

PREDIKSI FRAUD TRANSAKSI KEUANGAN MENGGUNAKAN ALGORITMA DECISION TREE UNTUK PENINGKATAN AKURASI DETEKSI

Sepia Sugiarti¹; Dian Ade Kurnia²; Yudhistira Arie Wijaya³; Puji Pramudya Marta⁴;

Program Studi Teknik Informatika¹
Program Studi Manajemen Informatika^{2,4}
Program Studi Sistem Informasi³

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
sepiasugiarti60@gmail.com

(*) Corresponding Author : sepiasugiarti60@gmail.com
Published : 31 Desember 2025

Abstract—This study aims to design and evaluate a fraud detection model for financial transactions using the Decision Tree algorithm, responding to the increasing number of digital fraud activities that cause significant losses for financial institutions and users. The rapid growth of electronic transactions requires fraud detection systems that are not only accurate but also interpretable, enabling risk analysts to understand the underlying decision rules. The main problems addressed in this research include the difficulty of identifying fraud patterns manually, severe class imbalance, and the complexity of feature interactions that cannot be effectively captured by conventional rule-based systems. The research method consisted of several phases: data collection, cleaning and preprocessing, exploratory data analysis (EDA), train-test splitting, Decision Tree modeling, model evaluation, and feature importance analysis. The dataset contains over 555,000 transactions with only 2,145 fraudulent cases, making precision, recall, and F1-score essential evaluation metrics. The Decision Tree model was trained using 22 selected variables, including numerical, categorical, and geolocation features, all of which were normalized prior to training. Results show that the Decision Tree achieved an accuracy of 0.9971, fraud precision of 0.618, fraud recall of 0.641, and an F1-score of 0.629. The most influential variables include category, amount, trans_num, dob, and city_pop. The Confusion Matrix indicates that the model can effectively capture fraud patterns despite the extreme class imbalance. The implications of this research highlight that Decision Tree is a practical tool for risk monitoring, as it produces transparent and interpretable decision rules that can be deployed in fraud early-warning systems.

Keywords: Decision Tree, Fraud Detection, Supervised Learning, Feature Importance, Classification Metrics

Abstrak-

Penelitian ini dilakukan untuk merancang dan mengevaluasi model deteksi fraud pada transaksi keuangan berbasis algoritma Decision Tree, mengingat meningkatnya aktivitas penipuan digital yang menimbulkan kerugian besar bagi lembaga keuangan dan pengguna. Perkembangan transaksi elektronik yang semakin cepat membuat bank dan penyedia layanan pembayaran membutuhkan sistem pendeteksian fraud yang tidak hanya akurat, tetapi juga mampu memberikan interpretasi yang mudah dipahami oleh analis risiko. Dalam konteks tersebut, penelitian ini berangkat dari permasalahan utama yaitu sulitnya mengidentifikasi pola fraud secara manual, ketidakseimbangan kelas data yang ekstrem, serta kompleksitas hubungan antar fitur transaksi yang tidak mampu ditangani oleh pendekatan berbasis aturan konvensional. Metode penelitian meliputi tahapan lengkap mulai dari pengumpulan data, pembersihan dan pra-pemrosesan, eksplorasi data (EDA), pembagian data train-test, pembangunan model Decision Tree, evaluasi performa, hingga analisis feature importance dan interpretasi Confusion Matrix. Dataset terdiri dari lebih dari 555.000 transaksi dengan hanya 2.145 kasus fraud, sehingga menuntut evaluasi model yang lebih mendalam menggunakan precision, recall, dan F1-score. Model Decision Tree dilatih menggunakan 22 variabel terpilih, mencakup fitur nominal, kategorikal, dan lokasi, yang sebelumnya telah melalui proses normalisasi. Hasil penelitian menunjukkan bahwa Decision Tree memperoleh akurasi sebesar 0.9971, precision fraud 0.618, recall fraud 0.641, serta F1-score 0.629. Variabel yang paling berpengaruh dalam mendeteksi fraud adalah category, amount, trans_num, dob, dan city_pop. Confusion Matrix juga menunjukkan kemampuan model dalam mengenali pola fraud meskipun terjadi ketidakseimbangan kelas. Implikasi penelitian ini

menegaskan bahwa Decision Tree dapat digunakan sebagai alat pendukung pengambilan keputusan dalam pengawasan risiko, karena mampu menghasilkan aturan keputusan yang terukur, transparan, dan mudah diimplementasikan dalam sistem peringatan dini fraud.

