



“An Imbalance-Aware Machine Learning Framework for Credit Card Fraud Detection”

Author1: [Mamta Indrasing Girase]

Department: [Electronics & Telecommunication]

Institute: [Gangamai College of Engineering, Nagaon, Ms, Dhule]

Email: [mamtagirase1210@gmail.com]

Author1: [Jagruti Sarang Patil]

Department: [Electronics & Telecommunication]

Abstract

The rapid growth of digital payment systems and online financial transactions has increased the exposure of financial institutions and customers to credit card fraud. Traditional rule-based fraud detection systems are often unable to adapt effectively to evolving fraudulent behavior and may generate a high number of false alarms. Machine learning (ML) provides a data-driven alternative by learning transaction patterns and identifying suspicious activities automatically.

This paper presents an imbalance-aware machine learning framework for credit card fraud detection using Logistic Regression, Decision Tree, Random Forest, and Extreme Gradient Boosting (XGBoost). The study is based on the widely used Credit Card Fraud Detection dataset containing 284,807 transactions, of which 492 are fraudulent, resulting in an extreme class imbalance of approximately 0.172%. The dataset contains anonymized PCA-transformed variables V1–V28 along with Time, Amount, and the binary Class target variable.

The proposed framework consists of data preprocessing, feature scaling, stratified train-test splitting, class-imbalance handling, model training, hyperparameter optimization, and evaluation using confusion matrix, precision, recall, F1-score, ROC-AUC, and Precision-Recall AUC. Particular emphasis is placed on precision and recall because conventional accuracy can be misleading for highly imbalanced fraud datasets. The framework also supports real-time deployment in which a transaction is processed, transformed into model-ready features, assigned a fraud probability, and forwarded to an appropriate decision layer.

The study provides a systematic methodology for comparing classical and ensemble machine learning algorithms and establishes a foundation for developing adaptive, scalable, and interpretable fraud detection systems.

Keywords: Credit Card Fraud Detection, Machine Learning, Class Imbalance, Random Forest, XGBoost, Logistic Regression, SMOTE, Fraud Detection, Precision, Recall, Explainable AI.

1. INTRODUCTION

The rapid adoption of digital payment technologies has transformed the way individuals and organizations conduct financial transactions. Credit cards are extensively used for online purchases, retail transactions, subscriptions, and international payments. Although digital payments provide convenience and speed, they also create opportunities for unauthorized transactions, identity theft, account takeover, card-not-present fraud, and other fraudulent activities.

The project material used as the basis for this research identifies credit card fraud as a significant cybersecurity and financial problem. Fraudulent transactions represent only a small fraction of total transactions, but their financial and operational consequences can be substantial. Traditional systems generally use manually defined rules such as transaction amount thresholds, unusual transaction timings, or location mismatches. However, static rules may become ineffective when fraudsters change their behavior.

Machine learning offers an adaptive approach in which algorithms learn patterns from historical transaction data. Classification models can estimate whether a transaction is legitimate or fraudulent, while anomaly-detection techniques can identify unusual transactions. Recent research also emphasizes the importance of addressing class imbalance, explainability, evolving fraud patterns, and real-time processing. A systematic review published in the *Journal of Big Data* highlights the growing use of artificial intelligence techniques for improving fraud detection systems.

The major difficulty in credit card fraud detection is that fraudulent transactions are extremely rare compared with legitimate transactions. The dataset considered in this study contains 284,807 transactions but only 492 fraud cases, corresponding to approximately 0.172% fraud.

Consequently, a model can obtain a very high overall accuracy while still failing to detect a meaningful number of fraudulent transactions. Therefore, fraud detection research must consider metrics such as precision, recall, F1-score, ROC-AUC, and particularly Precision-Recall AUC rather than relying on accuracy alone.

This paper develops an imbalance-aware comparative framework using Logistic Regression, Decision Tree, Random Forest, and XGBoost. The objective is not merely to maximize classification accuracy but to establish a practical methodology for detecting fraudulent transactions while controlling false positives.

2. PROBLEM STATEMENT

Credit card fraud detection is a challenging classification problem because fraudulent transactions are rare, fraud patterns evolve over time, and false decisions have different financial consequences.

The principal problems addressed in this research are:

1. Extremely imbalanced transaction classes.
2. Rapidly evolving fraudulent behavior.
3. High cost of false-negative predictions.
4. Customer dissatisfaction caused by excessive false-positive predictions.
5. Large-scale transaction processing requirements.
6. Limited interpretability of anonymized PCA-transformed features.
7. Requirement for rapid or real-time fraud classification.

