



Feature Importance in Credit Card Fraud Detection: A Comparative Analysis of Location and Temporal Features

Sachit Kumar Purohit¹, Pradip Kumar Sahu², Manas Ranjan Senapati³

¹UG Scholar, Dept. CSE, Veer Surendra Sai University of Technology, Burla, Sambalpur, Odisha, India-768018

^{2,3}Associate professor, Dept. of CSE, Veer Surendra Sai University of Technology, Burla, Sambalpur, Odisha, India-768018

Email ID: sachitkumarpurohit@gmail.com¹, pksahu_it@vssut.ac.in², mrsenapati_it@vssut.ac.in³

Abstract

Fraud detection in financial transactions is a critical challenge requiring robust machine learning techniques. In order to identify fraudulent credit card activity, this study assesses models such as Long Short-Term Memory (LSTM) networks, Random Forest, Decision Trees, and Logistic Regression. Key features such as transaction location (lat, long, merch_lat, merch_long), merchant details (zip, distance), and temporal data (unix_time) were crucial in identifying fraud patterns. Experimental results showed that tree-based models, particularly Random Forest, achieved superior performance with 99.94% accuracy, while LSTM effectively captured sequential data patterns. Random Forest's ability to handle feature interactions and imbalances made it the most reliable. Analysis of ROC curves highlighted models' learning behavior and generalization. This research emphasizes integrating spatial and temporal features to advance adaptive, real-time fraud prevention systems.

Keywords: Feature Selection; Long Short-Term Memory; ROC; SMOTE.

1. Introduction

The increase in digital transactions in recent years, especially as a result of e-commerce and online payment systems, has revolutionized financial operations. A significant increase in credit card fraud has unfortunately accompanied this growth, posing serious threats to both consumers and financial institutions. Reports indicated that losses from credit card fraud reached approximately \$35 billion by 2023 (E. Esenogho et al., 2022), highlighting the urgent need for effective detection mechanisms. Traditional fraud prevention strategies, such as data encryption and tokenization, while useful, have proven insufficient against increasingly sophisticated fraudulent tactics. Therefore, cutting-edge technologies like deep learning (DL) and machine learning (ML) are being used more and more to improve fraud detection skills. Machine learning has become essential in analyzing large datasets and identifying fraud patterns. Detecting fraudulent

transactions is effectively achieved using algorithms like decision trees, SVM and neural networks. However, the class imbalance in transaction datasets, where legitimate transactions outnumber fraudulent ones, presents a major challenge, often resulting in high false-positive rates and undermining consumer trust. Techniques such as SMOTE have been used to address this imbalance, while hybrid approaches combining multiple algorithms show promise in improving detection accuracy. Despite progress, challenges with interpretability and adapting to evolving fraud tactics remain. Ongoing research aims to refine models, integrate ML with DL techniques, and employ advanced preprocessing methods to enhance fraud detection systems in digital financial environments [1][2].

2. Related Works

SMOTE-ENN combined with boosted LSTM has been employed to address class imbalance in fraud

detection. This approach has enhanced model sensitivity and specificity, having outperformed traditional classification methods in accuracy and robustness (E. Esenogho et al.,2022). Addressing class imbalance has been recognized as crucial due to its impact on skewed results and elevated false-positive rates. Traditional models have been evaluated, revealing that accuracy alone has been found insufficient, necessitating appropriate metrics and scalable solutions (S. Makki et al.,2019). Genetic algorithm-based feature selection has been combined with classifiers like RF, DT, and ANN, enhancing performance by reducing dimensionality. The GA-RF model has been shown to outperform others on synthetic and real datasets, having improved accuracy and handling of imbalanced data (E. Ileberi et al.,2022) [3]. Advancements in big data, IoT, and real-time processing have been leveraged to enhance credit card fraud detection in digital payment environments. Deep learning models and cloud-based frameworks have been developed, yet a gap has been identified in addressing large-scale digital financial ecosystems. (A. Cherif et al.,2023). Machine learning approaches, including supervised, unsupervised, anomaly detection, and ensemble methods, have been applied to credit card fraud detection. Feature selection and resampling techniques have been used to address class imbalance, while deep learning models have been employed to capture complex patterns despite challenges in interpretability and adaptability (I.D. Mienye et al.,2024). Dynamic fraud detection models have been developed to adapt to evolving transaction behaviors by building individual cardholder profiles. Concept drift has been addressed through customer grouping and parameter tuning, while techniques like SMOTE and Matthews Correlation Coefficient have been applied to enhance classifier performance (V. N. Dornadula et al.,2019). Ensemble models combining AdaBoost with voting techniques have been shown to offer robustness against noisy data and fluctuations in data quality [5]. These hybrid models have demonstrated promising performance on real credit card datasets, achieving high MCC scores despite significant noise, proving their resilience for real-world applications (Kuldeep Randhawa et al.,2018).

