



ISSN: 2454-9940



**INTERNATIONAL JOURNAL OF APPLIED
SCIENCE ENGINEERING AND MANAGEMENT**

E-Mail :

editor.ijasem@gmail.com

editor@ijasem.org

www.ijasem.org

INTELLIGENT FRAUD PREVENTION USING MACHINE LEARNING-BASED RISK MANAGEMENT

GUIDE : P.RAMYA DEPT OF CSE(AI&DS) ELURU COLLEGE OF ENGINEERING
AND TECHNOLOGY

ROUTHU LAVANYA DEVI | SATISHCHNADRA CHOWDARY NAGALLA |
VEGIREDDY SAI DIVYA BHARGAVI | NADIPUDI VAMSI KRISHNA

COMPUTER SCIENCE WITH SPECILIZATION IN ARTIFICIAL INTELLENGENCE AND DATA SCIENCE,
ELURU COLLEGE OF ENGINEERING AND TECHNOLOGY, ELURU, ANDHRA PRADESH-534005

To Cite this Article

SEYEDEH KHADIJEH HASHEMI, SEYEDEH LEILI MIRTAHERI, SERGIO GRECO “Fraud Detection in Banking Data
by Machine Learning Techniques”, Issue 19 DECEMBER 2022, pp.26-12-22.

ABSTRACT:

This project aims to leverage machine learning, specifically Decision Tree Regressor and Classifier algorithms, to predict Air Quality Index(AQI) values using the dataset. AQI, a numerical indicator of air pollution levels, is influenced by various pollutants like particulate matter, sulphur dioxide, nitrogen dioxide etc. By cleaning and preprocessing the data, including handling non-standard formats and missing values, the project ensures data accuracy for modeling. Through regression analysis, the model evaluates values and fits the data, while classification methodologies categorically identify factors influencing air quality. By integrating Decision Tree algorithms, known for their efficiency and scalability, this project not only enhances our understanding of air quality dynamics but also serves as a catalyst for proactive measures to mitigate health risks associated with air pollution. Through predictive insights and policy recommendations, it contributes to fostering a sustainable and healthier environment, promoting public health, and advocating for policy changes to address the adverse effects of air pollution.

1. INTRODUCTION

Credit card fraud detection has become increasingly critical due to the surge in financial transactions facilitated by e-commerce and web-based banking platforms. Fraudulent activities have evolved over time, with fraudsters continuously adapting their methods to evade detection systems. The challenge lies in developing innovative techniques to enhance the accuracy and efficiency of fraud detection systems.

Fraud is defined as wrongful deception for financial or personal gain. Credit card fraud, particularly in digital transactions, involves illegal use of card details provided through telecommunication or websites. To combat fraud-related losses, two mechanisms are essential: fraud prevention, which proactively stops fraud, and fraud detection, which identifies fraudulent transactions post-attempt. Manual detection methods struggle with the vast volume of banking data, highlighting the importance of machine learning (ML)-based solutions.

ML algorithms, such as LightGBM, XGBoost, CatBoost, and logistic regression, play a pivotal role in analyzing large datasets efficiently. The paper proposes a majority-voting ensemble learning approach combining these algorithms to improve fraud detection performance on real-world unbalanced data. Bayesian optimization is suggested for hyperparameter tuning to address data imbalance. XGBoost is preferred for its fast training speed and regularization features, while CatBoost excels without intensive hyperparameter adjustments.

Deep learning techniques are utilized for fine-tuning parameters, ensuring the system adapts effectively. Evaluation metrics such as recall precision, ROC-AUC, F1-score, and MCC are applied to assess performance. The results demonstrate superior efficiency and accuracy compared to existing methods, fostering customer trust and minimizing losses for banks. The researchers further enhance accessibility by using publicly available datasets and sharing their source codes for broader use.

In conclusion, the paper emphasizes the integration of advanced ML techniques and hyperparameter optimization to create robust, efficient fraud detection systems. These systems safeguard financial institutions and their customers against evolving threats in the digital landscape.

2. RELATED WORKS

The Related Works section highlights significant contributions in the field of credit card fraud detection. Researchers have explored various innovative models and strategies to improve detection accuracy, efficiency, and cost-effectiveness.

Halvaiee and Akbari introduced the AIS-based fraud detection model (AFDM) utilizing the Immune System Inspired Algorithm (AIRS). This model enhances accuracy by 25%, reduces costs by 85%, and minimizes system response time by 40%. Bahnsen et al. proposed a transaction aggregation strategy based on periodic behavior analysis using the von Mises distribution. They developed a cost-based evaluation criterion to assess fraud detection models.

Randhawa et al. compared machine learning (ML) algorithms like Naive Bayes, decision trees, and neural networks, proposing a hybrid method with AdaBoost and majority voting. They also introduced data noise for robustness evaluation. Porwal and Mukund used clustering methods to detect fraud, focusing on data consistency to identify changes in user behavior. They demonstrated the effectiveness of precision-recall curves over ROC as evaluation metrics. Another study proposed a group learning framework addressing dataset imbalance through parallel base estimator training.

Itoo et al. tackled dataset imbalance using oversampling and tested ML algorithms like logistic regression, Naive Bayes, and K-nearest neighbors. Logistic regression outperformed others in metrics like accuracy and F1-score. Other researchers combined meta-learning ensemble techniques with a cost-sensitive learning paradigm, producing efficient results with acceptable AUC values.