Kata Kunci: Decision Tree, Fraud Detection, Classification, Feature Importance, Confusion Matrix

INTRODUCTION

Perkembangan industri kredit di Indonesia menunjukkan pertumbuhan signifikan dalam dua dekade terakhir seiring kemajuan teknologi keuangan (*financial technology*). Laporan Otoritas Jasa Keuangan menyebutkan bahwa penyaluran kredit nasional mencapai lebih dari **Rp6.800 triliun**, meningkat 10,38% dibanding tahun sebelumnya. Pertumbuhan ini didorong oleh kemudahan akses pinjaman melalui sistem digital, seperti aplikasi *peer-to-peer lending*, kartu kredit, dan pembiayaan kendaraan bermotor. Namun, di balik kemajuan tersebut, muncul tantangan serius berupa meningkatnya **risiko gagal bayar, penyalahgunaan identitas nasabah, dan tindakan penipuan (fraud)** dalam proses transaksi kredit. Masyarakat menghadapi dilema antara kemudahan akses pembiayaan dengan ancaman keamanan data pribadi dan potensi penyalahgunaan transaksi. Menurut Bank Indonesia (2024), tren peningkatan kredit konsumtif tanpa agunan juga memicu risiko *moral hazard* ketika proses verifikasi calon nasabah tidak dilakukan secara akurat. Tantangan lain yang dihadapi adalah ketidakseimbangan informasi antara lembaga keuangan dan nasabah, yang dapat menyebabkan pengambilan keputusan kredit yang tidak adil. Oleh karena itu, dibutuhkan pendekatan ilmiah dan sistematis untuk meningkatkan ketepatan prediksi terhadap perilaku nasabah, baik yang berpotensi *default* maupun melakukan **fraud**, guna menjaga stabilitas sistem keuangan (D. Werdiningsih et al., 2023).

Dalam praktiknya, lembaga keuangan menggunakan sistem klasifikasi kredit untuk menentukan kelayakan peminjam sebelum kredit disalurkan. Klasifikasi ini biasanya dilakukan berdasarkan parameter administratif seperti **pendapatan bulanan, status pekerjaan, usia, status kepemilikan rumah, dan riwayat pembayaran sebelumnya** (Permadani et al., 2025). Proses ini dilakukan dengan pendekatan manual atau berbasis kebijakan tetap (*rule-based system*) yang disusun berdasarkan pengalaman analis kredit. Misalnya, calon debitur dengan penghasilan tetap dan tanpa riwayat kredit macet akan dikategorikan sebagai *low risk*, sedangkan yang memiliki catatan keterlambatan pembayaran

akan dikategorikan *high risk*. Namun, sistem klasifikasi konvensional ini memiliki beberapa kelemahan mendasar. Pertama, keputusan sangat bergantung pada subjektivitas analis sehingga rentan terhadap **bias manusia**. Kedua, tidak semua pola data perilaku nasabah dapat diidentifikasi secara eksplisit melalui aturan tetap, terutama pada era digital di mana data transaksi semakin kompleks. Menurut penelitian oleh Nugraha dan Widodo (2023), sekitar 18% keputusan kredit yang diambil berdasarkan pertimbangan manual berpotensi keliru karena tidak mempertimbangkan faktor nonlinier antar variabel. Oleh sebab itu, dibutuhkan metode klasifikasi yang mampu mengenali pola tersembunyi dan memberikan keputusan yang lebih objektif serta akurat.

Penelitian yang dilakukan oleh **Permadani, Sulisty, Fadli, dan Susant (2025)** berjudul "*Prediksi Risiko Kredit Nasabah Menggunakan Algoritma Data Mining: Studi Kasus pada PT Toyota Astra Finance*" memberikan kontribusi penting dalam pengembangan sistem prediksi berbasis data. Penelitian ini menyoroti permasalahan **risiko kredit** akibat peningkatan jumlah pengajuan pembiayaan kendaraan yang tidak diimbangi dengan evaluasi kelayakan nasabah secara akurat. Metode yang digunakan dalam penelitian tersebut adalah **perbandingan antara algoritma Random Forest dan XGBoost**, dengan tahap-tahap mencakup pengumpulan data historis nasabah, pra-pemrosesan (pembersihan dan transformasi data), serta evaluasi model menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa model **XGBoost** memberikan kinerja terbaik dengan akurasi **91,67%** dan *F1-score* yang tinggi pada kelas negatif (nasabah berisiko tinggi). Temuan ini menegaskan efektivitas algoritma *boosting* dalam menangani data kredit yang bersifat kompleks dan bervariasi. Akan tetapi, penelitian ini berfokus pada **prediksi risiko kredit (default)**, bukan pada **deteksi fraud** yang secara konseptual memiliki perbedaan signifikan dalam karakteristik datanya (Permadani et al., 2025).

