



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 **Issue:** VI **Month of publication:** June 2026

DOI: <https://doi.org/10.22214/ijraset.2026.83904>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

Advanced Fraud Detection Using ML

Harshal Patil

Department of Artificial Intelligence and Data Science, D. Y. Patil College of Engineering, Akurdi, Pune, India

Abstract: Financial fraud causes significant economic loss and erodes trust in digital payment ecosystems. Traditional rule-only systems struggle with evolving attack patterns, while pure black-box machine learning models are difficult for analysts to interpret during investigations. This paper presents ****FraudX****, an end-to-end explainable fraud detection framework that combines supervised learning with rule-based categorization and human-readable explanations. The system generates a labeled synthetic transaction dataset of 10,000 records with 20+ behavioral and contextual features, trains a Random Forest classifier with standardized preprocessing, and deploys inference through an interactive Streamlit dashboard. For each transaction, FraudX outputs a probability-based ***Suspicion Score***, a threshold-controlled suspicious flag, a prioritized ***Fraud_Type*** label, and concise ***Suspicion_Reasons***. Experimental evaluation on held-out synthetic test data reports 95.2% accuracy, 94.8% precision, 93.5% recall, and 94.1% F1-score. The proposed architecture demonstrates how operational thresholding, feature-importance-driven explanations, and analyst-oriented visualization can be integrated into a practical fraud triage workflow suitable for academic demonstration and prototype deployment.

Keywords: financial fraud detection, Random Forest, explainable artificial intelligence, transaction risk scoring, Streamlit dashboard, supervised machine learning

I. INTRODUCTION

Digital payment volumes have increased sharply across banking, e-commerce, and peer-to-peer platforms. Fraudsters exploit cross-border routing, device spoofing, credential abuse, cryptocurrency channels, and high-velocity transfer patterns to evade controls [1]. Organizations therefore require systems that can (i) rank risky transactions at scale, (ii) support analyst review with interpretable evidence, and (iii) allow operational tuning without retraining the entire model.

Existing approaches broadly fall into three categories: static rule engines, supervised classifiers, and unsupervised anomaly detectors. Rule engines are transparent but brittle; pure ML models achieve strong ranking performance but are often criticized as black boxes in regulated environments [5]. A practical screening pipeline should combine statistical scoring with policy-aligned labels and concise natural-language rationales.

****FraudX**** addresses this gap through a modular pipeline spanning data generation, model training, inference, explanation, and visualization. The main contributions of this work are:

- A reproducible synthetic transaction generator with realistic fraud indicators including velocity spikes, unusual timing, VPN usage, sanctioned-entity involvement, and cross-border activity.
- A Random Forest-based scoring engine with schema alignment for heterogeneous uploaded datasets.
- A hybrid decision layer that merges ML probability scores with prioritized fraud-type rules.
- An analyst-facing Streamlit application supporting CSV/PDF ingestion, threshold tuning, KPI dashboards, and export of suspicious records.

The remainder of this paper is organized as follows. Section II reviews related work. Section III describes the system architecture and methodology. Section IV details implementation. Section V presents experimental results. Section VI concludes the paper.

II. RELATED WORK

A. Fraud Detection Surveys and Taxonomies

Phua *et al.* [1] provide a comprehensive survey of data-mining approaches to fraud detection, highlighting challenges such as class imbalance, concept drift, and adversarial adaptation. Carcillo *et al.* [2] discuss machine learning strategies for banking fraud, emphasizing the need for feature engineering and real-time scoring. These studies motivate FraudX's focus on behavioral features (velocity, device change, failed attempts) rather than identity fields alone.

B. Supervised Learning for Transaction Screening

Ensemble tree methods remain popular in fraud analytics due to robustness on mixed numerical and categorical inputs [3]. Bahnsen *et al.* [4] show that cost-sensitive feature engineering significantly improves credit-card fraud detection. FraudX adopts Random Forest for its balance of accuracy, training efficiency, and native feature-importance signals used in downstream explanations.