Altyeb et al. introduced a Bayesian-based hyperparameter optimization approach for tuning LightGBM, demonstrating strong results in metrics like ROC-AUC and F1-score on public datasets. Xiong et al. applied advanced feature engineering techniques, outperforming traditional models when tested on the IEEE-CIS fraud dataset. Viram et al. evaluated Naive Bayes and voting classifiers, showing superior performance of voting classifiers in metrics like accuracy. Verma and Tyagi analyzed supervised ML algorithms for imbalanced datasets, highlighting the efficiency of support vector classifiers.

These approaches emphasize the critical role of advanced techniques like ensemble learning, feature engineering, and hyperparameter optimization in fraud detection. They address key challenges, such as imbalanced datasets and evolving fraud patterns, providing robust solutions. Overall, the works collectively highlight ongoing advancements and the potential of ML-driven methodologies to combat financial fraud effectively

TABLE 1. Features of the credit-card fraud dataset that is used in this paper.

Variable Name	Description	Type
$V_1, V_2, \dots, V_{28}$	Transaction feature after PCA transformation	Integer
Time	Seconds elapsed between each transaction with the first transaction	Integer
Amount	Transaction Value	Integer
Class	Legitimate or Fraudulent	0 or 1

Dataset

The "creditcard" dataset utilized in this study comprises 284,807 transaction records collected over two days in 2013. Out of these, 492 transactions are fraudulent, representing only 0.172% of the data, making it highly imbalanced. The dataset includes numerical attributes derived through Principal Component Analysis (PCA).

These attributes are denoted as V_1 through V_{28} . Additionally, the dataset contains two untransformed features—"Time" and "Amount." While "Time" measures the time elapsed (in seconds) between each transaction and the first transaction, "Amount" reflects the monetary value of the transaction.

The "Class" variable indicates whether a transaction is fraudulent (1) or legitimate (0). Due to confidentiality, no background details about the dataset are available. The dataset is publicly accessible, offering a practical foundation for evaluating the proposed framework.

This dataset contains only numerical input variables resulting from a principle component analysis (PCA) transformation. Unfortunately, the original features and background information about the data are not given due to confidentiality and privacy considerations. PCA yielded the following principal components: $V_1; V_2; V_{28}$. The untransformed features with PCA are "time" and "amount." The "Time" column contains the time (in seconds) elapsed between each transaction and the first transaction in the dataset. The feature "Amount" shows the transaction amount. Feature "Class" is the response variable, and it takes the value 1 in case of fraud and 0 otherwise. The summary of the variables and features is presented in Table 1.

Data Pre-Processing

Handling imbalanced datasets poses a significant challenge in credit card fraud detection. Imbalance skews machine learning (ML) models, adversely affecting their ability to detect minority class instances accurately. Common approaches include undersampling and oversampling.

While undersampling reduces data volume, it results in data loss. Oversampling, on the other hand, duplicates minority class data, which often does not enhance model performance. Synthetic Minority Oversampling Technique (SMOTE) addresses these limitations by generating synthetic samples, but it may inadvertently increase false-positive rates—a major drawback in customer-centric industries like banking.

To mitigate these issues, the study employs class weight tuning as an alternative, effectively addressing data imbalance without introducing the limitations inherent in other methods.

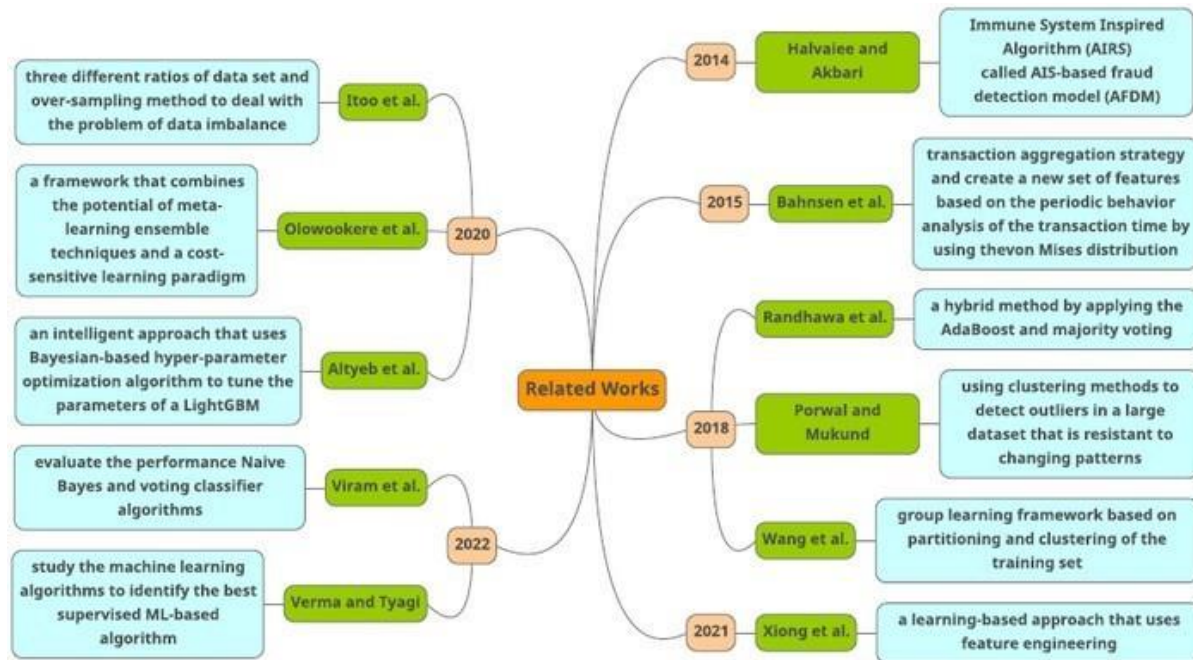


FIGURE 1. Summary of the related works on fraud detection in banking industry with machine learning techniques.

