



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 13 Issue: VII Month of publication: July 2025

DOI: <https://doi.org/10.22214/ijraset.2025.73150>

www.ijraset.com

Call:  08813907089

E-mail ID: ijraset@gmail.com

Comparative Study of Centralized Vs Federated Learning for Credit Card Fraud Detection

Koppula Dennis¹, K. Mythri Sridevi²

¹M.Tech, CSE Departmentt, UCEK, JNTU Kakinada, Andhra Pradesh, India

²Assistant Professor (c), CSE Department, UCEK, JNTUK Kakinada, Andhra Pradesh, India

Abstract: *The financial industry continues to face significant challenges due to credit card fraud, which causes significant financial losses worldwide. Traditional machine learning techniques usually depend on centralized data aggregation, raising issues with inter-institutional data sharing, user privacy, and regulatory compliance. This paper presents a privacy-aware framework that uses Federated Learning (FL) to identify fraudulent transactions in order to overcome these constraints. This setup maintains anonymity by having several simulated financial organizations (clients) train models privately on their own private datasets without disclosing raw data. To address the class imbalance common in fraud detection, each client employs a Random Forest classifier in combination with the Synthetic Minority Over-sampling Technique (SMOTE). Using the Flower architecture, the federated system is constructed and assessed throughout a number of communication cycles. According to the results, the FL- based strategy preserves data privacy while achieving accuracy on par with centralized methods that use models like Random Forest, Decision Tree, and Logistic Regression. The feasibility of federated learning for safe and scalable fraud detection in dispersed situations is highlighted by this study.*

Keywords: *credit card fraud, federated learning, random forest, SMOTE, distributed machine learning, Flower framework.*

I. INTRODUCTION

Global financial systems are seriously at risk from credit card theft, which has become much more likely due to the rise in digital financial activity. Strong fraud detection systems must be put in place by financial institutions immediately in order to stop illegal access and guarantee safe transactions. The dynamic nature of user activity and the increasing sophistication of fraudulent schemes make fraud detection an ever-changing challenge.

The stark class imbalance in transactional datasets, where legitimate transactions greatly outnumber fraudulent ones, is one of the main obstacles to credit card fraud detection. In traditional machine learning models, this imbalance frequently results in biased learning, which impairs the minority class's (fraud) detection performance. Furthermore, centralized data aggregation efforts are made more difficult by the fact that real-world data gathered by banks and financial service providers is frequently kept in separate silos.

Centralized learning, which involves combining transaction data from several sources onto a single server for model training, is a key component of traditional fraud detection techniques. Although this paradigm can provide very effective models, it presents serious issues with data sovereignty, security threats, privacy, and regulatory compliance. In order to preserve trust and comply with the law, financial institutions need to establish privacy-aware practices in light of strict privacy rules like the General Data Protection Regulation (GDPR).

Decentralized training of machine learning models among numerous clients is made possible by Federated Learning (FL), a promising paradigm that eliminates the need to move raw data to a central repository. Individual clients, like banks, use their own proprietary datasets to train models in this system. They only exchange encrypted model updates, such as weights or gradients, with a centralized server that aggregates the results. This method reduces the chance of data breaches, improves security, and guarantees data locality.

Federated Learning provides a scalable and privacy-preserving approach, as shown in Figure 1, which makes it ideal for fraud detection in delicate financial areas.

