

Credit Card Fraud Detection Using Machine Learning Techniques

¹ Jabulani Khumalo, ² Kim Min Joon

¹ University of South Africa, Pretoria, South Africa

² Pohang University of Science and Technology (POSTECH), Pohang, South Korea

Corresponding E-mail: jabulani.khumalo@interviauniversity.com

Abstract

Credit card fraud is a serious issue in the financial sector and can cause financial losses for customers, banks, and payment service providers. As digital transactions continue to grow, it has become increasingly important to detect fraudulent activity quickly and accurately. Traditional fraud detection systems often rely on predefined rules, which can make it difficult to recognize new or changing patterns of fraud. Machine learning offers another approach by learning patterns from previous transaction data and using those patterns to classify new transactions as legitimate or fraudulent.

This research compares four machine learning techniques for credit card fraud detection: Logistic Regression, Random Forest, Support Vector Machine, and XGBoost. The proposed process includes data preprocessing, handling class imbalance, preparing features, training the models, and evaluating their performance. Accuracy, precision, recall, F1-score, and ROC-AUC are considered as evaluation measures. Particular attention is given to precision and recall because fraud datasets are usually highly imbalanced, with fraudulent transactions making up only a small part of all transactions. Overall, simple models such as Logistic Regression can provide an understandable baseline, while ensemble methods such as Random Forest and XGBoost can capture more complicated transaction patterns. The study shows that machine learning can be useful for fraud detection when it is combined with suitable preprocessing, class-imbalance techniques, and careful evaluation.

Keywords: Credit Card Fraud, Machine Learning, Fraud Detection, Logistic Regression, Random Forest, SVM, XGBoost, Classification

1. Introduction

The rapid growth of online banking, e-commerce, mobile payments, and other digital financial services has made credit cards a common way to make payments. People can now

purchase goods and transfer money quickly without having to visit a bank or financial institution. While this convenience has many benefits, it has also created more opportunities for fraudulent activity. Credit card fraud occurs when someone uses another person's card or payment information without permission to make transactions.

Fraudulent transactions can result in financial losses and can create problems for both customers and financial institutions [1]. Banks and payment providers therefore need to monitor transactions continuously and identify suspicious activity. This is challenging because millions of legitimate transactions may take place every day, while fraudulent transactions usually make up only a small percentage of the total. A good fraud detection system needs to identify suspicious transactions without incorrectly blocking too many genuine customers.

Traditional fraud detection systems often depend on predefined rules. For instance, a transaction might be flagged if it happens in an unusual location, involves an unusually large amount, or occurs shortly after another transaction from a different location. Although rule-based systems can be helpful, they may not adapt easily when fraudsters change their methods [2]. This creates a need for more flexible approaches that can recognize patterns that are not covered by existing rules.

Machine learning provides one possible solution because it can learn patterns from historical transaction data. Instead of depending completely on manually created rules, a machine learning model can examine characteristics of previous transactions and learn the differences between legitimate and fraudulent behavior. Once trained, the model can be used to estimate whether a new transaction is likely to be fraudulent [3].

Different machine learning algorithms can be used for this task. Logistic Regression is relatively simple and easy to interpret, making it useful as a baseline. Support Vector Machine can identify more complex boundaries between legitimate and fraudulent transactions. Random Forest combines several decision trees and can capture nonlinear relationships between transaction features. XGBoost uses gradient boosting to create an ensemble model and can work effectively with structured financial data[4].

The main objective of this research is to analyze and compare these machine learning techniques for credit card fraud detection. The study focuses on preprocessing, class imbalance, model training, evaluation measures, and the strengths and limitations of the selected algorithms. It also highlights the importance of precision and recall because overall accuracy can be misleading when fraudulent transactions represent only a small part of the dataset.

2. Methodology

This research uses a supervised machine learning classification approach. The goal is to classify individual credit card transactions into two categories: legitimate and fraudulent.

The proposed methodology includes the following stages:

1. Dataset collection
2. Data preprocessing
3. Handling class imbalance
4. Feature preparation
5. Model training
6. Model evaluation
7. Comparative analysis

The four algorithms considered in this study are Logistic Regression, Support Vector Machine, Random Forest, and XGBoost.

