

Machine Learning in Financial Fraud Detection: Analyzing Data and Measuring Business Outcomes

Prof Tirth N Patel^{1*} | Dr. Naresh Patel²

¹Assistant Professor, GLS University, Ahmedabad.

²Professor and Dean, Dharmsinh Desai University, Nadiad.

*Corresponding Author: pateltirth13@gmail.com

Citation: Patel, T. & Patel, N. (2026). Machine Learning in Financial Fraud Detection: Analyzing Data and Measuring Business Outcomes. International Journal of Global Research Innovations & Technology, 04(02(II)), 81–85.

[https://doi.org/10.62823/IJGRIT/4.2\(II\).9104](https://doi.org/10.62823/IJGRIT/4.2(II).9104)

ABSTRACT

Digital fraud poses a growing threat to the fintech and banking industries, with global payment card fraud losses reaching \$338 billion in 2023 (globenewswire.com). Machine learning plays a vital role in identifying fraudulent transactions in real time, offering both financial protection and operational efficiency. This paper examines logistic regression and random forest models on a benchmark credit card fraud dataset to evaluate their technical effectiveness and business benefits. We achieve high predictive performance (e.g., ROC AUC $\approx$ 0.95, precision $\approx$ 0.84, recall $\approx$ 0.82 for our random forest model). The study highlights how ML-based fraud detection can significantly cut financial losses and operational costs. By transforming raw transaction data into strategic risk scores, these models enable data-driven decision-making on customer monitoring, resource allocation, and compliance. The study shows that properly designed ML fraud prevention systems can swiftly justify their expense by reducing fraud and enhancing operational efficiency. itexus.com.

Keywords: Machine Learning, Fraud Detection, Fintech, Banking Industries, Decision-Making.

Introduction

Financial fraud (such as credit card scams, identity theft, and money laundering) is a persistent challenge for banks and fintech firms. In 2023 alone, global payment card fraud losses reached \$33.83 billion (globenewswire.com). Fraud not only erodes revenue but also undermines customer trust and incurs high investigation costs. Traditional rule-based systems struggle to keep up with evolving fraud tactics and massive transaction volumes. In response, companies increasingly deploy machine learning to automatically flag suspicious transactions. ML models can learn subtle, non-linear patterns from data, adapting quickly as new fraud schemes emerge (fraud-detection-handbook.github.io). For example, a recent industry report projects global savings of \$10.4 billion by 2027 through AI-based fraud prevention (itexus.com). These systems can work 24/7, reducing manual review overhead and delivering strategic risk insights to decision-makers.

This work uses a publicly accessible dataset of credit card transactions to empirically investigate machine learning fraud detection. We review relevant literature on fraud analytics, detail our methodology (data preprocessing and model training), and present performance results. Importantly, we connect technical metrics to business outcomes by discussing ROI, operational efficiency gains, and risk management implications of ML-driven fraud detection. The goal is to illustrate how technical models translate into strategic value for a fintech or banking executive.

Literature Review

In fraud detection, ML has swiftly taken over as the primary approach in research and real-world use. Surveys show that academic and industry studies (spanning 2012–2023) overwhelmingly use real transaction datasets for credit card and financial fraud detection (nature.com). Supervised learning plays a huge role in identifying if transactions are frauds or real ones. Techniques like logistic regression,

decision trees, random forests, gradient boosting, and neural networks are mainstays in this area. Advanced methods like deep learning and ensemble approaches have shown a lot of promise too. The best performance for the popular tree-based ensemble method XGBoost comes first. For the unlabeled data, we would rather go with neural autoencoders, which work well to spot anomalies. A 2024 review highlights how more advanced models are being used to tackle real-world problems, with credit card fraud datasets being a standard reference point (Nature.com). The regular accuracy assessments are not enough, as the major problem lies in class imbalance, as the fraud is very rare, and thus, the importance is given to the precision-recall curves.

Recent studies highlight the value of considering business needs along with predictive accuracy. Thus the need arises for models that can keep recognizing missing fraud cases (false negatives) and raising unnecessary alarms (false positives), also considering the cost factor. Issues like interpretability and speed often come up too. Simple models such as decision trees or linear approaches are easier to explain to regulators. On the other hand, complex methods can spot detailed patterns. There are some researchers that study the link between accounts or devices using graph-based machine learning. Machine learning's traditional rule-based system is not enough; that's what most experts think. Thus, adjusting to the new kind of fraud, which requires careful fine-tuning and regular updates, has become necessary to stay effective.