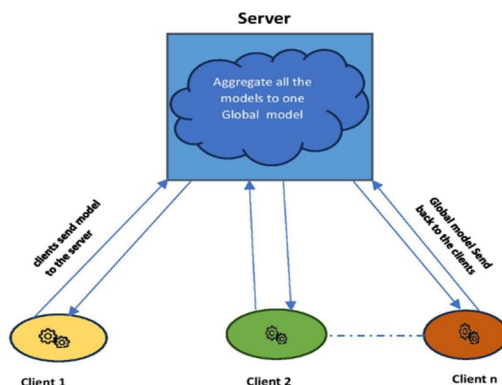


FIG 1: Federated Learning Architecture

This study proposes a Federated Learning-based system for detecting credit card fraud in which several clients work together to build local models using the Random Forest algorithm. To address the client's class disparity level, the Synthetic Minority Over-sampling Technique (SMOTE) is used, which improves the model's capacity to learn about minority fraud situations. The Flower framework, which facilitates adaptable, safe, and effective deployment in actual distributed contexts, is used to implement the suggested system. Finally, by using the Random Forest algorithm, this work provides a practical implementation of Federated Learning for credit card fraud detection. It offers a novel approach by applying the Synthetic Minority Over-sampling Technique (SMOTE) at the client level to lessen class imbalance and enhance the detection of fraudulent transactions. By comparing the performance of the proposed federated model to that of traditional centralized machine learning models, key metrics such as accuracy, precision, recall, and F1-score are evaluated in detail. Because it offers better privacy preservation while attaining equivalent accuracy, the results demonstrate that the federated approach is a promising choice for safe and scalable implementation in financial applications. Experimental findings validate the suggested FL approach's suitability in privacy-sensitive situations by demonstrating that it strikes a balance between detection accuracy and data privacy. areas such as digital finance and banking. Building reliable, scalable, and legally acceptable fraud detection systems that meet the demands of contemporary financial ecosystems is made possible by this research.

II. RELATED WORKS

Given the global increase in electronic transactions and the resulting increased susceptibility of systems to fraudulent activity, credit card fraud detection is a crucial area of research in financial cybersecurity. The use of Machine Learning (ML) algorithms to detect illegitimate or suspect transactions has been the subject of numerous studies. Among these, supervised learning methods like Random Forest (RF) have continuously shown excellent robustness and classification accuracy. Due to its ensemble structure, Random Forest performs very well with high-dimensional or noisy data, and its effectiveness is further enhanced when hyperparameters are appropriately adjusted for use cases including fraud detection. It is frequently compared to other well-known classifiers, such as XGBoost, Decision Trees, Support Vector Machines (SVM), K-Nearest Neighbors (KNN), and Naive Bayes, all of which have demonstrated differing levels of performance depending on on data characteristics [1][7][9][18][19].

Ensemble learning techniques have been frequently used in the literature to further improve performance. These tactics include hybrid pipelines, which combine algorithms like AdaBoost, Logistic Regression, and Multi-Layer Perceptrons (MLPs) to capitalize on the capabilities of varied learners, stacked models, and soft voting classifiers. In unbalanced and noisy datasets, this integration improves model stability, fraud detection rates, and generalizability [10][15][16]. Furthermore, comprehensive data preprocessing and feature engineering which includes methods like normalization, outlier removal, feature modification, and correlation analysis are crucial for the majority of high-performing systems. These actions are thought to be fundamental for enhancing the accuracy and convergence of ML models [11].

Sometimes, when paired with the right feature selection or used on domain-specific datasets, relatively basic models like SVM or Logistic Regression can perform better than sophisticated classifiers [11][12][13]. Additionally, these models have the benefit of being interpretable, which is crucial for financial system auditability and regulatory compliance. The class imbalance issue is a well-known topic that is frequently brought up in the literature. Because fraudulent transactions in real-world transactional datasets are so uncommon in comparison to valid ones, biased models frequently incorrectly categorize fraud situations. SMOTE (Synthetic Minority Over-sampling Technique) [1] and ADASYN (Adaptive Synthetic Sampling)

[15] are two common oversampling strategies used to counteract this. These techniques balance the dataset and increase model sensitivity by creating artificial samples of the minority class. More sophisticated techniques like Relocating Majority Decision Boundary (RMDD) [16] and Generative Adversarial Networks (GANs) [20] have surfaced to enhance anomaly detection performance and generate more realistic fraudulent cases, going beyond conventional oversampling.