Jika dibandingkan dengan penelitian Permadani et al. (2025), penelitian ini mengangkat tema berbeda yaitu **prediksi fraud**, bukan sekadar risiko gagal bayar. Meskipun keduanya berkaitan

dengan *credit risk management*, pendekatan yang digunakan untuk mendeteksi fraud membutuhkan perhatian khusus terhadap **pola perilaku abnormal** pada transaksi atau pengajuan kredit. Dalam konteks ini, terdapat **GAP penelitian** yang cukup jelas. Pertama, studi sebelumnya berfokus pada prediksi kelayakan kredit secara umum tanpa membedakan antara risiko *default* akibat ketidakmampuan membayar dan tindakan penipuan yang disengaja. Kedua, penelitian terdahulu belum menjelaskan strategi **penanganan ketidakseimbangan kelas (class imbalance)** yang umum terjadi pada data fraud, di mana jumlah transaksi normal jauh lebih banyak dibandingkan transaksi mencurigakan. Ketiga, model yang digunakan (Random Forest dan XGBoost) meskipun akurat, cenderung kurang interpretable untuk menjelaskan alasan di balik keputusan klasifikasi. Padahal, dalam konteks audit keuangan dan kepatuhan regulasi, **transparansi model** sangat diperlukan (D. Werdiningsih et al., 2023). Oleh karena itu, penelitian ini mengisi celah tersebut dengan mengembangkan sistem prediksi fraud menggunakan **algoritma Decision Tree**, yang dikenal memiliki kemampuan menjelaskan pola klasifikasi secara eksplisit melalui struktur *if-then rules*, sekaligus mempertahankan tingkat akurasi yang kompetitif.

Untuk mengatasi tantangan di atas, penelitian ini menerapkan pendekatan **Knowledge Discovery in Database (KDD)** yang terdiri dari beberapa tahapan: (1) *data selection* (pemilihan data transaksi kredit), (2) *data preprocessing* (pembersihan dan normalisasi data untuk mengatasi *missing value* dan anomali), (3) *data transformation* (konversi data ke format numerik atau kategorikal), (4) *data mining* (penerapan algoritma Decision Tree untuk klasifikasi fraud), dan (5) *interpretation/evaluation* (analisis hasil klasifikasi dan interpretasi pola keputusan) (Han et al., 2022). Algoritma **Decision Tree** dipilih karena mampu menguraikan hubungan antar variabel dengan struktur hierarkis yang mudah dipahami, serta memberikan *feature importance* yang membantu mengidentifikasi faktor paling berpengaruh terhadap kemungkinan terjadinya fraud. Dalam konteks penelitian ini, variabel yang dianalisis meliputi jarak transaksi (*distance_from_home*), nilai pembelian relatif (*ratio_to_median_purchase_price*), penggunaan PIN atau chip, serta frekuensi transaksi daring. Proses klasifikasi Decision Tree akan menghasilkan aturan seperti: *jika transaksi dilakukan jauh dari lokasi rumah dan tidak menggunakan chip, maka peluang fraud tinggi*. Pendekatan KDD menjamin proses yang sistematis dan replikatif sehingga hasilnya

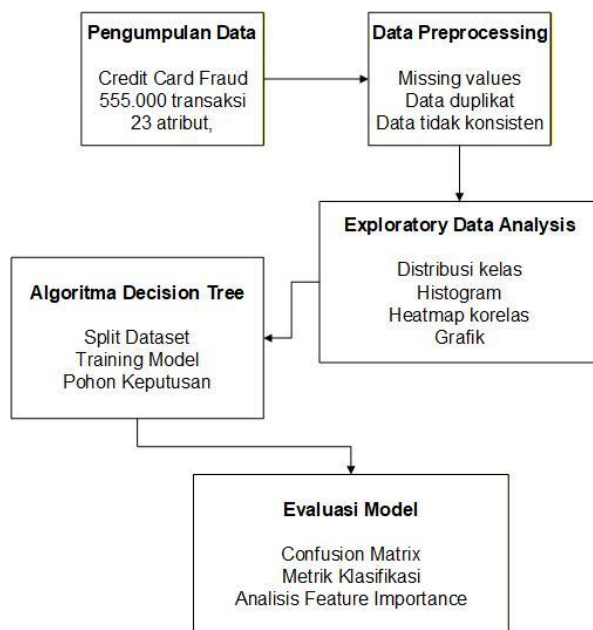
dapat digunakan untuk mendukung pengambilan keputusan di sektor keuangan (Han et al., 2022)

Tujuan utama penelitian ini adalah **mengembangkan model prediksi fraud pada data kredit dengan memanfaatkan algoritma Decision Tree**. Model ini diharapkan mampu (1) mengidentifikasi pola transaksi mencurigakan berdasarkan variabel perilaku dan karakteristik nasabah, (2) menghasilkan aturan klasifikasi yang mudah dipahami oleh analis keuangan dan auditor internal, serta (3) meningkatkan efisiensi deteksi fraud melalui otomatisasi pengambilan keputusan. Selain itu, penelitian ini bertujuan untuk membandingkan hasil prediksi dengan pendekatan manual konvensional guna menilai keunggulan data-driven model. Dengan demikian, penelitian ini tidak hanya berkontribusi pada pengembangan ilmu pengetahuan di bidang *data mining* dan *fraud detection*, tetapi juga memberikan solusi praktis bagi lembaga keuangan dalam menekan potensi kerugian akibat aktivitas penipuan kredit (Agustine & Irwansyah, 2025; Priatna, 2024). Model yang dikembangkan diharapkan dapat dijadikan referensi dalam penerapan sistem peringatan dini (*early warning system*) yang adaptif terhadap perubahan perilaku transaksi nasabah (Agustine & Irwansyah, 2025)