Incorporating deep neural networks (DNNs) has been shown to reduce false positives in fraud detection systems, yielding promising results. Optimized DNN configurations have demonstrated strong potential in distinguishing valid from false alerts, enhancing fraud capture rates and assisting human-driven fraud assessments (R. S. M. Carrasco et al.,2020). Hybrid ML techniques combining LSTM with attention mechanisms have been used to enhance fraud detection accuracy by prioritizing relevant data. Integrating SMOTE and UMAP for feature selection and dimensionality reduction has improved efficiency, capturing key consumer behavior patterns and distinguishing fraudulent transactions (I.Benchaji et al.,2021). A fractal-based technique, initially used for texture classification, has been applied to fraud detection by identifying self-similar patterns in transaction data. The Pixel Range Calculation (PRC) approach has outperformed advanced methods like the Gliding Box and Multi-Fractal Spectrum, demonstrating improved accuracy and reduced processing complexity (Abadhan Ranganath et al.,2022). Demonstrated with high accuracy, CNNs and deep learning models have been utilized to capture complex transaction patterns for fraud detection [4]. Suggested for future enhancement, advanced techniques like attention mechanisms and transfer learning have been considered to improve adaptability against evolving fraud tactics (F. K. Alarfaj et al.,2022). Explored for various tasks, LSTM models have been shown to effectively capture long-term dependencies in transaction sequences. Demonstrated as adaptable, they have been leveraged to enhance fraud detection by analyzing historical behavior patterns (M. Ma et al.,2022). Recent advancements in fraud detection have been achieved using machine learning and deep learning models, with a BiLSTM-MaxPooling-BiGRU model outperforming traditional classifiers. Undersampling and oversampling techniques have been applied, with the deep learning model achieving a superior AUC of 91.37%, demonstrating its ability to handle complex patterns and imbalanced data effectively (H. Najadat et al.,2020) [6]. Synthetic sampling techniques like SMOTE have been combined with algorithms such as Isolation Forest

and LOF to enhance fraud detection accuracy in imbalanced datasets. High accuracy scores of 99.74% and 99.66% have been achieved by Isolation Forest and LOF respectively, surpassing SVM, while parallel processing and null value handling have been emphasized for improving prediction reliability (S. Warghade et al.,2020) [7]. Feature selection techniques have been analyzed for fraud detection in web transactions, with class imbalance found to adversely affect their performance. Resampling strategies, including a novel Sampling Outlier method, have been evaluated, with the method having improved financial outcomes by 57.5% and demonstrated suitability under high classification costs (R. Lima et al.,2017).

3. Methodology

3.1 Dataset

This research uses a dataset of 1,852,394 transactions to analyze credit card fraud behaviors [16]. It includes transaction details, demographic information (such as cardholder name, gender, job), and geographic data (city, state, coordinates). These features support effective feature extraction for identifying fraud indicators. Eighty percent of the data was used for training machine learning models, while twenty percent was kept for testing. The classification task targets the `is_fraud` variable, labeling transactions as fraudulent or lawful.

3.2 Exploratory Data Analysis

To understand the dataset and identify patterns for detecting fraudulent transactions, Exploratory Data Analysis (EDA) was conducted. Key analyses included examining transaction amounts, trends across merchants, and demographic factors like gender and age. Preprocessing steps prepared the dataset for analysis. The `trans_date_trans_time` feature was simplified into a datetime format, deriving three new features: transaction hour, transaction day of the week, and transaction year and month, providing temporal insights into transaction patterns. An age-based analysis was made possible by deducting the cardholder's birthdate from the transaction date to obtain the age feature [8]. To enhance the dataset's usability for fraud detection modeling, categorical variables were transformed using one-hot encoding. The category variable,

indicating transaction types, was encoded into binary columns (e.g., groceries, dining, personal care, and travel) with 1 or 0 values. Additionally, age groups (<30, 30-45, 46-60) were created from the age variable and one-hot encoded, capturing the age distribution of cardholders. These transformations enriched the dataset, improving interpretability and aiding in fraud pattern detection. The distribution of fraudulent and non-fraudulent transactions by age group is shown in Figure 1.