3. RESEARCH OBJECTIVES

The main objectives of the proposed research are:

1. To develop a machine learning framework for detecting fraudulent credit card transactions.
2. To compare Logistic Regression, Decision Tree, Random Forest, and XGBoost algorithms.
3. To address the extreme class imbalance using suitable resampling and/or cost-sensitive learning techniques.

4. LITERATURE REVIEW

A. Credit Card Fraud

Credit card fraud can occur through several mechanisms, including card-present fraud, card-not-present fraud, account takeover, lost or stolen cards, identity theft, skimming, phishing, and application fraud. The supplied project report provides a detailed classification of these fraud types.

The diversity of fraud techniques makes static rule-based detection difficult. A rule designed to identify one fraud pattern may not detect another pattern. Moreover, excessive rules can increase false-positive rates and negatively affect genuine customers.

B. Traditional Fraud Detection

Traditional fraud detection systems commonly use:

- Transaction amount thresholds
- Geographic inconsistencies
- Unusual transaction timing
- Suspicious login behavior
- Manually defined business rules

Although such methods can be effective for known patterns, they are less flexible when fraudsters change their strategies. The supplied report specifically identifies static rules, high false positives, and inability to detect unseen fraud patterns as limitations.

C. Machine Learning-Based Fraud Detection

Machine learning models can learn statistical relationships from historical transaction data. Commonly investigated models include Logistic Regression, Decision Trees, Random Forest, Gradient Boosting, and neural networks.

Recent literature has expanded toward deep learning, graph-based learning, hybrid approaches, and explainable AI. A 2025 systematic review in the *Journal of Big Data* surveys AI-enhanced approaches for financial fraud detection and highlights the importance of advanced AI techniques in addressing evolving fraud behavior.

Recent work has also investigated graph neural networks for representing relationships among transactions and entities. For example, a 2025 *Scientific Reports* study proposed an adaptive graph-based approach for credit card fraud detection.

D. Class Imbalance

Class imbalance is one of the most important challenges in fraud detection. In the selected dataset, only 492 of 284,807 transactions are fraudulent.

Common techniques include:

- Random undersampling
- Random oversampling
- SMOTE
- ADASYN

E. Research Gap

Existing research demonstrates strong potential for machine learning-based fraud detection, but several challenges remain:

1. Extreme class imbalance can distort evaluation.
2. Accuracy can provide misleading conclusions.
3. Different resampling strategies may produce different outcomes.
4. Fraud patterns can change over time.
5. Anonymized PCA features limit business interpretation.
6. Many published studies use different train-test strategies, making direct comparison difficult.
7. Real-time deployment introduces latency and scalability constraints.
8. Explainability remains important for financial decision-making.

Therefore, a controlled comparative framework that evaluates multiple algorithms using consistent preprocessing and fraud-sensitive metrics is valuable.

5. DATASET DESCRIPTION

The research uses the publicly available Credit Card Fraud Detection dataset containing transactions made by European cardholders in September 2013. The dataset contains transactions from a two-day period and includes 284,807 records, of which 492 are fraudulent. The fraud class represents approximately 0.172% of all transactions.

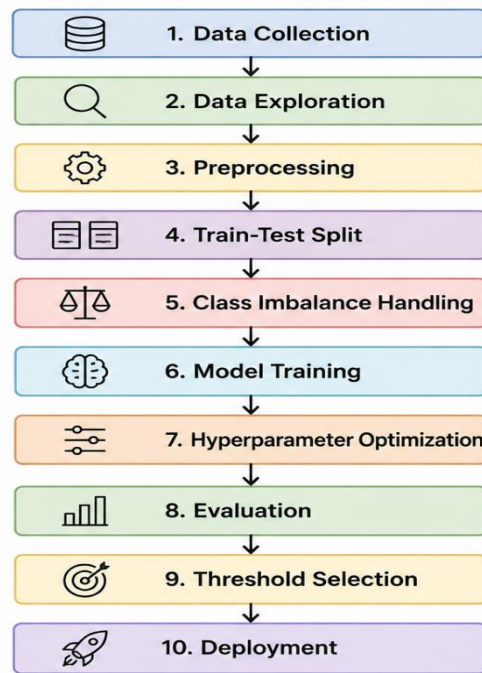
Table I. Dataset Characteristics

Parameter	Description
Total transactions	284,807
Fraudulent transactions	492
Legitimate transactions	284,315
Fraud ratio	Approximately 0.172%
Number of features	31
Target variable	Class
Legitimate class	Class = 0
Fraud class	Class = 1
PCA features	V1–V28
Non-PCA features	Time, Amount

The dataset contains anonymized numerical variables. Features V1 through V28 are PCA-transformed variables, while Time and Amount retain their original meanings. The Class attribute represents the target variable.