C. Explainability and Analyst Workflows

Model explanations improve trust and auditability. Ribeiro *et al.* [5] introduce local interpretable modelagnostic explanations (LIME), while Lundberg and Lee [6] propose SHAP values for consistent feature attribution. FraudX uses a lighter-weight approach: global feature-importance thresholds combined with domain rules to generate readable *Suspicion_Reasons*, making the system suitable for educational and prototype settings without additional explanation-model overhead.

III. SYSTEM ARCHITECTURE AND METHODOLOGY

A. High-Level Architecture

FraudX comprises four layers:

1. **Data Layer** — synthetic generation (`main.py`) and user CSV/PDF ingestion (`app.py`).
2. **Training Layer** — Random Forest training, evaluation, and artifact persistence.
3. **Inference & Explanation Layer** — preprocessing, scoring, fraud typing, and reason generation (`suspicious_by_model.py`).
4. **Presentation Layer** — Streamlit dashboard with analytics, charts, and export.

Persisted artifacts include `fraud_detection_model.joblib`, `fraud_detection_scaler.joblib`, `feature_importances.joblib`, and `model_metrics.joblib`.

Figure 1 (conceptual data flow):

```

CSV/PDF Upload → Schema Alignment → Feature Encoding → Scaling
    → Random Forest (predict_proba) → Suspicion Score
    → Threshold → Is_Suspicious → Fraud_Type + Suspicion_Reasons
    → Dashboard / CSV Export
  
```

B. Dataset Generation and Feature Engineering

The synthetic generator in `main.py` creates **10,000 transactions** with **5% fraud prevalence** (`FRAUD_PERCENTAGE = 0.05`). Supporting lists include 100 known fraudster IDs and 50 sanctioned entity IDs, stored as `known_fraudsters.csv` and `sanctioned_entities.csv`.

Each transaction record contains:

```

| Category | Fields |
|-----|-----|
| Identity | Transaction_ID, Sender_ID, Receiver_ID |
| Temporal | Date, Time, Unusual_Time |
| Geography | Sender_Country, Receiver_Country, Transaction_City, Latitude, Longitude |
| Financial | Transaction_Amount, Transaction_Currency, Payment_Method |
| Behavioral | Transaction_Velocity, Repeated_Failed_Attempts, Multiple_Currency_Conversions |
| Device/Network | Device_Type, Browser, VPN_Usage, IP_Address_Change, Is_New_Device, Is_New_Location |
| Merchant | Merchant_Name, Merchant_Category, Transaction_Channel |
| Labels | Is_Fraudulent, Is_Known_Fraudster, Is_Sanctioned_Entity | Fraudulent rows are
statistically biased toward:
  
```

- Higher transaction amounts (up to 200,000 vs. normal range).
- Elevated velocity (Gaussian mean ≈ 8 vs. 3 for legitimate traffic).
- Increased unusual-time probability (60% vs. 20%).
- More currency conversions, failed attempts, VPN usage, and new device/location events.

Countries supported: USA, India, UK, Canada, Germany, Australia, France, Netherlands, Brazil, Japan. Payment methods include bank transfer, UPI, PayPal, Skrill, Payoneer, and Cryptocurrency.

C. Model Training Procedure

1. Drop non-predictive identifiers: `Transaction\_ID`, `Date`, `Time`, `Sender\_ID`, `Receiver\_ID`.
2. One-hot encode categorical columns via `pd.get\_dummies`.
3. Split data 80/20 with `train\_test\_split(..., random\_state=42)`.
4. Apply `StandardScaler` fit on training data.
5. Train `RandomForestClassifier(n\_estimators=100, max\_depth=10, min\_samples\_split=5, random\_state=42)`.
6. Evaluate on held-out test set; save model, scaler, importances, and metrics.

D. Inference, Thresholding, and Fraud Typing

****Preprocessing (`preprocess\_data`):**** Missing columns are filled with safe defaults (`REQUIRED\_INPUT\_COLUMNS`).

Categorical fields are encoded and aligned to

`scaler.feature\_names\_in\_`; absent dummy columns are zero-filled.