Feature selection, which lowers computational complexity while maintaining predictive power, is another important topic of attention. To find the most pertinent subset of characteristics, methods based on metaheuristic algorithms such as Particle Swarm Optimization (PSO) [14] and Oppositional Cat Swarm Optimization (OCSO) [8] are employed, improving model accuracy and cutting down on training time. Federated Learning (FL), a privacy-preserving machine learning paradigm, has gained popularity as a result of growing awareness and concern over data privacy, particularly in light of laws like the General Data Protection Regulation (GDPR). With FL, several organizations (like banks or financial platforms) can work together to train a model without sharing raw data. Instead, utilizing protocols such as Federated Averaging (FedAvg), each client trains locally and only shares encrypted model updates with a central aggregator [2][3]. FL provides a workable answer to data sovereignty problems, allowing for decentralized training while maintaining user privacy.

Enhancing FL's flexibility in response to non-IID data, system heterogeneity, and communication limitations has been the main emphasis of recent developments. FedAvg-DWA (Distance-based Weighted Aggregation) and FedGAT-DCNN are noteworthy additions that use dilated convolutions and attention techniques to better accommodate diverse clientele [5]. Comparative studies have shown that FL models greatly improve privacy and lower data transport cost while frequently achieving accuracy levels comparable to centralized models [4][6].

Furthermore, hybrid approaches that combine supervised and unsupervised learning are becoming more popular due to their capacity to identify both established and new fraud tendencies. These include employing Support Vector Data Description (SVDD) to uncover anomalies in one-class classification issues [17] and integrating Random Forest classifiers with behavioral biometrics for improved user authentication [18]. Like hybrid or anomaly detection methods are particularly important during high-risk times, like as the COVID-19 pandemic, when transaction patterns may suddenly change and adaptive models are required [14].

GANs are being utilized more and more for feature augmentation and multi-feature fusion in addition to fraud detection, which enables models to learn more thorough and complex data representations [20]. Additionally, domain-specific fraud scenarios have led to the emergence of deep learning architectures that perform better in extremely specialized circumstances [8]. However, some difficulties still exist. In FL, communication overhead, client instability, and unpredictability in local data quality might impair performance. Additionally, in highly regulated industries like finance, model interpretability especially for ensemble or deep learning models—remains an issue. [13], [16], [19], and [20]. Models that can generalize across datasets from various institutions, regions, and time periods are also continuously needed [7][12][15].

The literature highlights the necessity of future research concentrating on creating scalable, privacy-aware, and real-time fraud detection systems in light of these challenges. This entails looking into hybrid machine learning architectures, making use of diverse and enhanced datasets, and making sure algorithms comply with legal and ethical requirements. [4, 7, 13, 14, 18, 19].

III. METHODOLOGY

A. Dataset Description

The study makes use of the publicly accessible creditcard.csv dataset, which consists of anonymized credit card transactions that a European bank collected over two days in September 2013. 284,807 transaction records make up this dataset; 492 of those transactions are classified as fraudulent, while the remaining 284,315 are classified as valid. The fact that fraudulent transactions only make up around 0.17% of the total dataset leads to a notable class imbalance.

Thirty features represent each transaction. Twenty-eight of these, designated V1 through V28, are anonymised principle components obtained by principle Component Analysis (PCA). Time and Amount are the final two features. The Time feature shows how many seconds have passed between a transaction and the dataset's first transaction. The transaction's monetary value is indicated via the Amount feature. A genuine transaction is indicated by a value of 0 for the target variable, Class, and a fraudulent transaction is indicated by a value of 1. In total, there are two integer attributes (Time and Class) and 29 floating-point attributes in the dataset. This dataset, which is made publically available via the Kaggle platform, is widely acknowledged as a standard benchmark in credit card fraud detection research [21].

Class Label	Transaction Type	Number of Samples
0	Legitimate Transaction	284,315
1	Fraudulent Transaction	492
Total		284,807

TABLE 1: Dataset Description

The Synthetic Minority Over-sampling Technique (SMOTE) [22] was used to rectify the dataset's class imbalance and increase the model's sensitivity to infrequent fraudulent occurrences. SMOTE generates new samples for the minority class by interpolating between existing examples of the minority class, which improves the balance of the dataset and the accuracy of the model in detecting fraudulent transactions.