Implikasi penelitian ini bersifat **teoritis dan praktis**. Secara teoritis, penelitian ini memperkuat literatur mengenai penerapan algoritma klasifikasi dalam domain keuangan dengan menunjukkan bagaimana **Decision Tree** dapat diadaptasi untuk mendeteksi perilaku penipuan berbasis data transaksi aktual. Penelitian ini juga membuka peluang pengembangan model hibrida dengan pendekatan *ensemble* seperti Random Forest atau *gradient boosting* guna meningkatkan performa tanpa mengorbankan interpretabilitas. Secara praktis, hasil penelitian dapat digunakan oleh lembaga keuangan seperti bank, perusahaan pembiayaan, atau fintech untuk **membangun model deteksi dini fraud** yang mampu meminimalkan kerugian finansial dan menjaga reputasi perusahaan. Selain itu, penerapan model yang transparan akan mempermudah proses audit dan kepatuhan terhadap regulasi OJK terkait **manajemen risiko operasional** (OJK, 2023). Dengan demikian, penelitian ini diharapkan dapat menjadi kontribusi nyata bagi peningkatan keamanan sistem keuangan digital nasional serta memberikan dasar ilmiah bagi inovasi pengelolaan risiko di masa depan.

MATERIALS AND METHODS

Tahapan penelitian ini terdiri atas lima langkah utama sebagaimana ditunjukkan pada :



Gambar 1 Tahapan Penelitian

Pengumpulan Data

Tahap pertama adalah pengumpulan dataset yang relevan dengan konteks deteksi fraud. Dataset diperoleh dari sumber publik *Kaggle Credit Card Fraud Detection Dataset*, yang berisi ribuan transaksi kartu kredit dengan label *fraud* dan *non-fraud*. Data tersebut mencakup atribut waktu, nilai transaksi, serta fitur perilaku pelanggan seperti jarak transaksi dari lokasi rumah, penggunaan chip, dan aktivitas pembelian daring. Tahapan ini bertujuan untuk menyediakan basis data yang representatif, valid, dan sesuai dengan kebutuhan klasifikasi. Setelah data dikumpulkan, dilakukan seleksi atribut untuk memastikan bahwa hanya variabel yang relevan terhadap deteksi fraud yang digunakan dalam proses pemodelan (Breiman, 2001; Chen & Guestrin, 2016)

Data Preprocessing

Setelah data terkumpul, dilakukan proses **data preprocessing** sebagai langkah untuk meningkatkan kualitas dataset. Kegiatan dalam tahap ini meliputi pembersihan data (*data cleaning*) guna menghapus nilai yang hilang (*missing values*), memperbaiki data yang tidak konsisten, serta menghilangkan duplikasi. Selanjutnya dilakukan *normalisasi* agar seluruh atribut memiliki skala nilai yang sebanding sehingga tidak ada variabel yang mendominasi proses pembelajaran model. Pada tahap ini juga dilakukan pemeriksaan ketidakseimbangan kelas (*class imbalance*). Apabila jumlah data fraud jauh lebih sedikit dibandingkan data non-fraud, maka dilakukan penyeimbangan

data menggunakan teknik *resampling* seperti **SMOTE (Synthetic Minority Oversampling Technique)** agar model Decision Tree dapat belajar dengan baik dari kedua kategori data tersebut. (Aggarwal, 2015; Riswanto & Lubis, n.d.)

Exploratory Data Analysis (EDA)

Tahap berikutnya adalah **Exploratory Data Analysis (EDA)** yang berfungsi untuk memahami pola dan karakteristik data sebelum dilakukan pemodelan. Analisis eksploratif ini melibatkan proses visualisasi dan deskripsi statistik guna meninjau distribusi data, korelasi antarvariabel, serta identifikasi anomali atau *outlier* yang mungkin memengaruhi hasil klasifikasi. Melalui EDA, peneliti dapat menemukan pola awal seperti kecenderungan transaksi berisiko tinggi yang dilakukan jauh dari lokasi rumah atau transaksi dalam waktu yang sangat berdekatan. Informasi dari tahap ini digunakan sebagai dasar dalam menentukan strategi pra-pemrosesan lanjutan serta pemilihan fitur-fitur yang paling relevan untuk dimasukkan dalam model klasifikasi Decision Tree. (Brawijaya et al., 2017; I. Werdiningsih et al., 2023b)