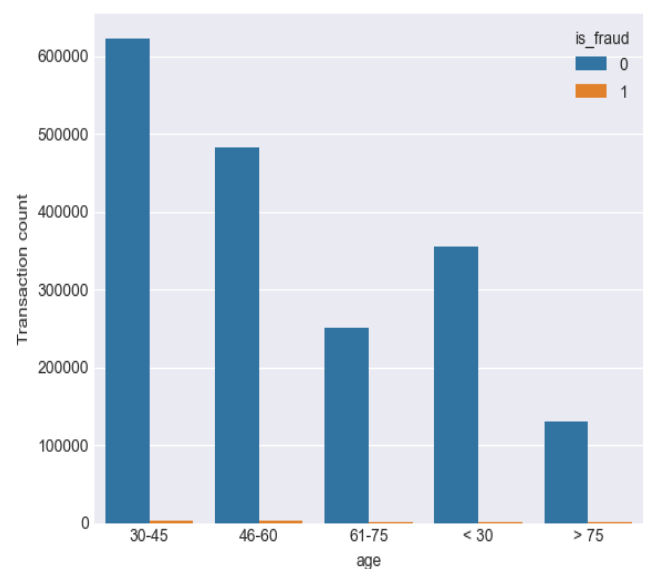


Figure 1 Distribution of Fraud and Non-Fraud Transactions Across Age Categories

The gender variable was one-hot encoded to create binary columns distinguishing male and female cardholders, enabling the inclusion of demographic characteristics without ordinal bias. Similarly, the `day_of_week` variable, derived from transaction timestamps, was one-hot encoded to capture transaction frequency across different days, aiding in detecting time-based fraud patterns [9]. The dataset also included a state feature representing the U.S. state of each transaction. Analysis revealed that a significant share of both fraudulent and non-fraudulent transactions originated from states with high transaction volumes, including Texas (TX), New York (NY), Pennsylvania (PA), and California (CA). Figure 2 illustrates the Transaction count in different states [9].

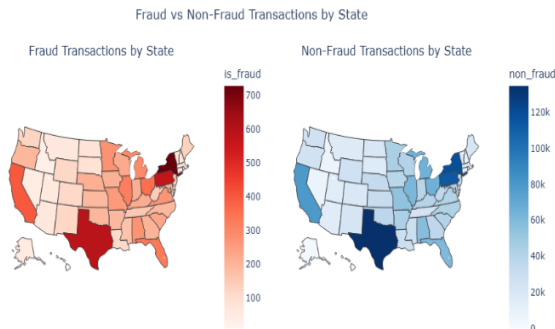


Figure 2 Fraud Vs Non-Fraud Transaction Count Across U.S. States

3.3 Feature Selection

Feature selection (FS) is crucial in machine learning, especially when datasets have many features that may affect model performance. Random Forest was used for feature selection due to its ability to rank features by their contribution to reducing impurity. The dataset was split into 80% training and 20% testing to evaluate model performance. Important features for predicting fraud were identified by ranking feature importance scores, improving model efficiency and reducing overfitting. Location-based features such as zip, lat, lang, merch_lat, and merch_long were among the top predictors [10]. The unix_time feature also emerged as significant. To capture spatial relationships, a new feature, distance, was created using the great_circle function from the geopy library. This distance measured the proximity between transaction and merchant locations using latitude and longitude coordinates in eq 1.

$$\text{distance} = \text{great_circle}((\text{lat}, \text{long}), (\text{merch_lat}, \text{merch_long})).\text{kilometers} \quad (1)$$

3.4 Resampling Techniques

Biased models arise from imbalanced datasets, which are problematic in credit card fraud detection due to the small proportion of fraudulent transactions. Resampling techniques like SMOTE help balance class distributions for better model performance. SMOTE generates synthetic samples for the minority class by interpolating between existing samples and their nearest neighbours, as shown in Eq. (2). This technique enhances diversity in the dataset, reducing overfitting and improving the model's

generalizability. SMOTE helps avoid the issues of naive oversampling by introducing meaningful variations. The synthetic sample generation process is defined by the equation:

$$[x]_{\text{new}} = x_i + \delta \cdot (x_j - x_i) \quad (2)$$

where, x_i : Minority class sample.

x_j : Neighbor of x_i .

δ : Random scalar in the range [0,1].

3.5 Methods

Decision Tree: For both classification and regression problems, supervised learning algorithms like as decision trees have been employed. A decision tree-like structure has been created by dividing the data into subsets according to feature values. Every internal node represented a feature-based choice, every branch represented a decision outcome, and every leaf node represented a value (in regression) or a class label (in classification) (E. Ileberi et al., 2022; I.D. Mienye et al., 2024;). Because scaling or normalization are not required, decision trees may handle both numerical and categorical data with little preprocessing. This simplicity makes them easy to implement with raw data. However, they are prone to overfitting when models become overly complex with deep branches. Regularization techniques like pruning, which removes unnecessary branches, help improve generalization [11]. Despite the risk of overfitting, decision trees remain popular for classification and regression tasks due to their versatility and effectiveness. The decision tree algorithm has been started by initializing the root node with the entire dataset. The best feature and threshold have been selected to split the data based on impurity metrics like Gini index or entropy. The data has been partitioned into left and right child nodes and recursively split further. This process has continued until conditions such as maximum depth, few samples, or homogeneous subsets have been met. For classification, each leaf node has been given a class label; for regression, it has been given an average value [12]. On the basis of fresh data, the final tree model has generated predictions. **Logistic Regression:** The logistic (sigmoid) function is used to a linear combination of input data in logistic regression, a statistical model for binary classification that estimates the likelihood of class