6. PROPOSED METHODOLOGY

The proposed research framework consists of the following stages:



A. Data Preprocessing

The supplied project report states that the dataset contains no missing values. The proposed preprocessing pipeline therefore includes:

1. Duplicate checking.
2. Missing-value verification.
3. Class-distribution analysis.
4. Feature scaling.
5. Stratified train-test splitting.

The Amount and Time attributes should be scaled because their numerical ranges differ from the PCA-transformed variables. StandardScaler or MinMaxScaler can be considered.

B. Train-Test Splitting

A stratified split should be used to preserve the fraud-to-legitimate ratio between training and testing data.

For example:

- Training set: 80%
- Testing set: 20%

The test set should remain representative of the original imbalanced distribution.

An important methodological rule is that oversampling techniques such as SMOTE should be applied **only to the training data**, never before splitting the dataset. This prevents information leakage from synthetic training samples into the evaluation set.

. CLASS IMBALANCE HANDLING

A. Random Undersampling

Random undersampling reduces the number of legitimate transactions.

Advantage

It is computationally efficient.

Limitation

Potentially useful legitimate transaction information may be discarded.

B. SMOTE

The Synthetic Minority Oversampling Technique generates synthetic minority-class observations.

Conceptually, for a minority observation (x_i) and one of its nearest minority neighbors (x_j), a synthetic observation can be represented as:

$$[x_{\text{new}} = x_i + (x_j - x_i)]$$

where:

$$[0]$$

SMOTE can improve minority representation, but synthetic observations can also introduce noise if the underlying minority distribution is complex.

C. ADASYN

ADASYN generates more synthetic examples in regions where minority samples are more difficult to classify.

D. Class Weighting

Instead of modifying the training dataset, the learning algorithm can assign a higher penalty to errors involving the minority fraud class.

This approach is particularly attractive because it preserves the original transaction distribution.

E. Anomaly Detection

Fraud can also be treated as an anomaly-detection problem. Candidate techniques include:

- Isolation Forest
- One-Class SVM
- Local Outlier Factor

These approaches are discussed in the supplied report.

7. MACHINE LEARNING MODELS

A. Logistic Regression

Logistic Regression provides a strong interpretable baseline.

The probability of fraud can be expressed as:

$$[P(y=1|x)=]$$

where:

$$[z=_0+_1x_1+_2x_2++_nx_n]$$

The model is simple, computationally efficient, and useful for establishing a baseline.

B. Decision Tree

A Decision Tree recursively divides the feature space using decision rules.

Advantages include:

- Easy interpretation
- Ability to model nonlinear relationships
- No requirement for feature linearity

However, deep trees can overfit the training data.

C. Random Forest

Random Forest is an ensemble model consisting of multiple decision trees.

For a classification problem, multiple trees produce predictions and the final class is determined using an aggregation mechanism.

The supplied report identifies Random Forest as an effective model for tabular data and notes that balanced class weights can help address class imbalance.

D. XGBoost

Extreme Gradient Boosting builds an ensemble of weak learners sequentially, with subsequent learners focusing on previous errors.