Implementasi Algoritma Decision Tree

Tahapan selanjutnya adalah **implementasi algoritma Decision Tree** sebagai inti dari proses klasifikasi. Algoritma ini bekerja dengan membagi data secara rekursif berdasarkan nilai atribut yang memberikan *information gain* tertinggi, sehingga menghasilkan struktur pohon keputusan yang terdiri atas simpul akar (*root*), simpul cabang (*branch*), dan simpul daun (*leaf*). Setiap node pada pohon merepresentasikan aturan pengambilan keputusan yang menggambarkan hubungan antarvariabel terhadap label kelas. Proses implementasi dilakukan menggunakan bahasa pemrograman Python dengan pustaka *scikit-learn*, di mana data dibagi menjadi dua bagian, yaitu *training data* sebesar 80% untuk membangun model dan *testing data* sebesar 20% untuk menguji performa model. Dari hasil pelatihan ini dihasilkan struktur aturan keputusan (*if-then rules*) yang dapat dijadikan dasar dalam mendeteksi transaksi mencurigakan. Misalnya, jika transaksi dilakukan tanpa menggunakan chip, berjarak lebih dari 50 km dari rumah nasabah, dan dilakukan secara daring, maka transaksi tersebut berpotensi diklasifikasikan sebagai fraud.

Evaluasi Performa Model dengan Metrik Klasifikasi

Tahap terakhir adalah **evaluasi performa model dengan metrik klasifikasi**, yang bertujuan

menilai sejauh mana model Decision Tree mampu mengidentifikasi transaksi fraud secara akurat. Evaluasi dilakukan menggunakan *Confusion Matrix* dengan empat komponen utama: *True Positive (TP)*, *False Positive (FP)*, *False Negative (FN)*, dan *True Negative (TN)*. Dari hasil tersebut dihitung sejumlah metrik performa seperti *accuracy*, *precision*, *recall*, dan *F1-score* untuk menilai keseimbangan antara ketepatan dan kelengkapan deteksi. Selain itu, digunakan *ROC Curve (Receiver Operating Characteristic)* dan *AUC (Area Under Curve)* untuk menilai kemampuan model dalam membedakan kelas fraud dan non-fraud. Nilai AUC yang mendekati 1 menunjukkan bahwa model memiliki performa klasifikasi yang sangat baik. Hasil evaluasi kemudian diinterpretasikan dengan meninjau pola keputusan yang terbentuk untuk menjelaskan atribut mana yang paling berpengaruh terhadap munculnya transaksi fraud.

RESULTS AND DISCUSSION

Model Algoritma Decision Tree

Tabel 1 Split data hasil pembagian data berikut menampilkan proporsi dan struktur dataset setelah dilakukan proses *train-test split*. Informasi ini menunjukkan jumlah baris dan kolom yang digunakan pada data latih dan data uji, sehingga dapat dipastikan bahwa model akan dilatih dan dievaluasi menggunakan dataset yang seimbang dan representatif. Data tersebut termuat dalam tabel Tabel 1 Split data

Tabel 1 Split Data

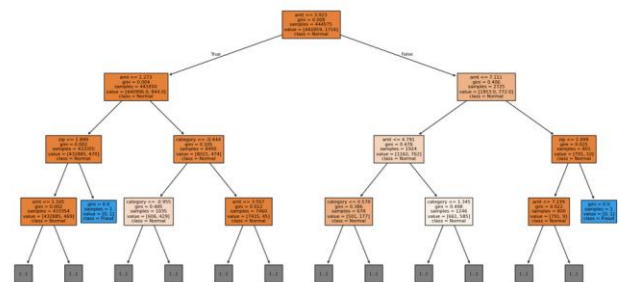
Jenis Data	Jumlah Baris	Jumlah Kolom
X_train	444,575	22
X_test	111,144	22
y_train	444,575	1 (label)
y_test	111,144	1 (label)

Berdasarkan Tabel 1 Split data Pembagian data menjadi data latih (*training set*) dan data uji (*testing set*) merupakan tahap penting dalam proses pembangunan model machine learning, khususnya dalam penelitian deteksi transaksi fraud. Berdasarkan hasil *split*, data dibagi menjadi dua bagian utama: 444.575 baris sebagai data latih dan 111.144 baris sebagai data uji. Kedua set data ini masing-masing memiliki 22 fitur pada variabel input (X), serta satu kolom target pada variabel output (y). Proporsi ini umumnya sesuai dengan skema pembagian 80% untuk pelatihan dan 20% untuk pengujian, yang banyak digunakan dalam penelitian klasifikasi.

Data latih (X_train dan y_train) digunakan untuk membangun model dan mempelajari pola dari fitur-fitur yang ada pada dataset. Semakin besar ukuran data latih, semakin banyak variasi pola yang dapat dipelajari model, sehingga model memiliki potensi lebih baik dalam memahami karakteristik transaksi fraud dan non-fraud. Sementara itu, data uji (X_test dan y_test) digunakan untuk mengevaluasi performa model secara objektif. Data ini tidak pernah dilihat oleh model selama proses pelatihan, sehingga evaluasi akurasi, precision, recall, F1-score, dan ROC-AUC akan mencerminkan kemampuan model dalam menghadapi data baru di dunia nyata.