2.1 Research Approach

The research follows a supervised classification approach in which the model learns from previously labelled transactions [5]. The aim is to use the available transaction features to determine whether a new transaction should be classified as legitimate or fraudulent.

2.2 Dataset

A publicly available credit card transaction dataset can be used to evaluate the selected models [6]. A commonly used fraud detection dataset contains a large number of transactions, while only a small proportion of them are labelled as fraudulent.

The available features may include transaction amount, time, and transformed numerical variables that represent characteristics of the transactions. The target variable identifies whether a transaction is legitimate or fraudulent.

A major challenge with this type of dataset is class imbalance. Legitimate transactions are normally far more common than fraudulent ones. Because of this, a model could achieve high overall accuracy simply by predicting most transactions as legitimate while failing to detect actual fraud[7].

For example, when fraudulent transactions form only a small percentage of the data, a model that predicts almost everything as legitimate may still appear highly accurate. This is why class imbalance needs to be considered during both model development and evaluation.

2.3 Data Preprocessing

Data preprocessing is an important step because the quality of the input data can affect the final performance of the models. The dataset should first be checked for missing values, duplicate records, unusual observations, and inconsistent information.

Transaction amounts may have a different numerical range from other features. Scaling or normalization may therefore be needed, especially for algorithms such as Logistic Regression and SVM. Putting the features on comparable scales can help these algorithms work more effectively[8].

The dataset should then be separated into training and testing sets. The training data is used to build the models, while the testing data is kept separate so that performance can be measured on transactions the models have not seen before.

A stratified split can be useful because it keeps a similar proportion of legitimate and fraudulent transactions in both subsets. Cross-validation can also be used during model development to obtain a more reliable estimate of performance [9].

2.4 Handling Class Imbalance

Class imbalance is one of the main challenges in credit card fraud detection. When fraudulent transactions are much less common than legitimate transactions, a model can become biased toward the majority class.

One way to address this problem is oversampling, where additional examples of the minority class are created or existing minority observations are replicated. SMOTE is one commonly used technique that generates synthetic examples of the minority class.

Another option is undersampling, which reduces the number of observations from the majority class. This can make the classes more balanced, but it may also remove useful information from legitimate transactions[10].

Class weighting is another approach. Instead of changing the dataset, the algorithm gives more importance to errors involving fraudulent transactions. This encourages the model to pay greater attention to the minority class.

Whichever approach is selected, imbalance handling should be performed only on the training data. Applying oversampling before splitting the dataset can cause information leakage and may result in unrealistically strong evaluation results.

2.5 Logistic Regression

Logistic Regression is a commonly used classification algorithm that estimates the probability of an observation belonging to a particular class. In this case, it can be used to estimate the probability that a credit card transaction is fraudulent [11].

One of the main advantages of Logistic Regression is its simplicity and interpretability. It provides a useful baseline for comparing more complex models and generally requires relatively few computational resources, which can be helpful when working with large transaction datasets[12].

Its main limitation is that fraudulent behavior can involve complicated relationships among several transaction characteristics. Logistic Regression may not capture these nonlinear patterns effectively unless additional feature engineering is carried out.

2.6 Support Vector Machine

Support Vector Machine is a supervised learning algorithm that tries to find a suitable boundary between different classes. For fraud detection, the aim is to separate legitimate transactions from fraudulent ones.

SVM can use kernel functions to model nonlinear relationships. This can be useful when fraudulent and legitimate transactions cannot be separated by a simple linear boundary.

However, SVM can require considerable computational resources when working with very large datasets. Its performance is also influenced by parameter selection and feature scaling, so careful preprocessing and tuning are important[13].

2.7 Random Forest

Random Forest is an ensemble learning method that combines multiple decision trees. Each tree learns from different subsets of the data and features, and their predictions are combined to produce the final result.

One of the main strengths of Random Forest is its ability to capture nonlinear relationships and interactions between transaction features. This can be useful because fraudulent behavior does not always follow a simple pattern.

Random Forest is also relatively robust because the final prediction is based on many trees rather than a single tree. In addition, it can provide information about feature importance. However, larger Random Forest models can become more difficult to interpret than simpler models such as Logistic Regression.

2.8 XGBoost

XGBoost is a gradient-boosting algorithm that builds decision trees one after another. Each new tree tries to improve on the errors made by the previous trees. In this way, the model gradually becomes better at classification.