membership. Multi-class problems have been handled through methods like One-vs-Rest. Probabilities have been outputted by the model, enabling threshold-based decision-making, and a linear relationship has been assumed between independent variables and the log-odds of the dependent variable. Its simplicity and interpretability have made it a versatile tool for classification tasks (E. Ileberi et al.,2022;). The logistic regression algorithm has been used as a supervised learning method for binary classification. Weights and bias have been initialized, and predicted probabilities have been computed using the sigmoid function. The binary cross-entropy loss function calculates gradients for parameter updates via gradient descent. Weights and bias are iteratively adjusted based on the gradients and learning rate until convergence or a maximum iteration limit is reached [13]. The resulting model predicts probabilities or classifies new data points based on a threshold. Random Forest: In order to increase accuracy and generalization, Random Forest, an ensemble learning technique, mixes several decision trees using bagging and random feature selection. This approach has reduced the risk of overfitting compared to a single decision tree, making it more robust to unseen data. Both categorical and numerical data have been handled effectively, demonstrating versatility across various datasets (E. Ileberi et al.,2022;). Each decision tree in the Random Forest algorithm's ensemble was trained using a bootstrap sample that was produced by sampling with replacement. At each split, a random subset of features has been chosen to partition the data, introducing diversity among trees. All of the trees' predictions have been included in the final model, which uses regression averaging or majority voting for categorization. This randomness has reduced overfitting while maintaining high performance, making Random Forest effective for complex datasets with varied features. Long Short-Term Memory (LSTM): By adding memory cells, Long Short-Term Memory (LSTM) networks a specific kind of RNN have been utilized to solve the vanishing gradient issue [15]. Long-term dependencies in transaction sequences are crucial for fraud detection tasks, and these memory cells have

made it possible for LSTMs to retain information across lengthy periods. Three gates the input gate, forget gate, and output gate have made up LSTMs and have regulated the information flow across the network. This has allowed LSTMs to retain only the most relevant information while discarding irrelevant data (I.Benchaji et al.,2021). This study uses a Long Short-Term Memory (LSTM) network to capture long-term dependencies in sequential transaction data for fraud detection. With dimensions for samples, sequence length (five transactions), and characteristics, the dataset was pre-processed into a three-dimensional format. The final transaction in each sequence served as the basis for the target labels. Up to 20 epochs of training were conducted with a batch size of 32, and 20% of the data was used for validation. To avoid overfitting, early stopping was used, and the optimal model weights were restored after five epochs of validation loss monitoring [14].

4. Results & Discussion

4.1 Results

As shown in Table 1, the effectiveness of different machine learning models in identifying credit card fraud has been assessed using common measures, such as accuracy, precision, recall, and F1-score. The findings show that, across all evaluation parameters, tree-based models Random Forest in particular have done noticeably better than other algorithms.

Table 1 Evaluation Metrics for Various Models

MODEL	Accuracy in %	Precision in %	Recall in %	F1-Score in %
Decision Tree	99.65	99.62	99.69	99.66
Logistic Regression	51.86	51.95	51.06	51.50
Random Forest	99.94	99.91	99.97	99.94
LSTM	99.72	99.99	99.44	99.72

4.2 Discussion

With 99.94% accuracy, 99.91% precision, 99.97% recall, and an F1-score of 99.94%, Random Forest performed best, demonstrating its ability to handle high-dimensional and unbalanced data. Leveraging

its strength with sequential data, LSTM came in second with 99.72% accuracy, 99.99% precision, 99.44% recall, and an F1-score of 99.72%. Decision Tree achieved 99.65% accuracy, with slightly lower metrics than Random Forest due to overfitting. Logistic Regression performed poorly with 51.86% accuracy and an F1-score of 51.50%, failing to capture nonlinear patterns [15]. The F1-score confirmed Random Forest as the top performer, followed by LSTM and Decision Tree. Logistic Regression's lower score highlighted its inadequacy for fraud detection, emphasizing the superiority of tree-based and neural network models. Figure 3, Figure 4, Figure 5 and Figure 6 depicts the F1-Score for each model.