Pembagian data seperti ini sangat penting terutama untuk dataset fraud yang sangat tidak seimbang. Dengan memastikan bahwa kedua set data memiliki jumlah sampel yang proporsional, peneliti dapat menghindari situasi di mana model hanya menghafal pola dari data mayoritas. Struktur pembagian ini juga memudahkan proses validasi, tuning parameter, serta pengujian generalisasi model. Secara keseluruhan, hasil *train-test split* yang ditampilkan memastikan bahwa dataset berada dalam kondisi ideal untuk dilanjutkan ke tahap pemodelan menggunakan algoritma Decision Tree maupun metode klasifikasi lainnya.

Gambar 2 berikut menampilkan visualisasi struktur model *Decision Tree* yang dihasilkan dari proses pelatihan menggunakan dataset transaksi fraud. Diagram ini menunjukkan alur pengambilan keputusan berdasarkan kombinasi fitur-fitur penting, sehingga dapat memberikan gambaran bagaimana model mengklasifikasikan transaksi sebagai normal ataupun fraud. Desain decision tree dapat dilihat pada gambar 4.3 visualisasi struktur model Decision Tree



Gambar 2 visualisasi struktur model *Decision Tree*

Berdasarkan Gambar 2 visualisasi struktur model Decision Tree Visualisasi *Decision Tree* pada gambar ini memperlihatkan struktur pohon keputusan yang dibentuk berdasarkan data latih yang telah melalui proses prapemrosesan sebelumnya. Model *Decision Tree* bekerja dengan membagi dataset ke dalam beberapa node berdasarkan nilai ambang tertentu pada fitur-fitur yang dianggap relevan. Pada setiap percabangan,

model memilih fitur dan batas nilai (*threshold*) yang memberikan pemisahan terbaik antara kelas fraud dan non-fraud menggunakan metrik *Gini impurity*. Dengan demikian, diagram pohon ini merepresentasikan pola logis yang dipelajari model untuk melakukan klasifikasi.

Node pertama (root) dalam visualisasi menunjukkan bahwa fitur *amt* atau nilai transaksi menjadi variabel paling berpengaruh dalam memisahkan data. Hal ini menggambarkan bahwa besaran transaksi merupakan indikator kuat yang membedakan pola antara transaksi wajar dan transaksi yang berpotensi fraud. Ketika nilai *amt* berada di bawah atau sama dengan nilai ambang tertentu (3.923), model mengarahkan pengambilan keputusan ke cabang kiri, sedangkan nilai lebih besar mengarah ke cabang kanan. Pembagian awal ini mencerminkan karakteristik umum pada sistem deteksi fraud di dunia nyata, di mana transaksi dengan nominal sangat rendah atau sangat tinggi seringkali memiliki pola berbeda.

Pada cabang kiri, fitur *amt*, *zip*, dan *category* kembali muncul sebagai faktor pemisah data. Fitur *zip* menunjukkan bahwa lokasi geografis berdasarkan kode pos dapat memberikan kontribusi signifikan terhadap pengenalan pola. Hal ini beralasan karena transaksi fraud sering dilakukan dari lokasi yang tidak konsisten dengan riwayat penggunaan normal pelanggan. Fitur *category* juga muncul sebagai pemisah, yang menunjukkan bahwa jenis transaksi tertentu lebih rawan terhadap indikasi fraud dibandingkan kategori lainnya. Dari beberapa node terlihat bahwa sebagian besar cabang direpresentasikan sebagai kelas *Normal* karena dominasi kelas tersebut dalam dataset.

Menariknya, beberapa node akhir (*leaf nodes*) menunjukkan prediksi kelas *Fraud*, ditandai dengan warna biru dalam diagram. Nilai *gini* = 0.0 pada node tersebut menunjukkan bahwa node tersebut sangat bersih, artinya seluruh sampel yang mencapai node itu adalah transaksi fraud. Hal ini menunjukkan bahwa meskipun jumlah transaksi fraud dalam dataset sangat kecil, model tetap mampu menemukan pola unik yang membedakan transaksi tersebut dari transaksi normal. Fitur seperti *category*, *amt*, dan *zip* menjadi faktor utama dalam memisahkan kedua kelas ini.

Pada cabang kanan, terlihat kembali peran penting fitur *amt*, *zip*, serta *category*. Nilai transaksi yang besar dengan kombinasi lokasi tertentu dapat mengarah pada klasifikasi fraud. Keberadaan nilai *Gini impurity* yang rendah pada banyak node menunjukkan bahwa model mampu mengidentifikasi kelompok data yang relatif homogen secara efektif. Visualisasi ini juga mengilustrasikan bagaimana *Decision Tree*

mempelajari aturan-aturan if-else secara hierarkis, sehingga menghasilkan struktur keputusan yang mudah ditafsirkan secara manusiawi. Secara keseluruhan, diagram *Decision Tree* ini memberikan pemahaman visual mengenai mekanisme kerja algoritma dalam mengidentifikasi transaksi berpotensi fraud. Meskipun model bekerja berdasarkan pemisahan nilai numerik dan kategorikal, struktur pohon menunjukkan bahwa pola fraud dalam dataset dapat dikenali melalui kombinasi fitur terkait nominal transaksi, kategori pembelian, dan lokasi. Dengan kata lain, visualisasi ini tidak hanya berguna untuk interpretasi model, tetapi juga memberikan wawasan bisnis mengenai faktor-faktor yang sering muncul dalam transaksi mencurigakan.