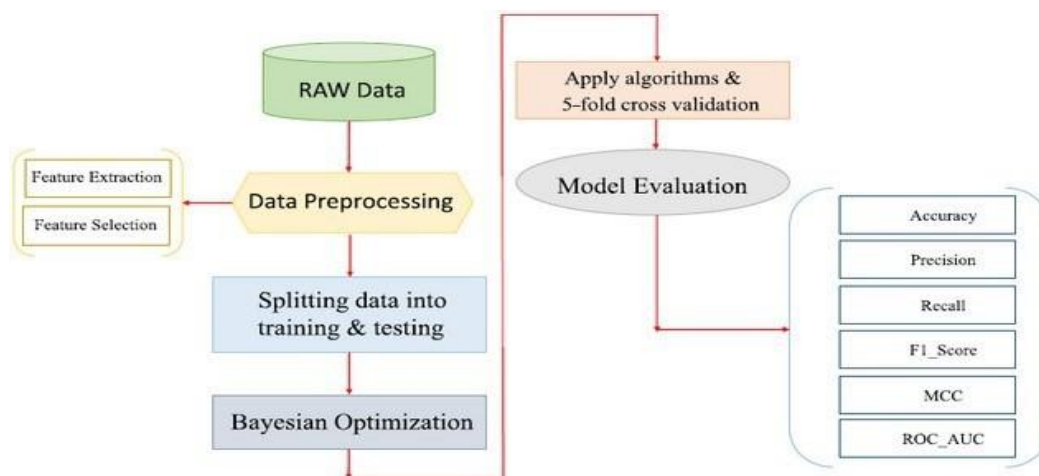


FIGURE 2. Proposed framework for credit card fraud detection.

Feature Extraction and Selection

Feature extraction involves transforming the “Time” variable to include transaction hour information, offering greater analytical insight. As the dataset does not provide feature descriptions apart from “Time” and “Amount,” feature selection plays a crucial role in improving detection accuracy.

The Information Gain (IG) method is employed for feature selection. IG quantifies the predictive power of each feature by analyzing the class (fraudulent or legitimate) associations, ultimately identifying the most relevant attributes for classification.

This technique reduces the dimensionality of the training data while ensuring computational efficiency. The study evaluates the top six features determined by IG for their contribution to fraud detection.

Algorithms and Hyperparameter Optimization

The framework incorporates three machine learning algorithms: Logistic Regression, LightGBM, and XGBoost, with hyperparameter tuning achieved through Bayesian optimization.

1. Logistic Regression:

This statistical method models the relationship between a binary-dependent variable and one or more independent variables.

Due to its limitations in handling imbalanced data, hyperparameter tuning is used to adjust class weights, enhancing performance in this context.

2. LightGBM:

Light Gradient Boosting Machine (LightGBM) is a highly efficient implementation of the Gradient Boosting Decision Tree (GBDT) framework. It is specifically designed to handle large-scale datasets and deliver high computational performance in scenarios involving complex prediction tasks. LightGBM differentiates itself by adopting a histogram-based algorithm, which partitions continuous features into discrete bins, thereby reducing memory usage and speeding up training processes.

This approach ensures that the algorithm can efficiently process vast amounts of data without being hindered by computational bottlenecks.

3. XGBoost:

eXtreme Gradient Boosting (XGBoost) is a widely used machine learning algorithm renowned for its speed, accuracy, and computational efficiency. It builds decision trees iteratively, correcting errors of previous models, a hybrid approach that refines predictions. XGBoost employs a parameter known as "max depth," which enables backward tree pruning to improve computational speed and performance.

To control overfitting, XGBoost uses a regularization technique called "formalization," ensuring that the model is not overly complex. The hyperparameters tuned for optimization include learning rate, maximum tree depth, and the number of trees, alongside applying class weights to handle imbalanced data.

The algorithm's parallel computing capability enhances efficiency, allowing for the simultaneous use of CPU resources during training, which makes it particularly effective for large datasets.

4. CatBoost:

Category Boosting (CatBoost) is an innovative gradient boosting algorithm that excels at managing highly unbalanced data and categorical features. Proposed by Prokhorenkova et al., CatBoost incorporates Bayesian estimators to address overfitting issues effectively. The algorithm is capable of handling data of diverse types and formats and does not require extensive pre-training, unlike other machine learning models.

Moreover, CatBoost features both CPU and GPU implementations, with the GPU version enabling much faster training speeds compared to other algorithms like XGBoost and LightGBM.

To further enhance its performance, random permutation of the dataset is utilized, combined with class weight hyperparameter tuning to manage data imbalance problems. CatBoost's computational efficiency makes it a competitive choice for large-scale fraud detection projects.