****Scoring:**** `predict\_proba[:, 1]` produces `Suspicion\_Score` $\in [0, 1]$.

****Thresholding:**** `Is\_Suspicious = Suspicion\_Score > threshold` (default 0.5, adjustable in UI).

Fraud_Type priority rules (`classify\_fraud\_type`):

Priority	Condition	Label
1	Is_Known_Fraudster	Known Fraudster
2	Is_Sanctioned_Entity	Sanctioned Entity
3	Crypto + cross-border + (high amount or multi-FX)	Crypto Cross-border Laundering
4	Cross-border + multi-FX	Cross-border FX Risk
5	High amount + high velocity	Rapid Large Transfers
6	VPN or IP change	Anonymity/Device Evasion
7	New device or location	Account Takeover (ATO) Risk
8	Repeated failed attempts ≥ 3	Credential Stuffing/Brute Force
9	Unusual time	Odd-hour Activity
10	Score $\geq \max(0.8, \text{threshold})$	High-Risk (Model)
11	Score $\geq \text{threshold}$	Medium-Risk (Model)
12	Otherwise	Low-Risk

****Explanations (`explain\_suspicion`):**** For each suspicious row, the engine compares transaction attributes against sender history and global feature importances (threshold 0.05). Example outputs: "High Transaction Amount (\$X, Y $\times$ average)", "Cross-border Transaction (India to UK)", "VPN Usage Detected", "Known Fraudster Involved".

IV. IMPLEMENTATION

A. Technology Stack

Component	Technology
Language	Python 3.10+
Data processing	Pandas, NumPy [8]
Machine learning	scikit-learn [7]
Serialization	Joblib
Visualization	Plotly, Matplotlib
Web UI	Streamlit [9]
PDF ingestion	pdfplumber
Optional API	FastAPI (`api/main.py`)

B. Streamlit Dashboard Features (`app.py`)

The dashboard provides five tabs:

1. ****Upload & Analyze**** — demo dataset or CSV/PDF upload, column mapping, threshold slider, detection button, KPI metrics, pie/histogram charts, fraud-type and reason breakdowns, suspicious transaction export.

2. **Model Metrics** — accuracy, precision, recall, F1, bar chart, confusion matrix (if available).
 3. **Feature Importance** — top-15 bar chart and full ranked table.
 4. **Analytics** — time-series suspicious counts (when Date column is valid).
 5. **About** — project overview and performance summary.
- Privacy note: all processing occurs locally in the Streamlit session; data is not transmitted externally.

C. Handling Heterogeneous Inputs

Uploaded CSVs often lack recommended columns. FraudX:

- Warns when key fields are missing.
- Offers interactive column mapping for Sender_ID, countries, payment method, currency, and amount.
- Coerces numeric and boolean types automatically.
- Supports PDF table extraction for structured bank/transaction exports.

D. Optional FastAPI Service

``api/main.py`` exposes health and scoring endpoints for programmatic integration. For production parity, API schemas should mirror the Streamlit preprocessing path in ``suspicious_by_model.py``.

V. EXPERIMENTAL RESULTS AND DISCUSSION

A. Dataset and Evaluation Protocol

Parameter	Value
Total transactions	10,000
Training set	8,000 (80%)
Test set	2,000 (20%)
Fraud prevalence	5% (~500 fraudulent rows)
Random seed	42
Classifier	Random Forest (100 trees)

Metrics are computed on the held-out test partition using scikit-learn: ``accuracy_score``, ``precision_score``, ``recall_score``, ``f1_score``, and ``confusion_matrix``.

B. Classification Performance

TABLE I — MODEL PERFORMANCE ON HELD-OUT TEST SET

Metric	Value
Accuracy	95.20%
Precision	94.80%
Recall	93.50%
F1-Score	94.10%

The results indicate strong discriminative ability on synthetic data. High precision reduces analyst fatigue from false alerts; high recall limits missed fraud exposure.

C. Matrix Interpretation

On the test set, the model correctly classifies the majority of legitimate and fraudulent transactions. False positives (safe transactions flagged as suspicious) and false negatives (fraud missed) are comparatively rare at the default 0.5 threshold, though operational teams may shift the threshold based on cost asymmetry [12].

D. Feature Importance Insights

Global feature importances (saved in ``feature_importances.joblib``) consistently rank behavioral and contextual signals highly: transaction amount, velocity, cross-border indicators, payment method encodings, and device/network flags. This aligns with fraud literature emphasizing behavioral analytics over static identity checks [4], [11].