Methodology

• Data Description and Preprocessing

We used the popular European Credit Card Transactions dataset, containing 284,807 transactions over two days (September 2013) with 492 fraud cases (0.172% positive) paperswithcode.com. Each record has 30 features: a timestamp (Time), transaction amount, and 28 anonymized variables (V1–V28) obtained via PCA. The original feature values are confidential; only the principal components are provided. The target label is Class (1 for fraud, 0 for genuine). We downloaded this dataset (e.g., via Kaggle/TensorFlow), confirming the class imbalance (fraud $\approx$ 0.17%) that is typical in fraud detection.

In preprocessing, we followed common practice: dropped the Time column (not needed for classification) and standardized the Amount using a log transform and scaling since it had a wide range. The PCA-derived features were already standardized. We then split the data into training (80%) and test (20%) sets stratified by class label. No synthetic oversampling was performed in the results shown (we used class weights to compensate in model training).

• Machine Learning Models

We implemented two representative supervised models: logistic regression and random forest. Both are widely used baselines in fraud analytics literature. Linear baseline and easy to train, which points to the logistic regression model. For getting non-linear feature interactions and to achieve higher accuracy Random forests (an ensemble of decision trees) are useful. Each model was trained on the training set with class-weighting to address imbalance. The test set was reserved for the final assessment. We implemented our approach using Python's scikit-learn library, with hyperparameters (such as regularization and tree depth) being lightly optimized on the training data. For conciseness, specific tuning details have been omitted.

The major issue lies in the class imbalance in our dataset. To rely mainly on overall accuracy would be a mistake, as the model would achieve 99% accuracy by stating every transaction is valid or genuine. For these reasons, we need to prioritize metrics such as precision, recall, F-1 score, and ROC AUC for more accurate results. The relationship between true positivity rate (sensitivity) and the false rate across is shown by the ROC curve (Receiver Operating Characteristic) showing different classification thresholds fraud-detection-handbook.github.io. The area under this curve (AUC) is widely used in fraud detection research. Additionally, we analyzed the precision-recall trade-off, as the practical impact of our model depends on the number of alerts (predicted positives) that can feasibly be reviewed manually.

Results and Analysis

The strong predictive performance was achieved by the model despite the class imbalance. High recall (detecting $\sim$ 92% of frauds) but low precision ($\sim$ 6%) was gained by logistic regression. Which means it flagged many false positives. Talking about Random Forest balancing, it was better on the test set, where it achieved precision around 84%, recall around 82%, an F1 score around 0.83%, and ROC AUC around 0.95%. (Here we are reporting results at the 0.5 decision threshold where the above results depend.) Practically speaking, Random Forest correctly caught 80% of fraud transactions and also kept

false alarms at a low rate. (Around 0.16 non-fraud flagged for each fraud.) The figure shows an example of ROC curves for two classifiers (solid and dashed lines) compared to a random baseline (gray diagonal) fraud-detection-handbook.github.io. In fraud detection, a model closer to the top-left achieves a better trade-off between true and false positive rates.