5. Majority Voting:

Majority voting is an ensemble learning technique that combines the predictive power of multiple classifiers, thereby reducing errors and producing more reliable results compared to single algorithms. It is implemented through parallel training and evaluation of classifiers, using the unique advantages of each algorithm. The final prediction outcome is determined through two strategies: hard voting and soft voting.

Hard voting relies on predicted class labels to vote for the majority law, whereas soft voting uses probabilistic predictions summed through the "Argmax" function, yielding class labels based on the highest averaged probability across classifiers.

This technique, especially when applied to well-calibrated classifiers, enhances fraud detection accuracy significantly.

7. Deep Learning:

Deep learning is a subset of machine learning algorithms featuring multi-layered artificial neural networks (ANNs). It is particularly suited for complex tasks like fraud detection in financial transactions, utilizing big data efficiently.

Inspired by biological neurons, deep learning networks consist of interconnected processing units that discover hierarchical representations. Each layer learns intermediate concepts to refine predictions. In this framework, a sequential model is used to construct the ANN, incorporating dense layers and activation functions like Relu and Sigmoid.

Relu outputs zero for non-positive values, while Sigmoid scales outputs between zero and one, ideal for binary classification tasks. The model's weights are initialized using "kernelinitializer," ensuring efficient random weight assignment.

To address data imbalance, class weights are set to a 1:4 ratio between majority and minority classes, optimizing model performance and processing speed. Hyperparameters such as the number of layers, neurons, epochs, and batch size are fine-tuned using Bayesian optimization to enhance speed and efficiency. For training, Adam's optimizer and binary cross-entropy are employed to refine the learning process.

Layer(Type)	Output Shape	Param No.
dense (Dense)	(None, 86)	2752
dense-1 (Dense)	(None, 44)	3828
dense-2 (Dense)	(None, 22)	990
dense-3 (Dense)	(None, 1)	23

TABLE 3. Details of our deep learning model used in the paper are provided. The total parameters are set to 7593, and all are trainable.

Evaluation Metrics:

The evaluation of the proposed credit card fraud detection model is performed using a cross-validation test. Specifically, a stratified 5-fold validation approach is applied to ensure reliable performance comparisons on the unbalanced dataset. This method involves dividing the dataset randomly into five equal subsets while maintaining proportional representation of samples across both legitimate and fraudulent transaction categories. During each step of the validation process, one subset (representing 20% of the total dataset) is reserved as validation data to evaluate the model's performance, while the remaining four subsets (80% of the dataset) are used for training purposes.

This process is repeated across all five subsets, and the final evaluation result is computed as the average performance of the model across these folds. This ensures robustness and minimizes biases caused by data distribution discrepancies. To fairly compare the proposed model's effectiveness, various metrics are employed to evaluate classification performance comprehensively. These include Accuracy, Precision, Recall, F1-Score, Matthews Correlation Coefficient (MCC), and ROC-AUC diagrams. Fraudulent transactions (positive samples) and legitimate transactions (negative samples) are classified based on the following definitions:

True Positive (TP): Fraudulent transactions correctly classified as fraudulent.

False Positive (FP): Legitimate transactions incorrectly classified as fraudulent.

True Negative (TN): Legitimate transactions correctly classified as legitimate.

False Negative (FN): Fraudulent transactions mistakenly classified as legitimate.

Accuracy:

Accuracy quantifies the total number of correct predictions made by the model, expressed as the ratio of TP and TN to the total samples in the dataset:

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}}$$

While widely used, accuracy can be misleading for highly imbalanced datasets, as detecting even a single fraudulent transaction may disproportionately inflate the score.

Recall:

Recall measures the model's ability to identify actual fraudulent transactions effectively. It is defined as:

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}$$

Precision:

Precision evaluates the reliability of positive predictions, measuring the proportion of correctly identified fraudulent transactions out of all predicted positives:

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}$$

F1-Score

The F1-Score provides a balanced measure by calculating the harmonic mean of Precision and Recall. It emphasizes both false positives and false negatives:

$$\text{F1-Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Matthews Correlation Coefficient (MCC):

MCC offers a balanced evaluation of classification performance, considering TP, TN, FP, and FN, particularly for imbalanced datasets:

$$\text{MCC} = \frac{\text{TP} \times \text{TN} - \text{FP} \times \text{FN}}{(\text{TP} + \text{FP})(\text{TN} + \text{FN})(\text{TP} + \text{FN})(\text{TN} + \text{FP})}$$

ROC-AUC:

This graphical plot represents the True Positive Rate (TPR) against the False Positive Rate (FPR) across different thresholds.

While it demonstrates the model's ability to separate classes effectively, it is not suitable for fraud detection due to its focus solely on positive values.

Precision-Recall Curve:

This graph visualizes precision rates on the y-axis and recall rates on the x-axis, enabling detailed comparisons of classifier performance. Since no single indicator can simultaneously describe TP, TN, FP, and FN, MCC is preferred for its balanced evaluation in two-class problems, even when the dataset classes are disproportionately sized.