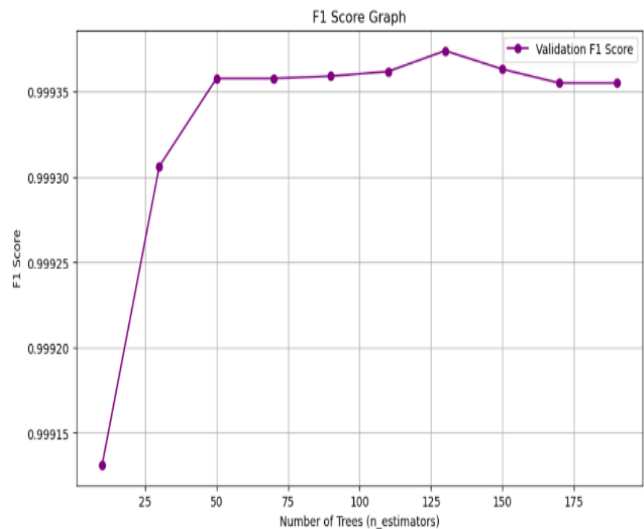


Figure 5 F1-Score: Random Forest.

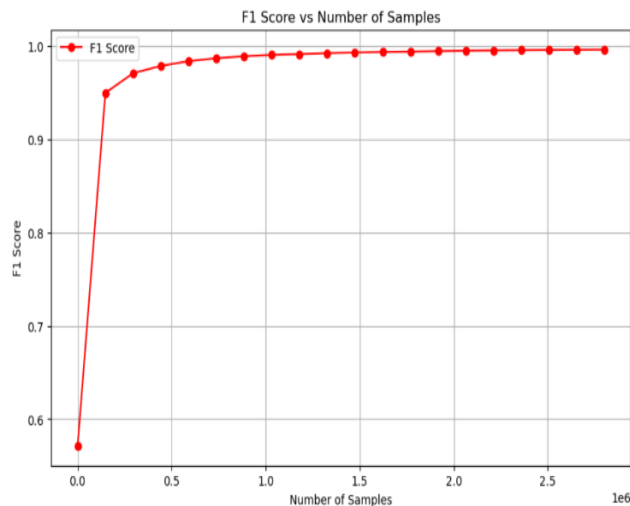


Figure 3 F1-Score: Decision Tree

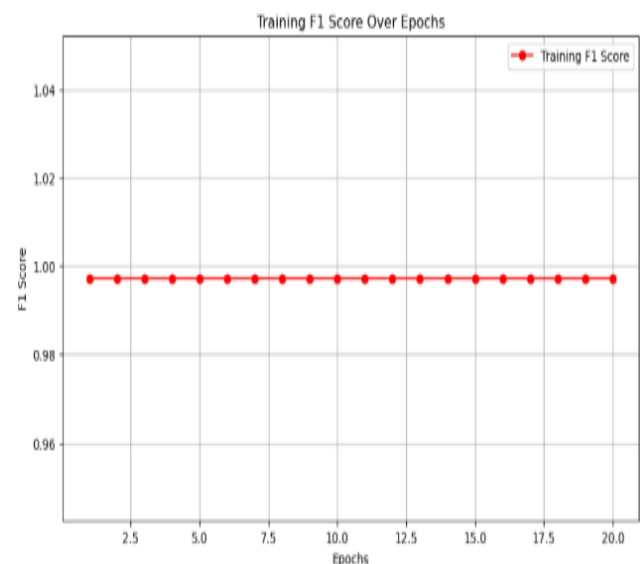


Figure 6 F1-Score: LSTM

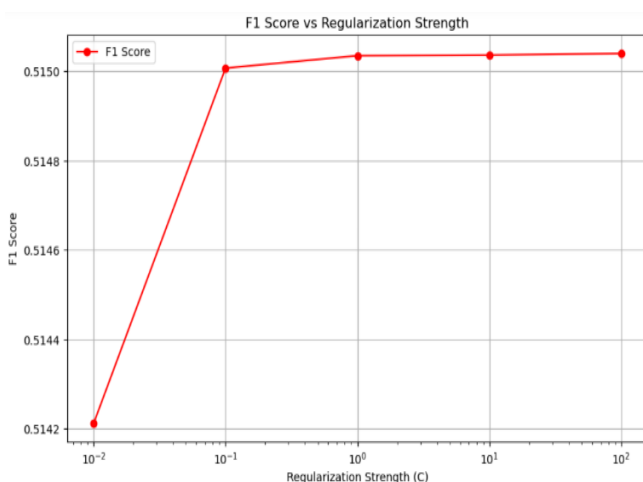


Figure 4 F1-Score: Logistic Regression

The confusion matrices highlighted the model's strengths and weaknesses in classifying legitimate and fraudulent transactions. The Decision Tree was observed to have overfitted, while Logistic Regression was hindered by limitations in handling non-linear patterns and class imbalance. Random Forest was demonstrated to have performed better, effectively balancing precision and recall, while LSTM was noted for having captured sequential dependencies with strong precision. Figure 7, Figure 8, Figure 9, and Figure 10 depicts the confusion matrix for each model.

