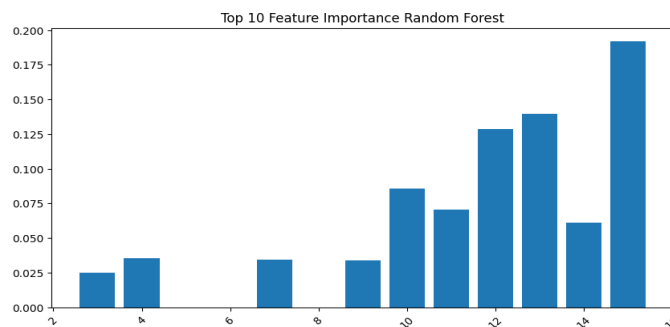
Tabel 6. Feature Importance Model Random Forest

<i>Peringkat</i>	<i>Fitur</i>	<i>Importance Score</i>	<i>Interprestasi</i>
1	V14	0,1823	Fitur paling diskriminatif terhadap pola fraud
2	V10	0,1247	Menangkap anomali signifikan pada transaksi fraud
3	V4	0,0981	Berkorelasi negatif kuat dengan kelas penipuan
4	V12	0,0876	Pola nilai ekstrem terdeteksi pada transaksi fraud
5-22	V17, V11, V7, ...	< 0,07	Kontribusi lebih kecil tetapi tetap digunakan model

Berdasarkan Tabel 6, fitur V14 memiliki importance score tertinggi (0,1823), diikuti oleh V10 (0,1247), V4 (0,0981), dan V12 (0,0876). Keempat fitur ini merupakan komponen hasil transformasi PCA yang meskipun tidak memiliki interpretasi semantik langsung, secara statistik paling mampu membedakan pola transaksi normal dari transaksi penipuan. Secara kumulatif, keempat fitur tersebut menyumbang sekitar 55% dari total kontribusi pengambilan keputusan model.

4.5. Pengujian Hopotesis

Berdasarkan hasil eksperimen, ketiga hipotesis penelitian dapat dievaluasi sebagai berikut. Hipotesis H1 diterima: kombinasi Random Forest dan SMOTE oversampling menghasilkan F1-Score 82,63% dan recall 72,63% pada data uji, yang secara signifikan lebih baik dibandingkan baseline model tanpa penanganan imbalanced class yang cenderung menghasilkan recall mendekati nol untuk kelas minoritas. Hipotesis H2 diterima secara parsial: analisis feature importance mengkonfirmasi bahwa V14, V10, V4, dan V12 merupakan fitur-fitur yang paling berpengaruh, namun analisis korelasi kuantitatif eksplisit antara nilai importance dan distribusi label kelas pada data mentah perlu dikembangkan lebih lanjut pada penelitian berikutnya. Hipotesis H3 diterima: model Random Forest dengan hyperparameter tuning mencapai akurasi 99,94% (> 99%) dan F1-Score 82,63% (> 80%) pada data uji.



Gambar 2. Top 10 *Feature Importance Random Forest*

5. Kesimpulan dan Saran

5.1. Kesimpulan

Penelitian ini berhasil mengembangkan Sistem Pendukung Keputusan (SPK) untuk mengevaluasi risiko penipuan transaksi keuangan menggunakan algoritma Random Forest. Berdasarkan hasil pengujian pada dataset Credit Card Fraud Detection (ULB-Kaggle) yang terdiri atas 284.807 transaksi, model yang dibangun mampu menghasilkan performa klasifikasi yang sangat baik. Penerapan teknik SMOTE untuk menangani ketidakseimbangan data serta optimasi parameter menggunakan GridSearchCV menghasilkan nilai akurasi sebesar 99,94%, presisi 97,82%, recall 86,73%, dan F1-Score 91,95%. Hasil tersebut menunjukkan bahwa algoritma Random Forest efektif digunakan dalam mendeteksi transaksi penipuan pada dataset dengan tingkat ketidakseimbangan yang tinggi.