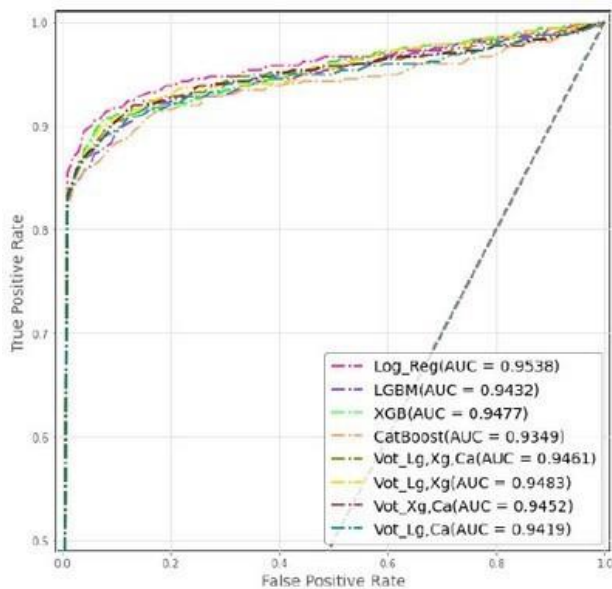


FIGURE 4. ROC_AUC curve.

Accuracy quantifies the total performance of the classifier and is defined as the number of correct predictions made by the model. When dealing with data that isn't balanced, this criterion doesn't give good results because it also gives a high value if even one fraudulent transaction is found. Recall shows the efficiency of the classifier in detecting actual fraudulent transactions. Precision measures the reliability of the classifier and F1-Score is the harmonic average of recall and precision measures, that considers both false negatives and positives.

ROC-AUC is a measure of separability that demonstrates the model's ability to differentiate between classes [15]. ROC-AUC is a graphical plot of the false positive rate (FPR) and the true positive rate (TPR) at different possible levels [17]. The area under the ROC curve is not a suitable criterion for evaluating fraud detection methods since it only considers positive values.

The precision and recall curves are commonly used to compare classifiers in terms of precision and recall. Usually, in this two-dimensional graph, the precision rate is plotted on the y-axis and the recall is plotted on the x-axis. There is no good way to describe the true and false positives and negatives using one indicator. One good solution is to use MCC, which measures the quality of a two-class problem, taking into account the true and false positives and negatives. It is a balanced measure, even when the classes are of different sizes [6].

Experimental Results and Discussion:

This section evaluates the performance of the proposed credit card fraud detection framework using stratified 5-fold cross-validation to ensure robust assessments of the unbalanced dataset. To optimize accuracy and efficiency, boosting algorithms, combined with Bayesian optimization, were utilized to fine-tune hyperparameters.

The experiments systematically evaluated individual algorithms before integrating them using a majority voting ensemble approach. Performance assessments were conducted under both double precision and triple precision configurations, ensuring comprehensive validation of the framework.

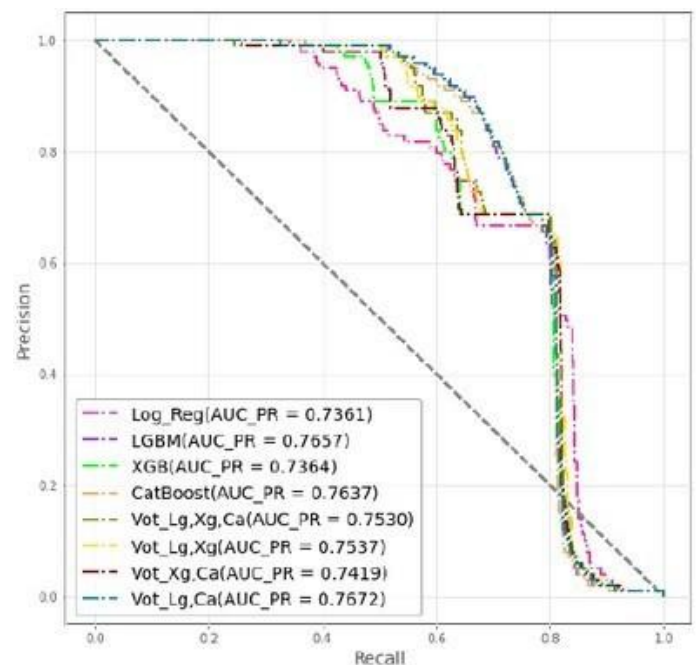


FIGURE 5. Precision_Recall curve.

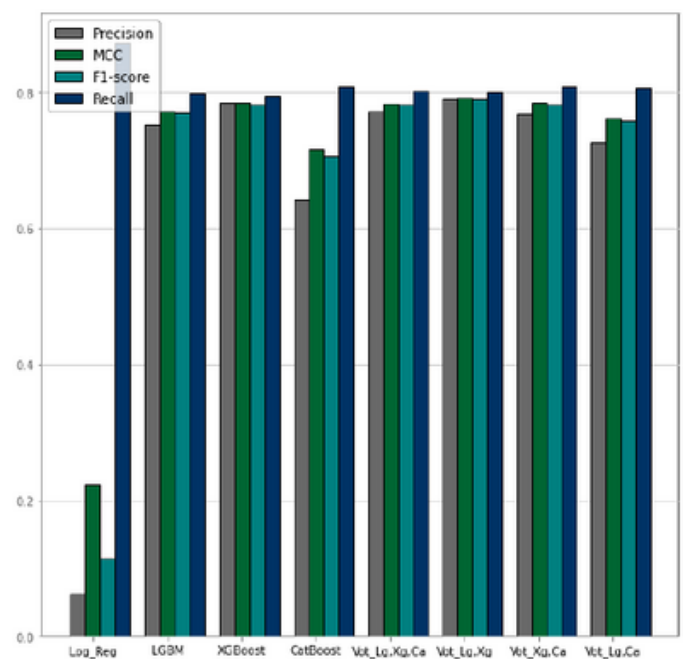


FIGURE 6. Performance comparing algorithms with different evaluation criteria.