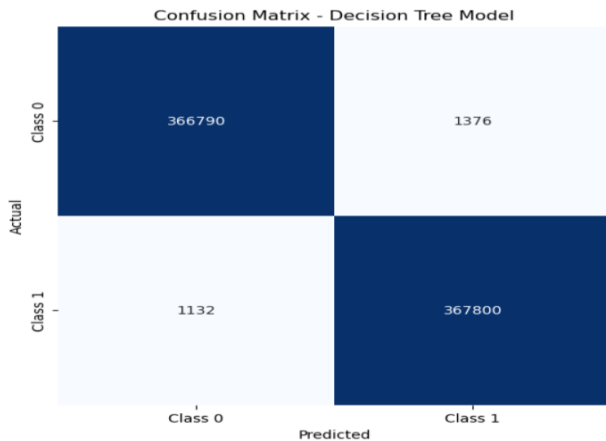


Figure 7 Confusion Matrix: Decision Tree

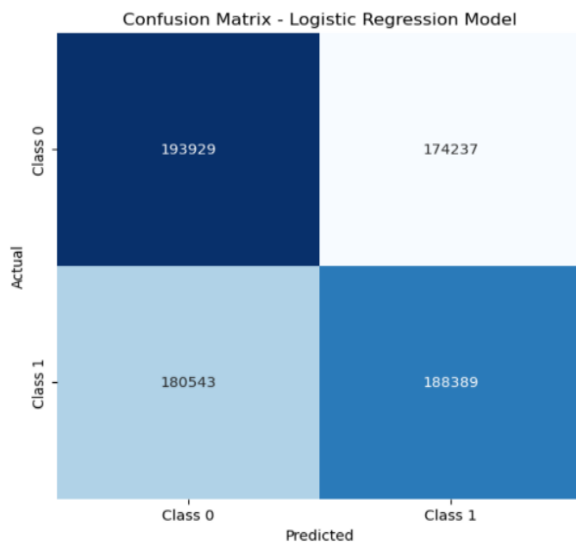


Figure 8 Confusion Matrix: Logistic Regression

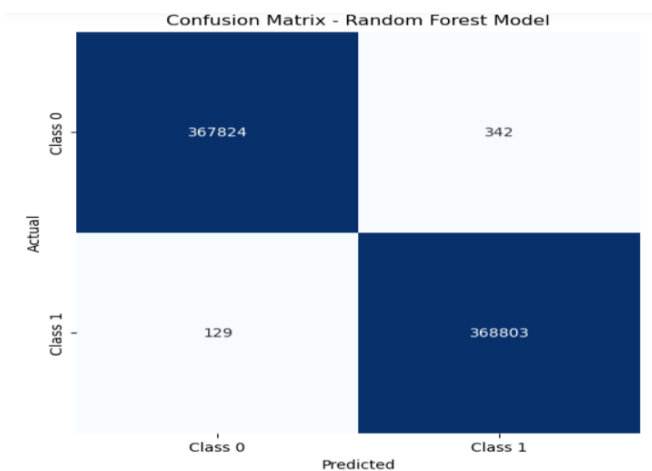


Figure 9 Confusion Matrix: Random Forest

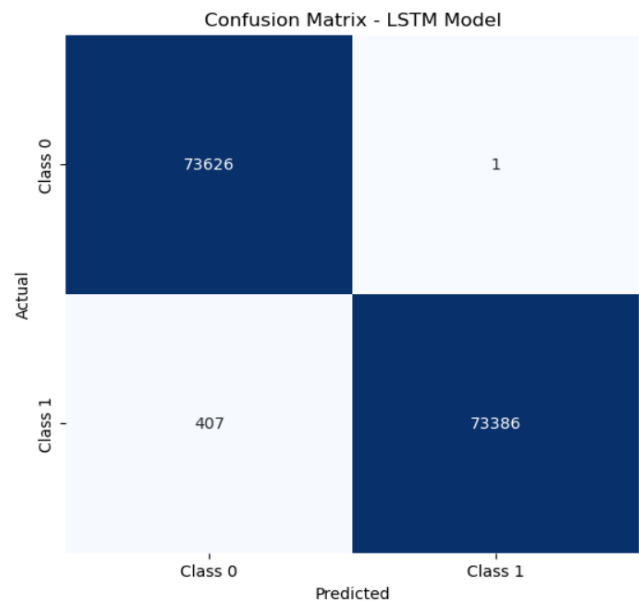


Figure 10 Confusion Matrix: LSTM

The Receiver Operating Characteristic (ROC) curves highlighted the models' ability to distinguish between legitimate and fraudulent transactions. Decision Tree, Random Forest, and LSTM demonstrated near-perfect curves, showcasing their superior performance, while Logistic Regression exhibited a weaker curve due to its limitations in modeling non-linear patterns. Figure 11, Figure 12, Figure 13 and Figure 14 depicts ROC curve for each model.