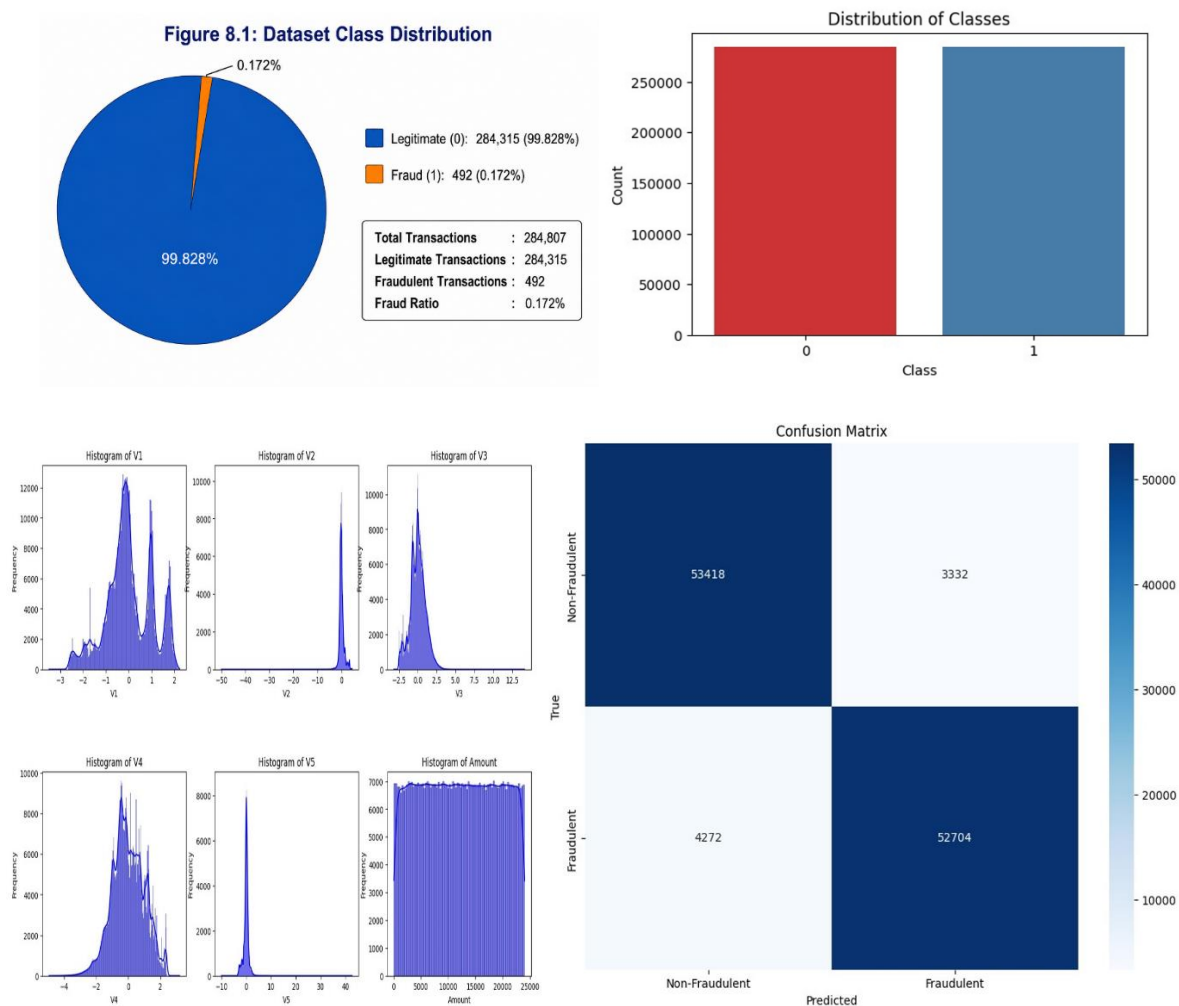
A simplified objective can be represented as:

$$[Obj=_i L(y_i,)+_k(f_k)]$$

where (L) represents the loss function and () represents model complexity regularization.

XGBoost is suitable for tabular classification and provides parameters that can be adjusted for imbalanced classification.

8. RESULTS



9. PROPOSED MODEL COMPARISON

Table II. Comparison of Algorithms

Algorithm	Main Advantage	Major Limitation	Expected Role
Logistic Regression	Simple and interpretable	Limited nonlinear modeling	Baseline
Decision Tree	Easy to understand	Overfitting	Interpretable model
Random Forest	Robust ensemble	Higher computation	Strong candidate
XGBoost	Powerful boosting model	Requires tuning	High-performance candidate

The comparison is based on the characteristics of the models described in the supplied research material.

10. EXPERIMENTAL EVALUATION FRAMEWORK

The supplied files describe the methodology but do not contain experimentally measured model results. Therefore, numerical performance values should be generated only after implementing and executing the proposed models.

The final experimental table should follow the structure below.

Table III. Experimental Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC	PR-AUC
Logistic Regression	To be measured	To be measured	To be measured	To be measured	To be measured	To be measured
Decision Tree	To be measured	To be measured	To be measured	To be measured	To be measured	To be measured
Random Forest	To be measured	To be measured	To be measured	To be measured	To be measured	To be measured
XGBoost	To be measured	To be measured	To be measured	To be measured	To be measured	To be measured

11. ADVANTAGES OF THE PROPOSED SYSTEM

Financial Advantages

- Potential reduction in fraud losses.
- Faster identification of suspicious transactions.
- Reduced manual investigation effort.
- Improved transaction monitoring.

Technical Advantages

- Data-driven decision-making.
- Ability to model nonlinear relationships.
- Support for large-scale transaction processing.
- Potential for real-time prediction.
- Flexible model selection.