XGBoost is well suited to structured datasets and can learn complex relationships between transaction features. It also includes parameters that can be adjusted to control the complexity of the model and help reduce overfitting[14, 15].

The main challenge is that XGBoost can require careful parameter tuning. If it is not configured properly, the model may learn patterns that work well on the training data but do not generalize well to new transactions[16].

3. Results and Discussion

The comparison of Logistic Regression, SVM, Random Forest, and XGBoost shows that each algorithm approaches fraud detection differently.

Logistic Regression is a useful baseline because it is simple, relatively fast, and easier to interpret. It can identify general relationships between transaction features and the likelihood of fraud. However, its simpler decision boundary may limit its ability to recognize complicated fraud patterns.

SVM provides more flexibility by allowing the model to create complex classification boundaries. This can be helpful when legitimate and fraudulent transactions have similar characteristics. The downside is that SVM can require more computational resources, especially with large datasets[8].

Random Forest combines multiple decision trees and can capture nonlinear relationships and interactions between features. It is generally more flexible than Logistic Regression and can also provide feature-importance information. At the same time, its greater complexity can make individual predictions harder to explain.

XGBoost uses sequential decision-tree learning and can model complex patterns in structured transaction data. Because each new tree focuses on previous errors, the model can gradually improve its predictions. However, it needs careful tuning and validation.

Fraud detection models should not be judged by accuracy alone. Since fraudulent transactions are usually much less common than legitimate transactions, accuracy can give a misleading picture of how well a model is working.

Precision measures how many of the transactions predicted as fraudulent are actually fraudulent. High precision means that the system produces fewer false alarms.

Recall, or sensitivity, measures how many of the actual fraudulent transactions are detected. High recall is important because failing to identify fraud can result in financial losses.

F1-score combines precision and recall into a single measure and can be useful when both false positives and false negatives matter. ROC-AUC can also be used to assess how well the model separates fraudulent transactions from legitimate ones across different classification thresholds.

The comparison suggests that ensemble methods such as Random Forest and XGBoost can be useful when transaction patterns are complicated. However, a more complex model is not automatically the best choice. Financial institutions also need to consider interpretability, processing speed, false-alarm rates, and practical operating requirements.

Another important issue is the balance between detecting fraud and avoiding false alarms. A system that flags too many legitimate transactions can inconvenience customers and create unnecessary work for investigators. On the other hand, a system that is too conservative may allow fraudulent transactions to go through. The decision threshold should therefore match the practical needs of the organization.

Real-time performance is also important because credit card transactions happen continuously and may need to be assessed within a very short time. A model that performs well but takes too long to make a prediction may not be practical for real-world use. Both predictive performance and computational efficiency should therefore be considered.

The quality of the training data also affects the reliability of the system. Fraudsters can change their methods over time, so a model trained on older transactions may become less effective as new patterns appear. Regular monitoring, retraining, and validation can help maintain performance.

Privacy and security are also important when working with financial transaction data. Such data can contain sensitive information, so appropriate safeguards should be used during collection, storage, processing, and model development.

Overall, machine learning can improve fraud detection by identifying patterns that may be difficult to capture with traditional rule-based systems. However, successful implementation depends on good preprocessing, appropriate handling of class imbalance, reliable evaluation, and continuous monitoring.

4. Conclusion

This research examined the use of machine learning techniques for credit card fraud detection and compared four algorithms: Logistic Regression, Support Vector Machine, Random Forest, and XGBoost. The study discussed the importance of preprocessing, handling class imbalance, preparing features, training models, and using suitable evaluation measures.

Logistic Regression provides a simple and interpretable baseline, while SVM offers more flexible classification. Random Forest can model nonlinear relationships by combining multiple decision trees, while XGBoost provides a powerful boosting-based approach for complex structured transaction data. The study also shows why accuracy alone is not a suitable measure for fraud detection. Since fraudulent transactions normally represent only a small part of all transactions, precision, recall, F1-score, and ROC-AUC give a more useful picture of model performance. Ensemble methods such as Random Forest and XGBoost have strong potential for identifying complicated fraud patterns. However, the final choice of model should depend not only on predictive performance but also on computational efficiency, interpretability, false-positive rates, and real-time requirements. Machine learning can therefore be an effective part of modern credit card fraud detection systems. Future improvements could include real-time learning, adaptive models, larger and more diverse transaction datasets, and advanced anomaly-detection methods. With regular monitoring and proper validation, machine learning can help financial institutions identify suspicious transactions more effectively, reduce financial losses, and improve the security of digital payments.