Model	Accuracy	AUC	Recall	Precision	F1-score	MCC
Log_Reg	0.97477	0.9578	0.8730	0.0617	0.1143	0.2248
LGBM	0.99919	0.9472	0.7990	0.7534	0.7699	0.7727
XGB	0.99923	0.9517	0.7949	0.7862	0.7830	0.7864
CatBoost	0.99880	0.9390	0.8096	0.6431	0.7066	0.7158
Vot_Lg, Xg, Ca	0.99924	0.9501	0.8033	0.7720	0.7825	0.7847
Vot_Lg, Xg	0.99927	0.9522	0.8012	0.7901	0.7901	0.7925
Vot_g, Ca	0.99923	0.9492	0.8097	0.7681	0.7823	0.7852
Vot_Lg, Ca	0.99912	0.9459	0.8075	0.7260	0.7581	0.7620

TABLE 4. Performance evaluation of algorithms.

Model	Accuracy	AUC	Recall	Precision	F1-score	MCC
Keras	0.9994	0.9401	0.8222	0.8043	0.8132	0.8129

TABLE 5. Deep learning model results.

The precision-recall curve is illustrated in Fig. 5 and shows the system performance in a more precise manner compared with the ROC-AUC curve. However, the results cannot be cited because false negatives are far from the view of this diagram. As Fig. 5 shows, the highest value belongs to the combination of the CatBoost and LightGBM algorithms with a value of 0.7672, and the lowest value belongs to logistic regression and is 0.7361.

Comparing the precision, recall, and F1-score as well as the MCC, the algorithms used are shown in Fig. 6. The best performance is related to the combination of lightGBM and XGBoost algorithms, which have an MCC value of 0.79 and an F1-score of 0.79. In individual algorithms, XGBoost has the highest values. According to the digits obtained in Table 5, deep learning has achieved better performance compared with individual algorithms and majority voting ensemble learning. The MCC and F1-score metrics have values of 0.8129 and 0.8132, respectively. The area under the ROC curve in the deep learning method is illustrated in Fig. 7 and shows a value of 0.9401.

Model	Accuracy	AUC	Recall	Precision	F1-score
Method presented in [17]	0.984	0.909	0.406	0.973	0.569
Proposed LightGBM	0.9992	0.947	0.799	0.753	0.769
Proposed Approach	0.9993	0.952	0.801	0.79	0.79

TABLE 6. Performance comparison of the proposed approach and the method presented in [17].

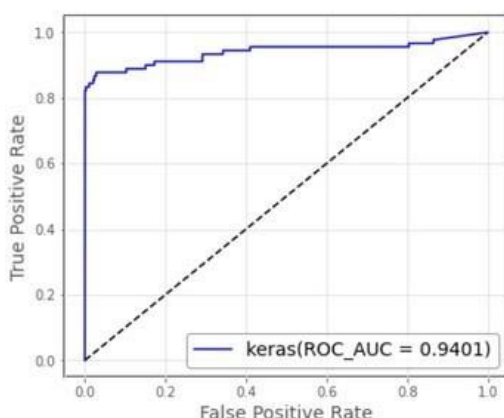


FIGURE 7. ROC curve of deep learning.

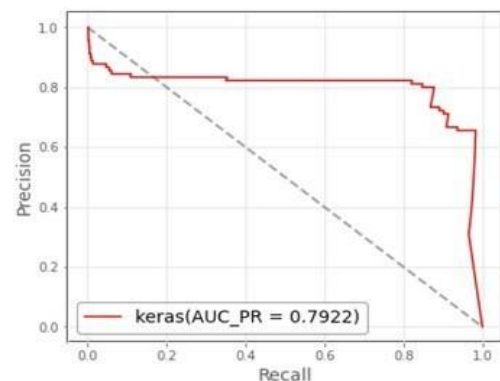


FIGURE 8. Precision- recall curve of deep learning.

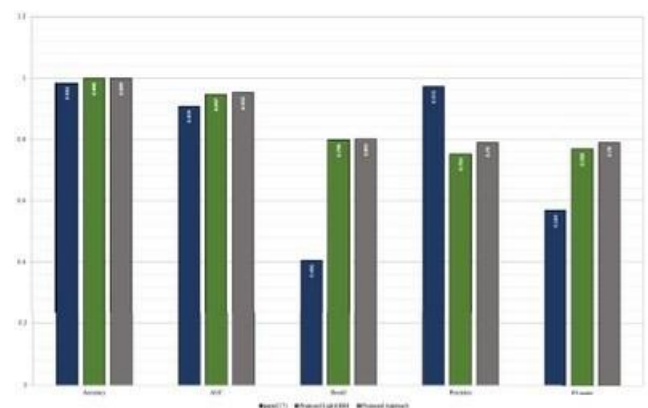


FIGURE 9. Performance comparison of the proposed approach with the paper [17] based on the different evaluation criteria.