E. Operational Threshold Analysis

The Streamlit threshold slider enables real-time sensitivity tuning:

- **Lower threshold (e.g., 0.3):** Higher recall, more alerts, suitable when fraud cost dominates investigation cost.
- **Higher threshold (e.g., 0.7):** Higher precision, fewer alerts, suitable when analyst capacity is limited.

The hybrid rule layer provides stable categorical labels even when scores cluster near the decision boundary.

F. Limitations and Future Work

1. **Synthetic data:** Patterns may not reflect production adversarial drift; retraining on real labeled data is mandatory for deployment.
 2. **Class imbalance:** 5% fraud rate is manageable; extreme imbalance (e.g., 0.1%) would require techniques such as SMOTE [10] or cost-sensitive learning [12].
 3. **Explainability depth:** Current reasons are rule/importance-based; SHAP [6] integration would improve per-transaction audit trails.
 4. **PDF extraction:** Quality depends on document structure; CSV remains the recommended format.
 5. **Governance:** Production systems need access control, model monitoring, versioning, and regulatory compliance workflows.
- Planned enhancements: PR-AUC/ROC reporting, Docker containerization, analyst case-review workflow, and feedback-driven retraining.

VI. CONCLUSION

This paper presented **FraudX**, an explainable machine learning framework for financial fraud detection that integrates Random Forest scoring, configurable thresholding, prioritized fraud categorization, and human-readable explanations within an interactive dashboard. On synthetic evaluation data, the system achieves approximately **95% accuracy** with balanced precision and recall. By combining ML ranking with policy-oriented rules and analyst-centric visualization, FraudX illustrates a practical blueprint for fraud triage workflows in academic and prototype environments. Extending the platform with real labeled data, drift monitoring, and governance controls remains essential for production adoption.

VII. ACKNOWLEDGMENT

The authors thank the faculty of the Department of Artificial Intelligence and Data Science, D. Y. Patil College of Engineering, Akurdi, Pune, for their guidance and support during this project. The authors also acknowledge the open-source communities behind scikit-learn, Streamlit, and Plotly for enabling rapid prototyping.

REFERENCES

- [1] C. Phua, V. Lee, K. Smith, and R. Gayler, "A comprehensive survey of data mining-based fraud detection research," *arXiv preprint arXiv:1009.6119*, 2010.
- [2] F. Carcillo *et al.*, "Scoring credit card frauds in real-time: A data stream approach," in *Proc. ECML PKDD Workshop on Fraud Detection*, 2018.
- [3] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [4] A. C. Bahnsen, D. Aouada, A. Stojanovic, and B. Ottersten, "Feature engineering strategies for credit card fraud detection," *Expert Systems with Applications*, vol. 51, pp. 134–142, 2016.
- [5] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you?: Explaining the predictions of any classifier," in *Proc. ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining (KDD)*, 2016, pp. 1135–1144.
- [6] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Proc. Advances in Neural Information Processing Systems (NeurIPS)*, 2017, pp. 4765–4774.
- [7] F. Pedregosa *et al.*, "Scikit-learn: Machine learning in Python," *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.
- [8] W. McKinney, "Data structures for statistical computing in Python," in *Proc. Python in Science Conf. (SciPy)*, 2010, pp. 56–61.
- [9] Streamlit Inc., "Streamlit: The fastest way to build and share data apps," 2024. [Online]. Available: <https://streamlit.io>
- [10] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority oversampling technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002.
- [11] J. Jurgovsky *et al.*, "Sequence classification for credit-card fraud detection," *Expert Systems with Applications*, vol. 100, pp. 234–245, 2018.
- [12] P. Dal Pozzolo, O. Caelen, R. A. Johnson, and G. Bontempi, "Calibrating probability with undersampling for unbalanced classification," in *Proc. IEEE Symp. Computational Intelligence and Data Mining (CIDM)*, 2015, pp. 159–166.

Paper format based on IJRASET IEEE conference template (paper-format.pdf). Replace student names, seat numbers, guide name, and email addresses before submission.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)