References:

- [1] J. Barach, "Enhancing Ransomware Resilience in Cloud-Based HR Systems through Moving Target Defense," *Computers, Materials and Continua*, vol. 86, no. 2, pp. 1-23, 2025, doi: <https://doi.org/10.32604/cmc.2025.071705>.
- [2] A. Mohammed, "Artificial Intelligence-Driven Cybersecurity Crimes: Emerging Threats, AI Bots, and Advanced Defense Strategies," *Euro Vantage journals of Artificial intelligence*, vol. 3, no. 1, pp. 46-65, 2026, doi: <https://doi.org/10.65923/c2srx21>.
- [3] D. Bodra and S. Khairnar, "Comparative performance analysis of modern NoSQL data technologies: Redis, Aerospike, and Dragonfly," *arXiv preprint arXiv:2510.08863*, 2025.
- [4] T. Vafeiadis, K. I. Diamantaras, G. Sarigiannidis, and K. C. Chatzisavvas, "A comparison of machine learning techniques for customer churn prediction," *Simulation Modelling Practice and Theory*, vol. 55, pp. 1-9, 2015.
- [5] A. Qatawneh, "The influence of data mining on accounting information system performance: a mediating role of information technology infrastructure," *Journal of Governance and*

-
- Regulation/Volume*, vol. 11, no. 1, pp. 141-151, 2022, doi: <https://doi.org/10.22495/jgrv11i1art13>.
- [6] A. Mohammed, "Transforming SOC Operations: Harnessing the Power of AI and ML for Enhanced Threat Detection," *INTERNATIONAL JOURNAL OF RESEARCH CULTURE SOCIETY Monthly Peer-Reviewed, Refereed, Indexed*, vol. 8, pp. 11-21, 2024, doi: 10.2017/IJRCS/ICCDMP01.
- [7] S. Malik, S. Harous, and H. El-Sayed, "Comparative analysis of machine learning algorithms for early prediction of diabetes mellitus in women," in *International Symposium on Modelling and Implementation of Complex Systems*, 2020: Springer, pp. 95-106.
- [8] Q. Truong, M. Nguyen, H. Dang, and B. Mei, "Housing price prediction via improved machine learning techniques," *Procedia Computer Science*, vol. 174, pp. 433-442, 2020.
- [9] A. M. Qatawneh, "The effect of electronic commerce on the accounting information system of Jordanian banks," *International Business Research*, vol. 5, no. 5, p. 158, 2012, doi: 10.5539/ibr.v5n5p158.
- [10] E. Ahmed, "Student performance prediction using machine learning algorithms," *Applied computational intelligence and soft computing*, vol. 2024, no. 1, p. 4067721, 2024.
- [11] J. Barach, "Federated Learning for Privacy-Preserving Employee Performance Analytics," *IEEE Access*, vol. 13, 2025, doi: 10.1109/ACCESS.2025.3591360.
- [12] S. Ara, A. Das, and A. Dey, "Malignant and benign breast cancer classification using machine learning algorithms," in *2021 international conference on artificial intelligence (ICAI)*, 2021: IEEE, pp. 97-101.
- [13] M. Amrane, S. Oukid, I. Gagaoua, and T. Ensari, "Breast cancer classification using machine learning," in *2018 electric electronics, computer science, biomedical engineerings' meeting (EBBT)*, 2018: IEEE, pp. 1-4.
- [14] M. Beseiso, "Enhancing student success prediction: A comparative analysis of machine learning technique," *TechTrends*, vol. 69, no. 2, pp. 372-384, 2025.
- [15] X. Gu *et al.*, "Digital twin technology for intelligent vehicles and transportation systems: A survey on applications, challenges and future directions," *IEEE Communications Surveys & Tutorials*, 2025.
- [16] S. Khairnar "EDGE COMPUTING FOR IOT DEVICES: A COMPREHENSIVE FRAMEWORK FOR DISTRIBUTED DATA PROCESSING AND REALTIME ANALYTICS," *Journal of Integrated Research*, vol. 2, no. 1, 2021.