The diagram of the Precision-Recall curve is shown in Fig. 8, and shows the value as 0.7922.

The evaluation results of the proposed approach using different pre-processing and class weight hyperparameter tuning to deal with the problem of data unbalance compared to the paper [17] are shown in Fig. 9. The results show improvement of both methods compared to the method presented in [17].

According to the Table 6, it is shown that the proposed methods outperform the intelligence method presented in [17] using common metrics and a public dataset.

Conclusion and Future Work:

This study addressed the critical issue of credit card fraud detection using real-world, unbalanced datasets. We proposed a machine learning-based approach to improve detection performance and reduce computational costs. The experiments were conducted on a publicly available "credit card" dataset containing 28 features, with fraud data constituting only 0.17% of the total records. Two primary methodologies were explored:

LightGBM with hyperparameter tuning and the integration of a majority voting ensemble. Additionally, deep learning techniques were leveraged to enhance the system's predictive accuracy further.

The LightGBM model was optimized using class weight tuning, enabling it to handle the significant class imbalance more effectively without relying on conventional sampling methods. Common evaluation metrics, such as accuracy, precision, recall, F1-score, and AUC, were employed to quantify the model's performance. The experimental results indicated a substantial improvement: the LightGBM model demonstrated a 50% enhancement in fraud detection rates and a 20% increase in the F1-score compared to methods recently presented in the literature [17].

Furthermore, the majority voting ensemble significantly improved the detection accuracy by leveraging the strengths of multiple machine learning algorithms. Deep learning methods demonstrated superior robustness, with Matthews Correlation Coefficient (MCC) values and F1-scores reaching 0.8129 and 0.8132, respectively. These results confirm that deep learning, when combined with optimized algorithms like LightGBM and XGBoost, provides a scalable and efficient solution for fraud detection. The MCC metric, tailored for unbalanced datasets, emerged as a stronger and more reliable evaluation criterion compared to other metrics.

Future Work

To expand upon this research and address its limitations, we propose the following directions for future work:

Exploration of Additional Hybrid Models: Investigate new combinations of machine learning and deep learning models. Special attention should be given to the integration of CatBoost, with a focus on tuning additional hyperparameters, such as the number of trees, to optimize performance further.

Improved Hardware Utilization: Due to hardware constraints during this study, experiments were conducted with limited computational resources. Employing advanced hardware, such as GPUs or TPUs, can accelerate training processes and enhance the accuracy of future models, enabling more comprehensive experimentation.

Dynamic Feature Engineering: Extend the feature extraction techniques by incorporating external factors such as user behavioral patterns, geographical information, or real-time transaction tracking, which may enhance the fraud detection capability.

Ethical and Explainable AI: Develop interpretable machine learning models that provide clear explanations of fraud detection decisions, ensuring ethical implementation and compliance with financial regulations.

Large-Scale Benchmark Datasets: Validate the proposed methodologies on larger and more diverse datasets from various domains. This would ensure the generalizability and scalability of the approach across multiple financial systems.

Real-Time Fraud Detection: Explore advancements in real-time fraud detection by improving inference speeds and integrating detection systems into live financial environments. Adaptive frameworks that can evolve dynamically with changing fraud patterns would further improve accuracy.

By implementing these directions, future studies can address current limitations, optimize fraud detection techniques further, and contribute to creating more robust and scalable solutions to protect financial ecosystems.

REFERENCES

- [1] J. Nanduri, Y.-W. Liu, K. Yang, and Y. Jia, "Ecommerce fraud detection through fraud islands and multi-layer machine learning model," in Proc. Future Inf. Commun. Conf., in Advances in Information and Communication. San Francisco, CA, USA: Springer, 2020, pp. 556–570.
- [2] I. Matloob, S. A. Khan, R. Rukaiya, M. A. K. Khatkhat, and A. Munir, "A sequence mining-based novel architecture for detecting fraudulent transactions in healthcare systems," IEEE Access, vol. 10, pp. 48447–48463, 2022.
- [3] H. Feng, "Ensemble learning in credit card fraud detection using boosting methods," in Proc. 2nd Int. Conf. Comput. Data Sci. (CDS), Jan. 2021, pp. 7–11.
- [4] M. S. Delgosha, N. Hajiheydari, and S. M. Fahimi, "Elucidation of big data analytics in banking: A four-stage delphi study," J. Enterprise Inf. Manage., vol. 34, no. 6, pp. 1577–1596, Nov. 2021.
- [5] M. Puh and L. Brkić, "Detecting credit card fraud using selected machine learning algorithms," in Proc. 42nd Int. Conv. Inf. Commun. Technol., Electron. Microelectron. (MIPRO), May 2019, pp. 1250–1255.
- [6] K. Randhawa, C. K. Loo, M. Seera, C. P. Lim, and A. K. Nandi, "Credit card fraud detection using AdaBoost and majority voting," IEEE Access, vol. 6, pp. 14277–14284, 2018.
- [7] N. Kumaraswamy, M. K. Markey, T. Ekin, J. C. Barner, and K. Rascati, "Healthcare fraud data mining methods: A look back and look ahead," Perspectives Health Inf. Manag., vol. 19, no. 1, p. 1, 2022.