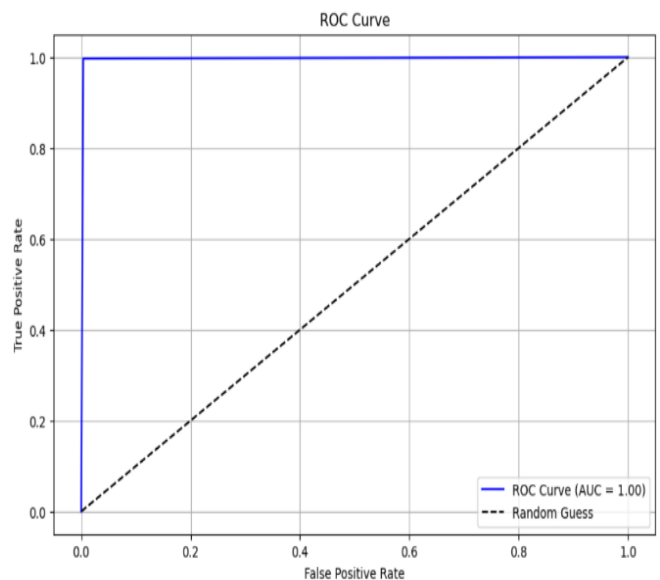


Figure 11 Roc Curve: Decision Tree

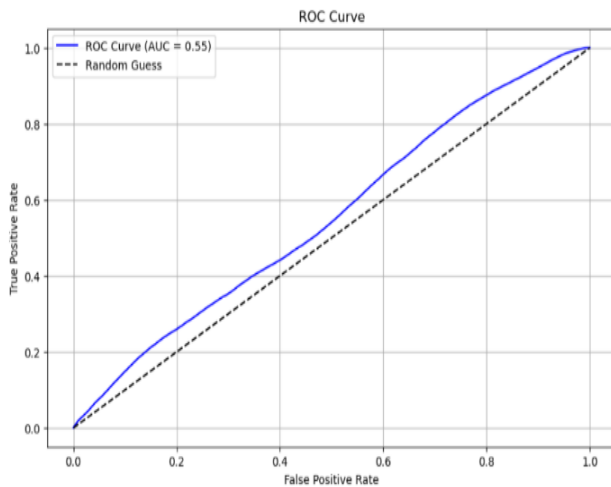


Figure 12 Roc Curve: Logistic Regression

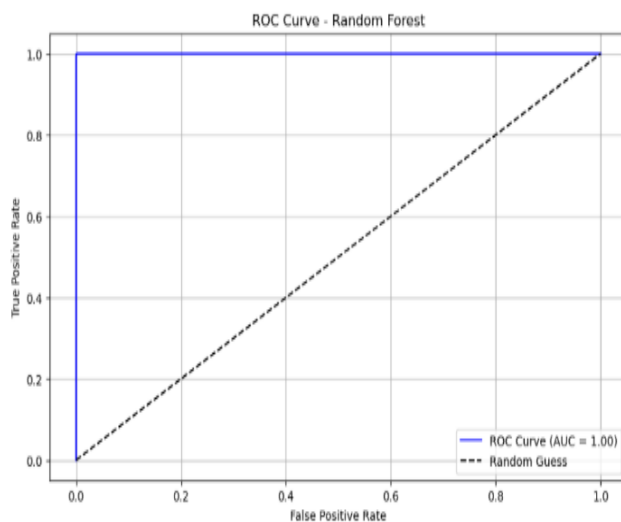


Figure 13 Roc Curve: Random Forest

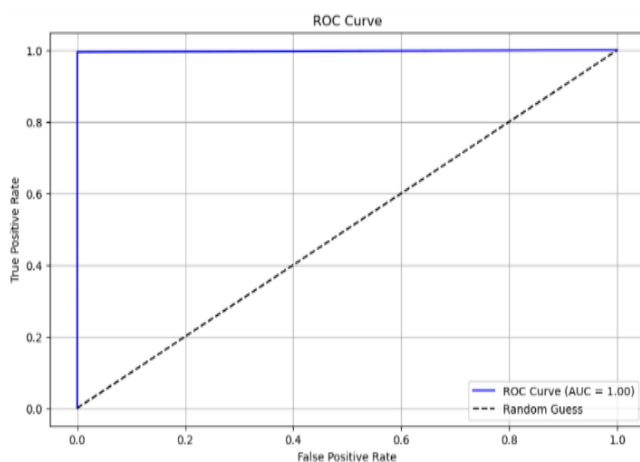


Figure 14 Roc Curve: LSTM

Conclusion

This study highlights the importance of carefully selecting features to enhance fraud detection models. Geographic location features, including 'zip', 'lat', 'long', 'merch_lat', and 'merch_long', were critical in identifying suspicious patterns by capturing spatial relationships. The 'distance' feature, which measures the separation between transaction and merchant locations, helped in pinpointing outliers that may indicate potential fraud. Temporal features such as 'unix_time', representing transaction timestamps, proved vital for detecting time-based patterns. These features were particularly effective when used with models like LSTM, which excel in handling sequential data. Incorporating contextual information, such as transaction categories, user profiles, and device identifiers, could further improve model performance. Exploring advanced ensemble techniques, hybrid models, and transfer learning approaches may enhance the model's ability to generalize across different fraud patterns. Integrating real-time processing capabilities and adaptive models that learn from emerging fraud behaviors will also be essential for maintaining detection effectiveness. Additionally, future research could focus on privacy-preserving methods, such as federated learning, to achieve a balance between fraud detection and data confidentiality.

References

- [1]. E. Esenogho, I. D. Mienye, T. G. Swart, K. Aruleba, and G. Obaido, "A Neural Network Ensemble with Feature Engineering for Improved Credit Card Fraud Detection," *IEEE Access*, vol. 10, pp. 16400–16407, Jan. 31, 2022.
- [2]. S. Makki, Z. Assaghir, Y. Taher, R. Haque, M.-S. Hacid, and H. Zeineddine, "An Experimental Study With Imbalanced Classification Approaches for Credit Card Fraud Detection," *IEEE Access*, vol. 7, pp. 69200–69209, Jul. 8, 2019.
- [3]. E. Ileberi, Y. Sun, and Z. Wang, "A machine learning-based credit card fraud detection using the GA algorithm for feature selection," *Journal of Big Data*, vol. 9, Feb. 24, 2022.
- [4]. A. Cherif, A. Badhib, H. Ammar, S. Alshehri,