The dataset went through a number of preparation stages before the model was trained. To make sure that all values fell within a similar range, the Time and Amount features were scaled using Min-Max normalization. Larger magnitude features can not substantially affect the model's learning process because to this normalization. In order to preserve the data's quality and integrity, the dataset was further examined for duplicate entries and missing values, which were eliminated. The dataset was horizontally divided into subsets, each of which was assigned to a distinct client in order to replicate a realistic federated learning environment. By ensuring that no raw data was exchanged between clients, this partitioning complied with federated learning's privacy preservation guidelines.

Preprocessing Step	Purpose
Missing Value Check	Ensure completeness and accuracy of data
Duplicate Removal	Eliminate repeated entries that could introduce bias during training
Min-Max Scaling	Normalize Time and Amount features for balanced model input
SMOTE Oversampling	Address class imbalance and enhance fraud detection capability

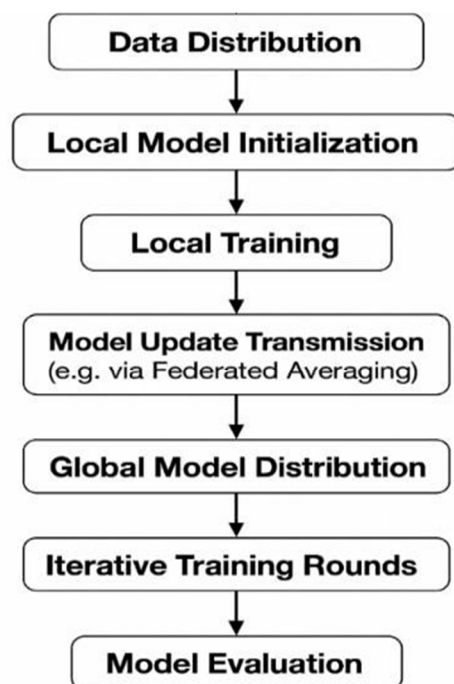
TABLE 2: Preprocessing Workflow

B. Federated Learning Setup

This study used a Federated Learning (FL) framework with the Flower library (version 1.18.0) to enable collaborative model training while safeguarding the privacy of sensitive financial data. Two client nodes, each representing a separate financial institution, were intended to be a part of the experimental configuration. The entire dataset was horizontally partitioned in a non-identically distributed (non-IID) fashion, with each client having access to a distinct and private piece of the transaction data.

A Random Forest classifier was locally trained on each client using the corresponding dataset. Three rounds of communication were used to coordinate the training process with a central server. The server did not get raw data from the local models that were trained on the clients. Rather, only the parameters of the model were sent. The Federated Averaging (FedAvg) technique was then used to centrally aggregate these parameters [23]. An iterative approach was then used to disseminate the aggregated global model to the clients for additional updates. By keeping data on the client side, this FL method not only protected data privacy but also made it possible to profit from cooperative model training across dispersed data sources.

C. Federated Learning Process Flow



- 1) **Data Distribution:** The credit card transaction dataset was partitioned horizontally among two simulated clients, representing distinct financial institutions. Each client retained exclusive access to its respective subset of the dataset, thereby maintaining complete data privacy. This horizontal distribution strategy ensured that no raw transaction data was shared or centralized, effectively mimicking a real-world scenario where data resides within isolated systems of different banks or organizations.
- 2) **Local Model Initialization:** Each client independently set up a local Random Forest Classifier with identical hyperparameters and structural configurations. This consistency in model setup across all clients ensured standardized training processes. However, because each client used only its own data, the models adapted uniquely to the specific data patterns present at their respective sites.
- 3) **Local Training with Class Imbalance Handling:** Individual datasets were used for training at each client node. The SMOTE (Synthetic Minority Oversampling Technique) approach was used to rectify the class imbalance commonly observed in fraud detection datasets. SMOTE improved the balance before training by artificially increasing the amount of fraud cases in the dataset. This action improved the model's capacity for generalization. To protect the privacy of the data, all training was done locally.
- 4) **Secure Model Update Transmission:** Once training was completed locally, each client shared only the model updates such as decision paths, node metrics, and voting outputs with a central federated server. No raw data or sensitive transaction information was transmitted, which significantly minimized privacy risks and potential data leaks.
- 5) **Global Model Aggregation via FedAvg:** Using the Federated Averaging (FedAvg) method, the central server gathered and combined the model changes from every client. FedAvg developed a new global model by calculating a weighted average of the customer parameters. This approach preserved the privacy of each client's data while utilizing the collective knowledge of distributed models.
- 6) **Global Model Distribution for Continuous Learning:** Once the global model was updated, it was sent back to all clients. This sharing mechanism allowed clients to gain insights derived from other institutions' data without direct data access. Each client then used the global model as the base for its next local training cycle, supporting ongoing model improvement through shared learning.
- 7) **Iterative Communication Rounds:** The federated training process was repeated over three communication cycles. In each round, local training, update transmission, aggregation, and redistribution were performed. This repetition helped the model evolve steadily, improving its ability to detect fraudulent activity by learning from the variety of decentralized datasets.