- [8] E. F. Malik, K. W. Khaw, B. Belaton, W. P. Wong, and X. Chew, "Credit card fraud detection using a new hybrid machine learning architecture," *Mathematics*, vol. 10, no. 9, p. 1480, Apr. 2022.
- [9] K. Gupta, K. Singh, G. V. Singh, M. Hassan, G. Himani, and U. Sharma, "Machine learning based credit card fraud detection: A review," in *Proc. Int. Conf. Appl. Artif. Intell. Comput. (ICAAIC)*, 2022, pp. 362–368.
- [10] R. Almutairi, A. Godavarthi, A. R. Kotha, and E. Ceesay, "Analyzing credit card fraud detection based on machine learning models," in *Proc. IEEE Int. IoT, Electron. Mechatronics Conf. (IEMTRONICS)*, Jun. 2022, pp. 1–8.
- [11] N. S. Halvaie and M. K. Akbari, "A novel model for credit card fraud detection using artificial immune systems," *Appl. Soft Comput.*, vol. 24, pp. 40–49, Nov. 2014.
- [12] A. C. Bahnsen, D. Aouada, A. Stojanovic, and B. Ottersten, "Feature engineering strategies for credit card fraud detection," *Expert Syst. Appl.*, vol. 51, pp. 134–142, Jun. 2016.
- [13] U. Porwal and S. Mukund, "Credit card fraud detection in e-commerce: An outlier detection approach," 2018, arXiv:1811.02196.
- [14] H. Wang, P. Zhu, X. Zou, and S. Qin, "An ensemble learning framework for credit card fraud detection based on training set partitioning and clustering," in *Proc. IEEE SmartWorld, Ubiquitous Intell. Comput., Adv. Trusted Comput., Scalable Comput. Commun., Cloud Big Data Comput., Internet People Smart City Innov. (Smart-World/SCALCOM/UIC/ATC/CBDCOM/IOP/SCI)*, Oct. 2018, pp. 94–98.
- [15] F. Ito, M. Meenakshi, and S. Singh, "Comparison and analysis of logistic regression, Naïve Bayes and knn machine learning algorithms for credit card fraud detection," *Int. J. Inf. Technol.*, vol. 13, no. 4, pp. 1503–1511, 2021.
- [16] T. A. Olowookere and O. S. Adewale, "A framework for detecting credit card fraud with cost-sensitive meta-learning ensemble approach," *Sci. Afr.*, vol. 8, Jul. 2020, Art. no. e00464.
- [17] A. A. Taha and S. J. Malebary, "An intelligent approach to credit card fraud detection using an optimized light gradient boosting machine," *IEEE Access*, vol. 8, pp. 25579–25587, 2020.
- [18] X. Kewei, B. Peng, Y. Jiang, and T. Lu, "A hybrid deep learning model for online fraud detection," in *Proc. IEEE Int. Conf. Consum. Electron. Comput. Eng. (ICCECE)*, Jan. 2021, pp. 431–434.
- [19] T. Vairam, S. Sarathambekai, S. Bhavadharani, A. K. Dharshini, N. N. Sri, and T. Sen, "Evaluation of Naïve Bayes and voting classifier algorithm for credit card fraud detection," in *Proc. 8th Int. Conf. Adv. Comput. Commun. Syst. (ICACCS)*, Mar. 2022, pp. 602–608.
- [20] P. Verma and P. Tyagi, "Analysis of supervised machine learning algorithms in the context of fraud detection," *ECS Trans.*, vol. 107, no. 1, p. 7189, 2022.
- [21] J. Zou, J. Zhang, and P. Jiang, "Credit card fraud detection using autoencoder neural network," 2019, arXiv:1908.11553, Mar. 2021.

Author's profile



P. RAMYA is an Assistant professor in Eluru college of engineering and technology. Her blend of programming expertise, analytical thinking, and innovation makes her a valuable asset in technology and data-driven decision-making. Dedicated to continuous growth, she thrives in addressing complex challenges.
Mail.id: ramyaparasa99@gmail.com



ROUTHU LAVANYA DEVI is an AWS Academy Graduate skilled in C, Python, and web development. She excels in deep learning, machine learning, cloud computing, networking. Her blend of programming expertise, analytical thinking, and innovation makes her a valuable asset in technology and data-driven decision-making. In the department of cse-artificial intelligence and data science
Mail.id: lavanyadevirouthu@gmail.com



Satish Chandra Chowdary Nagalla is an AWS Academy Graduate skilled in C, Python, and web development. He excels in deep learning, machine learning, cloud computing, networking. With a blend of programming expertise, analytical thinking, and innovation, he is a valuable asset in technology and data-driven decision-making. In the department of cse-artificial intelligence and data science
Mail.id: satishchowdary1477@gmail.com



VEGIREDDY SAI DIVYA BHARGAVI is an AWS Academy Graduate skilled in C, Python, and web development. She excels in deep learning, machine learning, cloud computing, networkings. Her blend of programming expertise, analytical thinking, and innovation makes her a valuable asset in technology and data-driven decision-making. In the department of cse-artificial intelligence and data science
Mail.id: vsaidivyaabhargavi@gmail.com



NADIPUDI VAMSI KRISHNA is an AWS Academy Graduate skilled in C, Python, and web development. He excels in deep learning, machine learning, cloud computing, networking. With a blend of programming expertise, analytical thinking, and innovation, he is a valuable asset in technology and data-driven decision-making. In the department of cse-artificial intelligence and data science
Mail.id: vk6636650@gmail.com