- M. Kalkatawi, and A. Imine, "Credit Card Fraud Detection in the Era of Disruptive Technologies: A Systematic Review," *Journal of King Saud University - Computer and Information Sciences*, vol. 35, no. 1, pp. 145-174, Jan. 2023.
- [5]. I. D. Mienye and N. Jere, "Deep Learning for Credit Card Fraud Detection: A Review of Algorithms, Challenges, and Solutions," *IEEE Access*, vol. 12, July 2024.
- [6]. V. N. Dornadula and G. Sa, "Credit Card Fraud Detection using Machine Learning Algorithms," *Procedia Computer Science*, vol. 165, pp. 631-641, 2019.
- [7]. Kuldeep Randhawa, Chu Kiong Loo, Manjeevan Seera, Chee Peng Lim, and Asoke K. Nandi, "Credit Card Fraud Detection Using AdaBoost and Majority Voting," vol.6, pp. 14277-14284, January 3, 2018.
- [8]. R. S. M. Carrasco and M.-Á. Sicilia-Urban, "Evaluation of Deep Neural Networks for Reduction of Credit Card Fraud Alerts" *IEEE Access*, vol. 8, pp. 186421–186432, Sep. 2020.
- [9]. I. Benchaji, S. Douzi, B. El Ouahidi, and J. Jaafari, "Enhanced credit card fraud detection based on attention mechanism and LSTM deep model," *Journal of Big Data*, vol. 8, no. 1, art. 151, 2021.
- [10]. Abadhan Ranganath, Manas Ranjan Senapati, Pradip Kumar Sahu "A novel pixel range calculation technique for texture classification" *Multimedia Tools and Applications*, vol. 81, pp. 17639-17667, 7 March 2022.
- [11]. F. K. Alarfaj, I. Malik, H. U. Khan, N. Almusallam, M. Ramzan, and M. Ahmed, "Credit Card Fraud Detection Using State-of-the-Art Machine Learning and Deep Learning Algorithms," *IEEE Access*, vol. 10, pp. 32429–32442, Apr. 2022.
- [12]. M. Ma, C. Liu, R. Wei, B. Liang, and J. Dai, "Predicting machine's performance record using the stacked long short-term memory (LSTM) neural networks," *Journal of Applied Clinical Medical Physics*, Vol 23, Issue 3, March 2022.
- [13]. H. Najadat, O. Altit, A. Abu Aqouleh, and M. Younes, "Credit Card Fraud Detection Based on Machine and Deep Learning," , Apr. 7-9, 2020.
- [14]. S. Warghade, S. Desai, and V. Patil, "Credit Card Fraud Detection from Imbalanced Dataset Using Machine Learning Algorithm," *Int. J. Comput. Trends Technol.*, vol. 68, no. 3, pp. 21- 27, Mar. 2020.
- [15]. R. Lima and A. C. M. Pereira, "Feature selection approaches to fraud detection in e-payment systems," in *Lecture Notes in Business Information Processing*, vol. 287, pp. 90-103, Springer, Feb. 2017.