- 8) **Comprehensive Model Evaluation:** A different centralized dataset that had not been utilized for training was used to test the final version of the global model after all training cycles were finished. Accuracy, precision, recall, F1-score, and area under the curve (AUC) were among the evaluation criteria. These metrics provided a comprehensive assessment of the model's fraud detection capabilities while minimizing false positives and false negatives.

D. Algorithms Used

1) Logistic Regression (LR)

One kind of supervised machine learning method that is typically employed for binary classification is logistic regression. It models the relationship between a dependent variable and one or more independent factors using a logistic (sigmoid) function to forecast the likelihood of an outcome, such as whether a transaction is fraudulent (1) or lawful (0).

The log-odds of the event are expressed by the model as follows:

$$\ln\left(\frac{p}{1-p}\right) = \alpha_0 + \alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_n x_n$$

Where:

- p is the estimated probability of fraud,
- x_1 to x_n are the input variables,
- α_0 is the intercept term,
- α_i are the coefficients for each input variable.

The weighted sum of the inputs is converted by the sigmoid function into a probability score that ranges from 0 to 1, signifying the likelihood of fraudulent conduct. The benchmark model for centralized fraud detection in this work is logistic regression.

2) Decision Tree (DT)

For classification tasks, a decision tree is a structure that resembles a flowchart. It creates a tree-like structure by dividing the dataset into subgroups according to particular feature values. Branches show potential outcomes, leaf nodes assign final class labels like fraud or non-fraud, and each internal node reflects a judgment based on a feature.

It makes use of assessment metrics like:

Entropy:

$$H(X) = -\sum p(x_i) \log_2 p(x_i)$$

Gini Index:

$$G(X) = 1 - \sum p(x_i)^2$$

Information Gain

$$IG(X, Y) = E(X) - E(X | Y)$$

These help determine the best feature for splitting the data. Decision Trees require minimal data preprocessing and are effective for interpretable classification tasks.

3) Random Forest (RF)

In order to increase overall prediction accuracy, Random Forest, an ensemble learning technique, builds several decision trees and aggregates their results. A randomly chosen subset of features is assessed at each decision node, and each tree is trained using a distinct bootstrap sample of the dataset. This procedure improves resilience and adds variance among trees.

To make a prediction for input X , the model uses a majority voting mechanism across all decision tree

$$(X) = \arg \max\{\text{Fraud votes}, \text{Not Fraud votes}\}$$

Random Forest is known for reducing overfitting and increasing the generalization ability of models. It is especially effective on datasets with class imbalance such as fraud detection scenarios and is further enhanced when used in combination with techniques like SMOTE.

In this study, the Random Forest algorithm is implemented in both centralized and federated environments.

4) Federated Learning Implementation

A distributed machine learning technique called federated learning (FL) enables several clients, including banks and other financial institutions, to work together to train a common model while maintaining the privacy of their local data. Every client in the study's design maintains access to its own credit card transaction dataset and participates in model

training without sending raw data. This approach, referred to as cross-silo FL, is especially beneficial in sectors that demand high privacy and data governance. It ensures compliance and protects sensitive data while enabling shared learning across institutions.