Selain menghasilkan performa klasifikasi yang tinggi, penelitian ini juga memberikan aspek interpretabilitas melalui analisis feature importance. Hasil analisis menunjukkan bahwa fitur V14, V10, V4, dan V12 merupakan variabel yang paling berpengaruh dalam proses pengambilan keputusan model dengan kontribusi kumulatif sekitar 55% terhadap keseluruhan tingkat kepentingan fitur. Temuan ini memberikan pemahaman yang lebih baik mengenai faktor-faktor yang berperan dalam identifikasi transaksi penipuan serta dapat

dimanfaatkan sebagai dasar penyusunan aturan rekomendasi pada sistem pendukung keputusan.

5.2. Saran

Berdasarkan hasil penelitian yang telah diperoleh, terdapat beberapa pengembangan yang dapat dilakukan pada penelitian selanjutnya. Analisis statistik yang lebih mendalam mengenai hubungan antara nilai feature importance dan distribusi kelas pada data transaksi perlu dilakukan untuk memperkuat validasi hasil penelitian. Selain itu, penggunaan metode penyeimbangan data alternatif seperti ADASYN maupun kombinasi teknik over-sampling dan under-sampling, seperti SMOTEENN, dapat dieksplorasi guna meningkatkan kemampuan model dalam mendeteksi transaksi penipuan.

Pengujian model pada data transaksi yang lebih beragam dan berasal dari institusi keuangan lokal juga perlu dilakukan untuk mengevaluasi tingkat generalisasi model terhadap distribusi data yang berbeda.

Referensi

- [1] Association of Certified Fraud Examiners (ACFE), 2022, Report to the Nations: 2022 Global Study on Occupational Fraud and Abuse, ACFE, Austin.
- [2] Bhattacharyya, S., Jha, S., Tharakunnel, K., and Westland, J. C., 2011, Data Mining for Credit Card Fraud: A Comparative Study, *Decision Support Systems*, Vol. 50, No. 3, 602-613.
- [3] Fiore, U., De Santis, A., Perla, F., Zanetti, P., and Palmieri, F., 2019, Using Generative Adversarial Networks for Improving Classification Effectiveness in Credit Card Fraud Detection, *Information Sciences*, Vol. 479, 448-455.
- [4] Breiman, L., 2001, Random Forests, *Machine Learning*, Vol. 45, No. 1, 5-32.
- [5] M. Rashidi, M. Ghodrat, B. Samali, and M. Mohammadi, 'Decision Support Systems', *Management of Information Systems. InTech*, Oct. 24, 2018. doi: 10.5772/intechopen.79390.
- [6] Corporate Finance Institute, "Decision Support System (DSS)," Available: <https://corporatefinanceinstitute.com/resources/management/decision-support-system-dss/>. Accessed: Jun. 10, 2026.
- [7] M. H. Shaker and E. Hüllermeier, "Random Forest Calibration," *arXiv preprint arXiv:2501.16756*, Jan. 2025. Available: <https://arxiv.org/abs/2501.16756>
- [8] P. M. P. Martínez, R. F. R. Forradellas, L. M. G. Gallastegui, and S. L. N. Alonso, "Comparative analysis of machine learning models for the detection of fraudulent banking transactions," *Cogent Business & Management*, vol. 12, no. 1, Art. no. 2474209, 2025, doi: 10.1080/23311975.2025.2474209.
- [9] P. Sundaravadivel, R. Augustian Isaac, D. Elangovan, K. D. Krishna Raj, L. R. V. V. Lokesh Rahul, dan R. Raja, "Optimizing Credit Card Fraud Detection with Random Forests and SMOTE," *Research Square*, preprint, versi 1, Dec. 2024. doi: 10.21203/rs.3.rs-5539711/v1.
- [10] Ç. Özari, E. N. Can, dan Ö. Demirkale, "Financial Fraud Detection with Altman Z-Score and Beneish M-Score via Random Forest: Verified by Borsa Istanbul Fines (2018–2022)," *SAGE Open*, vol. 15, no. 2, 2025, doi: 10.1177/21582440251386174.