Customer Advantages

- Improved transaction security.
- Potential reduction in unauthorized transactions.
- Reduced unnecessary transaction declines when thresholds are properly optimized.
- Increased confidence in digital payment systems.

The supplied report similarly identifies financial, technical, and customer-experience benefits.

12. FUTURE SCOPE

The supplied research material identifies several important future directions.

A. Deep Learning

Future research can investigate:

- Recurrent Neural Networks
- LSTM
- Autoencoders
- Transformer-based models

LSTM models can be useful when transaction sequences contain temporal information.

B. Explainable AI

SHAP and LIME can be investigated to explain why a transaction was assigned a high fraud probability.

This is especially valuable in financial environments where analysts need to understand model decisions.

C. Real-Time Cloud Deployment

The system can be deployed using cloud or edge-computing infrastructure to support rapid fraud decisions.

D. Hybrid Fraud Detection

A hybrid architecture combining rule-based systems with machine learning can provide complementary detection capabilities.

E. Graph Neural Networks

Transactions can be represented as relationships between cards, accounts, merchants, devices, and locations. Graph-based models may identify suspicious relationships that conventional tabular models cannot easily capture.

F. Continual Learning

Future systems should periodically learn from newly confirmed fraud cases while controlling model drift.

G. Cost-Sensitive Learning

Instead of treating every classification error equally, future models can assign different costs to false positives and false negatives.

The supplied dissertation specifically proposes deep learning, real-time deployment, Explainable AI, hybrid models, and blockchain integration as future directions.

13. CONCLUSION

Credit card fraud is a major challenge in modern digital financial systems because fraudulent transactions are rare, fraud patterns continuously evolve, and incorrect predictions can have significant financial and customer-experience consequences.

This research proposes an imbalance-aware machine learning framework for credit card fraud detection using Logistic Regression, Decision Tree, Random Forest, and XGBoost. The framework incorporates preprocessing, stratified data splitting, imbalance handling, hyperparameter optimization, threshold selection, and fraud-sensitive evaluation metrics.

The selected dataset contains 284,807 transactions and only 492 fraudulent transactions, making it an appropriate benchmark for studying extreme class imbalance.

The study emphasizes that accuracy alone is insufficient for evaluating fraud detection systems. Precision, recall, F1-score, ROC-AUC, and Precision-Recall AUC provide a more meaningful assessment of minority-class performance.

REFERENCES

- [1] A. Dal Pozzolo, O. Caelen, R. A. Johnson, and G. Bontempi, “Calibrating Probability with Undersampling for Unbalanced Classification,” in *2015 IEEE Symposium Series on Computational Intelligence*, 2015. The dataset documentation identifies this work as a recommended citation for the benchmark dataset.
- [2] M. L. G. — ULB, “Credit Card Fraud Detection Dataset,” Kaggle. The dataset contains 284,807 transactions and 492 fraudulent transactions from European cardholders.
- [3] I. Y. Hafez, A. Y. Hafez, A. Saleh, A. A. Abd El-Mageed, et al., “A systematic review of AI-enhanced techniques in credit card fraud detection,” *Journal of Big Data*, vol. 12, article 6, 2025.
- [4] S. Verma and J. Dhar, “Credit Card Fraud Detection: A Deep Learning Approach,” arXiv, 2024.
- [5] Y. Wang, “A Data Balancing and Ensemble Learning Approach for Credit Card Fraud Detection,” arXiv, 2025.
- [6] L. Ni, X. Li, Y. Zhou, H. Qi, X. Man, et al., “HMOA-GNN: adaptive adversarial GraphSAGE with hierarchical hybrid sampling and metric-optimized graph construction for credit card fraud detection,” *Scientific Reports*, vol. 15, article 43005, 2025.
- [7] P. Sundaravadivel, R. Augustian Isaac, D. Elangovan, D. KrishnaRaj, V. V. Lokesh Rahul, et al., “Optimizing credit card fraud detection with random forests and SMOTE,” *Scientific Reports*, 2025. **Note: this article was subsequently retracted on December 22, 2025 and therefore should not be used as a supporting research claim.**
- [8] Mienye, I. D., and Jere, N., “Deep learning for credit card fraud detection: A review of algorithms, challenges, and solutions,” *IEEE Access*, 2024.
- [9] Sarker, A., Yasmin, M. A., Rahman, M. A., Rashid, M. H. O., and Roy, B. R., “Credit Card Fraud Detection Using Machine Learning Techniques,” *Journal of Computer and Communications*, vol. 12, 2024.