To replicate a realistic multi-institution environment, the dataset was partitioned

horizontally into distinct, non-identically distributed (non-IID) segments. Each client node trained a local Random Forest classifier independently using only its local dataset. Despite identical model configurations, the learning outcomes differed across clients due to variations in data distribution. At each client, the Synthetic Minority Over-sampling Technique (SMOTE) was used locally to solve the class imbalance that is frequently present in fraud detection jobs. In order to increase model sensitivity, this made sure that infrequent fraudulent transactions were more accurately reflected in the training data.

Once local training was complete, clients transmitted only their model parameters (e.g; node structures, feature importance scores, and voting probabilities) to a central server. No actual transaction data was ever shared, preserving the confidentiality of client-held records.

The Federated Averaging (FedAvg) method was used by the central server to aggregate the models. By using a weighted average of the model parameters that each client provides, this approach determines how many training samples each client possesses. The following is the aggregation formula:

$$\theta^{(t+1)} = \sum_{k=1}^K \frac{n_k}{n} \cdot \theta_k^{(t)}$$

Where:

- $\theta^{(t)}$ denotes the model parameters from client k at iteration t ,
- n_k is the quantity of training samples that client k possesses.
- n is the total number of training samples for each client that is taking part.

After updating the global model, it was sent back to each client so they could use the newly aggregated model as the starting point for local training. Three communication rounds of this iterative process local training, model update transmission, global aggregation, and redistribution were conducted.

After the last round of communication, an independent hold-out test dataset that had not been used during any training phase was used to validate the trained global model.

Accuracy, precision, recall, F1-score, and AUC were among the primary classification metrics used to evaluate the model's performance. These metrics collectively offered a thorough assessment of the model's capacity to identify fraudulent transactions while reducing the number of false positives and false negatives.

This federated setup not only maintained data privacy but also demonstrated that collaborative learning can achieve performance levels comparable to traditional centralized training, even in the presence of non-IID data distributions and class imbalance.

IV. RESULTS AND DISCUSSION

This section presents the results of both centralized and federated machine learning approaches applied to credit card fraud detection. The evaluation emphasizes classification performance, particularly the models' ability to identify fraudulent transactions, which are frequently underrepresented because of imbalanced datasets. Accuracy and F1-score are important evaluation metrics; for a more thorough class-wise examination, confusion matrices and additional metrics like precision and recall are utilized (not displayed here).

A. Centralized Model Evaluation

Three conventional machine learning models Logistic Regression (LR), Decision Tree (DT), and Random Forest (RF) were trained using the entire, preprocessed dataset in order to create a baseline for performance comparison. Evaluating these models' ability to detect fraudulent transactions in situations of class imbalance was the main goal.

B. Comparative Analysis

To provide a holistic performance comparison, both centralized and federated models were evaluated and their accuracy and F1-score are summarized in the table below. Check Table 3.

Model	Accuracy	F1-Score
Logistic Regression	0.9241	0.06
Decision Tree	0.9476	0.48
Random Forest	0.9495	0.80
Federated RF (Client 1)	0.9800	0.83
Federated RF (Client 2)	0.9800	0.82