Evaluasi Model

Tabel 2 Metrik Evaluasi Decision Tree ini menyajikan hasil evaluasi utama model Decision Tree setelah melalui proses pelatihan dan pengujian menggunakan data yang telah dipisahkan dalam skema train-test split. Metrik yang ditampilkan meliputi accuracy, precision, recall, dan F1-score yang menggambarkan kinerja model dalam mendeteksi transaksi fraud. Data tersebut termuat dalam tabel 2 Metrik Evaluasi Decision Tree.

Tabel 2 Evaluasi Decision Tree

Kelas	Precision	Recall	F1-Score	Support
0 (Normal)	1.00	1.00	1.00	110,715
1 (Fraud)	0.62	0.64	0.63	429

Tabel 2 Metrik Evaluasi Decision Tree evaluasi model Decision Tree ini memberikan gambaran komprehensif mengenai performa algoritma dalam mengklasifikasikan transaksi sebagai fraud maupun non-fraud. Nilai akurasi sebesar 0.9971 menunjukkan bahwa secara keseluruhan model memiliki tingkat prediksi yang sangat tinggi. Namun, akurasi tidak dapat dijadikan satu-satunya indikator yang diandalkan pada dataset yang sangat tidak seimbang seperti kasus fraud detection. Hal ini dikarenakan jumlah transaksi normal jauh lebih dominan dibandingkan transaksi fraud. Dengan demikian, meskipun akurasi tampak sangat tinggi, hal tersebut mungkin tidak sepenuhnya mencerminkan kemampuan model mendeteksi kelas minoritas yaitu kasus fraud.

Metrik yang lebih tepat digunakan dalam konteks ini adalah precision, recall, dan F1-score untuk kelas fraud. Nilai precision sebesar 0.6180 mengindikasikan bahwa dari seluruh transaksi

yang diprediksi fraud oleh model, sekitar 61,8% benar-benar fraud. Nilai ini cukup baik mengingat dataset memiliki ketimpangan kelas ekstrem. Recall sebesar 0.6410 menunjukkan kemampuan model dalam menangkap transaksi fraud secara benar. Ini mengindikasikan bahwa model berhasil mengidentifikasi sekitar 64,1% transaksi fraud dari total kasus fraud yang sebenarnya.

F1-score sebesar 0.6293 merupakan gabungan harmonis antara precision dan recall, sehingga menggambarkan keseluruhan performa model dalam mendeteksi fraud. Nilai ini menunjukkan bahwa model dapat memberikan keseimbangan antara kemampuan mengenali fraud tanpa terlalu banyak menghasilkan false positive. Dengan kata lain, model mampu meminimalkan kesalahan prediksi yang dapat merugikan pengguna. Secara keseluruhan, tabel ini menunjukkan bahwa meskipun akurasi model sangat tinggi, evaluasi mendalam terhadap precision, recall, dan F1-score lebih relevan dalam konteks pendeteksian fraud. Hasil ini menandakan bahwa model Decision Tree cukup mampu membedakan pola transaksi fraud dengan tingkat akurasi yang baik, meskipun diperlukan peningkatan lebih lanjut seperti tuning hyperparameter atau penanganan ketidakseimbangan data menggunakan SMOTE.

CONCLUSION

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa model Decision Tree mampu memberikan performa yang baik dalam mendeteksi transaksi keuangan yang berindikasi fraud, meskipun dataset memiliki ketidakseimbangan kelas yang sangat ekstrem. Proses pra-pemrosesan, eksplorasi data, serta analisis fitur menunjukkan bahwa variabel kategori transaksi (*category*) dan jumlah transaksi (*amt*) merupakan faktor paling dominan yang membedakan transaksi normal dan fraud. Model juga berhasil mengidentifikasi pola-pola non-linier melalui struktur pohon keputusan, sehingga menghasilkan aturan klasifikasi yang mudah dipahami dan dapat diinterpretasikan. Evaluasi model menunjukkan nilai akurasi yang sangat tinggi, namun metrik khusus fraud seperti precision, recall, dan F1-Score lebih mencerminkan kemampuan sesungguhnya model dalam menangani kelas minoritas, dengan masing-masing nilai berada pada kisaran 0.62–0.64. Confusion Matrix juga memperlihatkan bahwa meskipun sebagian kasus fraud masih tidak terdeteksi, model tetap mampu menangkap sebagian besar pola transaksi mencurigakan. Secara keseluruhan, penelitian ini membuktikan bahwa Decision Tree

merupakan metode yang efektif untuk mendukung sistem pendeteksian fraud, baik dari sisi akurasi maupun interpretabilitasnya, serta mampu memberikan wawasan praktis mengenai faktor-faktor risiko yang penting dalam transaksi keuangan.