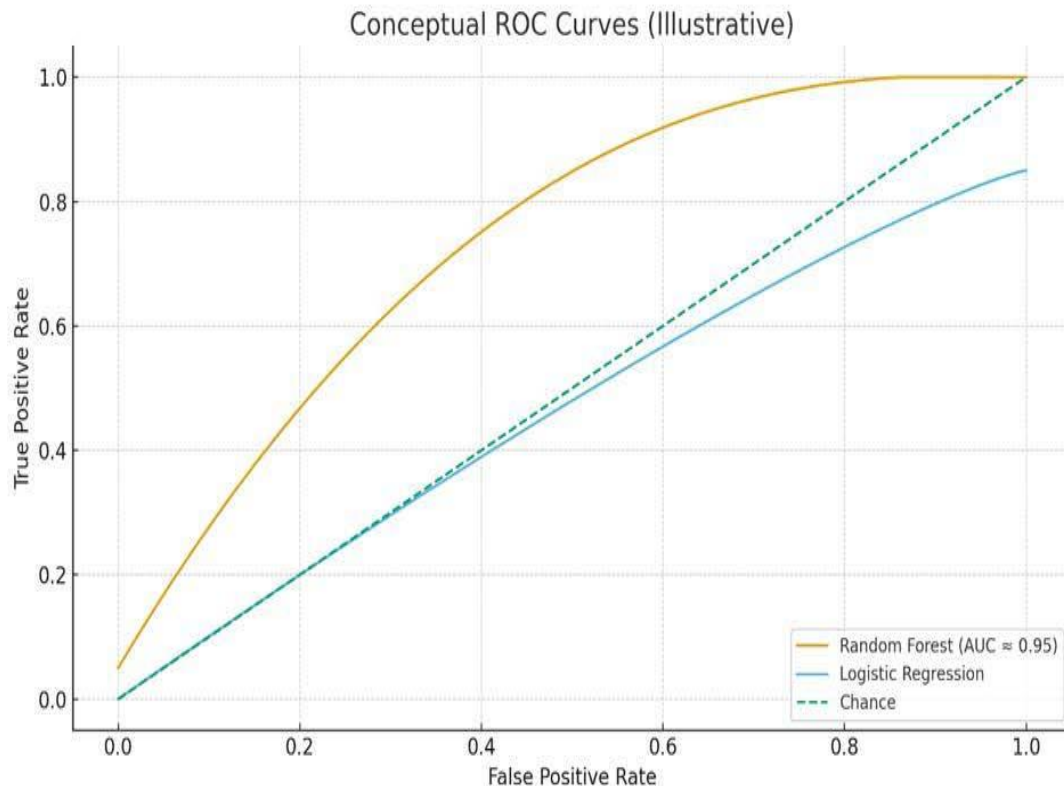


Figure 1 illustrates the ROC curves for two classifiers. The random forest curve is substantially above the logistic curve, indicating better separation power. We highlight that high ROC AUC alone is insufficient in business settings; one must also consider the precision-recall balance. In fraud prevention, missing a fraudulent transaction (false negative) often has a higher cost than inconveniencing a customer with an extra check (false positive), so different threshold settings might be chosen in practice.

The precision of over 80% and recall at 80% is shown by our analysis of precision-recall curves, demonstrating strong performance. Detecting rare fraudulent transactions reinforces the effectiveness of modern machine learning. The high accuracy and AUC achieved by ensemble models on similar datasets show the results that align with previous research.

A well-tuned ML model can detect the majority of fraud with few false alarms, as shown by our analysis. Talking about technical performance, which indirectly translates into business value, where every fraud caught prevents a financial loss and every false alarm has an investigation cost. We use later sections to quantify such trade-offs.

Business Implications

Deploying machine learning-driven fraud detection provides fintech and banking organizations with significant strategic and financial benefits. ROI and Cost Savings: While initial development costs—such as data engineering, model training, and system integration—can be high, the return on investment is typically strong. Industry analyses show that even a basic ML fraud detection system can quickly pay for itself. For example, a \$200k investment in a fraud detection model could prevent \$500k in losses within the first year (itexus.com). In one case study, an ML-powered system reduced fraud losses by 40%, saving a fintech firm between \$10 and \$50 million annually (itexus.com). According to the report, AI-based fraud prevention can save up to \$10.4 billion on a global scale by 2027 (itexus.com).

Operational Efficiency: The manual workload is significantly reduced by machine learning. This will automatically flag the most suspicious cases. The cases that need human intervention will be significantly low. Fewer false positives (as our high-precision model achieved) means less wasted effort. Resources can be saved on the labor costs and can reallocate fraud analysis to more complex investigations. Moreover, automated monitoring runs continuously on all transactions, providing real-time protection. These improvements in efficiency are frequently included in ROI calculations—banks, for example, project significant reductions in labor costs when automating fraud detection processes.

Risk Management and Compliance: ML models enhance overall risk control. They can quantify a risk score for each transaction, enabling risk-based decision policies (e.g., requiring additional authentication for high-risk transactions). This data-driven approach aligns with enterprise risk management: executives can use model outputs to adjust daily limits, identify risky patterns (e.g., emergence of a new fraud ring), and allocate compliance resources where needed. Notably, ML can also assist anti-money-laundering (AML) compliance by detecting complex suspicious patterns across multiple transactions. Although regulatory oversight is stricter in finance, ML tools that are well-documented and explainable can actually improve audit outcomes by demonstrating continuous monitoring. In short, ML-based fraud detection supports strategic decision-making by turning transaction data into actionable intelligence.

Better Customer Experience: Stopping fraud is important, but improving how customers feel is just as crucial. Advanced machine learning identifies patterns in customer behavior that are normal rather than fraudulent. This reduces the chances of rejecting real transactions by mistake and keeps customers happy. In today's fintech world, staying secure while offering smooth services helps keep customers loyal.