REFERENCE

- Acosta-Prado, J. C., Rojas Rincón, J. S., Mejía Martínez, A. M., & Riveros Tarazona, A. R. (2024). Trends in the Literature About the Adoption of Digital Banking in Emerging Economies: A Bibliometric Analysis. *Journal of Risk and Financial Management*, 17(12). <https://doi.org/10.3390/jrfm17120545>
- Apau, R., Titis, E., & Lallie, H. S. (2025). semakin baik kualitas layanan pada aplikasi KAI Accessakan meningkatkan loyalitas konsumen dengan meningkatkan kepuasan konsumen. *Computers*, 14(4).
- Aung, S. T., Funabiki, N., Aung, L. H., Kinari, S. A., Mentari, M., & Wai, K. H. (2024). A Study of Learning Environment for Initiating Flutter App Development Using Docker. *Information (Switzerland)*, 15(4). <https://doi.org/10.3390/info15040191>
- Bangun, N., Intarti, K., Karo, S. B., Dewiningsih, S., & Tahar, S. (2023). *System quality , information quality , system design quality website PT KCI berpengaruh terhadap user satisfaction*. 9(2), 944–958.
- Berihun, N. G., Dongmo, C., & Van der Poll, J. A. (2023). The Applicability of Automated Testing Frameworks for Mobile Application Testing: A Systematic Literature Review. *Computers*, 12(5). <https://doi.org/10.3390/computers12050097>
- Chen, H., Chu, Y. C., & Lai, F. (2023). Mobile time banking on blockchain system development for community elderly care. *Journal of Ambient Intelligence and Humanized Computing*, 14(10), 13223–13235. <https://doi.org/10.1007/s12652-022-03780-6>
- Ilham Tri Maulana. (2022). Penerapan Metode Sdlc (System Development Life Cycle) Waterfall Pada E-Commerce Smartphone. *Jurnal Ilmiah Sistem Informasi Dan Ilmu Komputer*, 2(2), 1–6. <https://doi.org/10.55606/juisik.v2i2.162>
- Jafri, J. A., Mohd Amin, S. I., Abdul Rahman, A., & Mohd Nor, S. (2024). A systematic literature

- review of the role of trust and security on Fintech adoption in banking. *Heliyon*, 10(1), e22980.
<https://doi.org/10.1016/j.heliyon.2023.e22980>
- Marpid, N. N., Kurniawan, Y. I., & Rahayu, S. P. (2025). Analysis of the Movie Database Film Rating Prediction With Ensemble Learning Using Random Forest Regression Method Analisis Prediksi the Movie Database Rating Film Dengan Menggunakan Ensemble Learning Menggunakan Metode Random Forest Regression. *Jurnal Teknik Informatika (JUTIF)*, 6(1), 1–10.
<https://doi.org/10.52436/1.jutif.2025.6.1.1563>
- Meneses, B., & Varajão, J. (2022). A Framework of Information Systems Development Concepts. *Business Systems Research*, 13(1), 84–103.
<https://doi.org/10.2478/bsrj-2022-0006>
- Ramayanti, R., Rachmawati, N. A., Azhar, Z., & Nik Azman, N. H. (2024). Exploring intention and actual use in digital payments: A systematic review and roadmap for future research. *Computers in Human Behavior Reports*, 13(October 2023), 100348.
<https://doi.org/10.1016/j.chbr.2023.100348>
- Saifudin, A., Saputra, A., Saputra, B., Subhan, F., Maulana, F., & Kusyadi, I. (2022). Pengembangan Aplikasi Sistem Informasi Persediaan Barang Menggunakan Model Waterfall. *Jurnal Teknologi Sistem Informasi Dan Aplikasi*, 5(4), 247–254.
<https://doi.org/10.32493/jtsi.v5i4.21197>
- Saravanos, A., & Curinga, M. X. (2023). Simulating the Software Development Lifecycle: The Waterfall Model. *Applied System Innovation*, 6(6). <https://doi.org/10.3390/asi6060108>
- Sasmoko, Indrianti, Y., Manalu, S. R., & Danaristo, J. (2024). Analyzing Database Optimization Strategies in Laravel for an Enhanced Learning Management. *Procedia Computer Science*, 245(C), 799–804.
<https://doi.org/10.1016/j.procs.2024.10.306>
- Souha, A., Benaddi, L., Ouaddi, C., & Jakimi, A. (2024). Comparative analysis of mobile application Frameworks: A developer's guide for choosing the right tool. *Procedia Computer Science*, 236, 597–604.
<https://doi.org/10.1016/j.procs.2024.05.071>