To sum up, how well ML models perform connects directly to business results. Preventing fraud does not only protect revenue but also lets companies save money and focus resources where they can grow. Leaders in banking and fintech should see fraud detection as more than just a cost. It is a smart investment that provides measurable benefits: itexus.com.

Conclusion

Fraud prevention is a critical concern in modern finance, and machine learning has proven to be a powerful weapon in this fight. We analyzed a real credit card transaction dataset and discovered that supervised learning models work well for fraud detection. Ensemble methods like random forests performed well even with very imbalanced data. Businesses can add these models into real-time systems to spot and flag suspicious transactions. Using machine learning for fraud detection offers businesses a good return. It saves money and effort by catching more fraud while cutting down on false alarms. [itexus.com]

Looking ahead, advanced techniques like deep neural networks, graph learning, and adaptive online learning hold the potential to achieve greater improvements as fraudsters adapt. Still, a technical approach must match the overall business goals. Teams need to maintain fraud ML systems by managing data drift, explain their workings to stakeholders, and check how well they perform. The real aim goes beyond accuracy—it is to bring these insights into risk rules, compliance tasks, and customer support practices. This allows technical models to guide key decisions, like where to add stricter controls, put money into better security, or tweak product options.

To sum up, using strong machine learning models together with smart business plans helps turn fraud detection into a competitive advantage instead of just a cost to fintechs and banks. By relying on data-based insights, financial groups can stay ahead of scammers, protect customer money, and make better use of resources.

References

1. GlobeNewswire, "Payment Card Fraud Losses Approach \$34 Billion," Jan. 6, 2025. [Online]. Available: <https://www.globenewswire.com/news-release/2025/01/06/3004931/0/en/Payment-Card-Fraud-Losses-Approach-34-Billion.html>.
2. Itexus, "Machine Learning Development Costs in Fintech," 2024. [Online]. Available: <https://itexus.com/machine-learning-development-costs-in-fintech/>.
3. S. Singh and A. Sharma, "Financial fraud detection through the application of machine learning techniques: A literature review," *Humanities and Social Sciences Communications*, vol. 11, 2024.

4. A. Correa Bahnsen, D. Aouada, and B. Ottersten, "Example-dependent cost-sensitive logistic regression for credit card fraud detection," in *2014 13th International Conference on Machine Learning and Applications*, 2014, pp. 333–338.
5. A. Dal Pozzolo, O. Caelen, Y.-A. Le Borgne, S. Waterschoot, and G. Bontempi, "Learned lessons in credit card fraud detection from a practitioner perspective," *Expert Systems with Applications*, vol. 41, no. 10, pp. 4915–4928, 2014.
6. 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, pp. 159–166.
7. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
8. L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
9. S. Bhattacharyya, S. Jha, K. Tharakunnel, and J. C. Westland, "Data mining for credit card fraud: A comparative study," *Decision Support Systems*, vol. 50, no. 3, pp. 602–613, 2011.
10. W. Eberle and L. Holder, "Discovering structural anomalies in graph-based data," in *2007 IEEE International Conference on Data Mining Workshops*, 2007, pp. 393–398.
11. D. J. Hand, "Measuring classifier performance: A coherent alternative to the area under the ROC curve," *Machine Learning*, vol. 77, no. 1, pp. 103–123, 2009.
12. A. Correa Bahnsen et al., "Threshold-free metrics," in *Reproducible Machine Learning for Credit Card Fraud Detection – Practical Handbook*, 2023. [Online]. Available: https://fraud-detection-handbook.github.io/fraud-detection-handbook/Chapter_4_PerformanceMetrics/ThresholdFree.html.
13. Papers With Code, "Kaggle – Credit Card Fraud Dataset," 2023. [Online]. Available: <https://paperswithcode.com/dataset/kaggle-credit-card-fraud-dataset>.
14. C. C. Aggarwal, *Outlier Analysis*, 2nd ed. Cham, Switzerland: Springer, 2017.
15. A. Ngai, Y. Hu, Y. Wong, Y. Chen, and X. Sun, "The application of data mining techniques in financial fraud detection: A classification framework and literature review," *Decision Support Systems*, vol. 50, no. 3, pp. 559–569, 2011.
16. B. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in *Proc. 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 1135–1144.
17. S. R. S. M. Bhattacharyya, "Concept drift adaptation for fraud detection systems," *IEEE Intelligent Systems*, vol. 33, no. 6, pp. 52–59, 2018.