TABLE 3: Centralized Vs Federated Model Accuracy And F1-Score COMPARISON

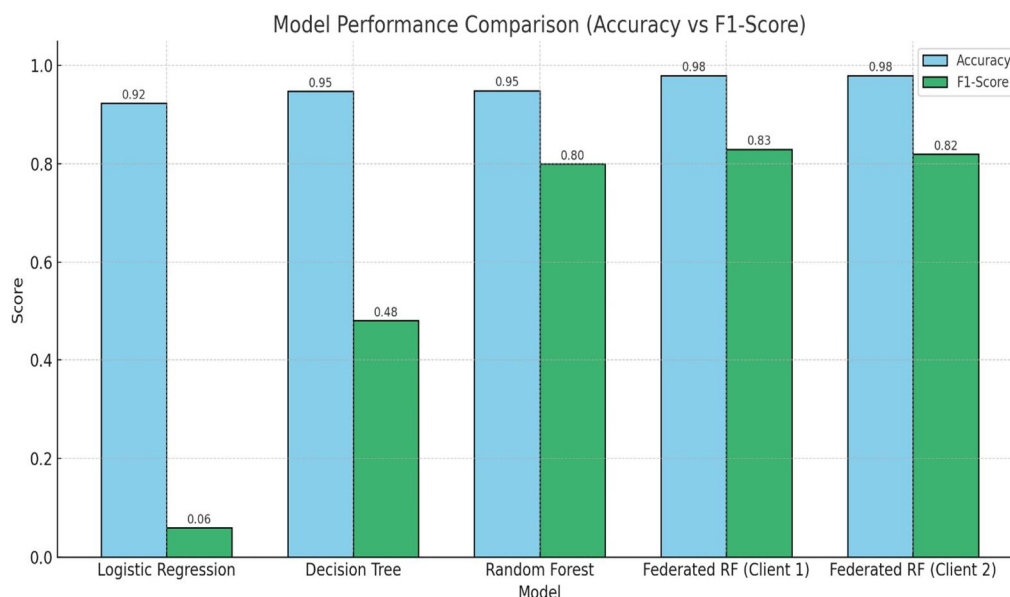


Fig 1: Combined Accuracy And F1-Score Comparison Between Centralized And Federated Learning Models

The results indicate that Federated Learning not only matches but in some cases slightly exceeds the performance of the centralized Random Forest model in terms of accuracy. Moreover, F1-scores are consistently higher in federated models compared to Logistic Regression and Decision Tree models, underscoring their improved ability to balance precision and recall.

These findings highlight the effectiveness of Federated Learning as a privacy-preserving, scalable alternative that can deliver near-centralized performance without requiring raw data to be shared making it particularly suitable for sensitive domains such as finance and healthcare.

C. Evaluation Metrics Employed

To evaluate the model's performance, we used several key metrics suitable for binary classification problems and imbalanced datasets.

Accuracy:

Measures how often the model's predictions are correct overall.

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

Precision:

Indicates the proportion of predicted fraud cases that are genuinely fraudulent.

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

Recall (Sensitivity or True Positive Rate):

Reflects the proportion of actual fraud cases that were successfully detected by the model.

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

F1-Score:

represents the precision and recall harmonic mean, which makes it very useful for evaluating models that were trained on unbalanced data.

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

These metrics collectively provide detailed insights into how well the model identifies fraudulent transactions and distinguishes them from legitimate ones crucial in skewed datasets.

D. Discussion

- The centralized model achieved slightly better results in terms of accuracy and recall. This is expected, as the centralized model had access to the entire dataset, enabling it to learn a broader representation of both fraud and non-fraud patterns.
- In contrast, the federated model ensured user data confidentiality by keeping training localized to client devices. Despite these privacy constraints, its performance remained closely aligned with the centralized model, proving its effectiveness in privacy-critical applications like finance and healthcare.
- One important aspect influencing performance differences was the non-IID (non-identically distributed) nature of data among clients. This reflects realistic conditions where institutions may have distinct user behavior and fraud patterns.
- Despite this data heterogeneity, the federated global model was able to nearly match the centralized model's performance, showcasing the robustness and adaptability of federated learning in distributed environments.
- The application of SMOTE (Synthetic Minority Over-sampling Technique) significantly improved recall and F1-scores in both centralized and federated settings. This highlights the importance of addressing class imbalance a common challenge in fraud detection tasks.

V. CONCLUSION AND FUTURE SCOPE

Using the Random Forest algorithm, this study compares centralized and federated learning strategies for detecting credit card fraud. Because centralized learning had access to the complete dataset, it was able to catch a wider range of transaction behavior, which contributed to its somewhat higher accuracy. Federated learning, on the other hand, maintained data privacy by storing user information locally on client devices. The federated technique yielded results comparable to the centralized model, even though it worked with decentralized and non-identically distributed (non-IID) data. By correcting class imbalance and raising the sensitivity of fraud detection, the use of SMOTE substantially enhanced the effectiveness of both strategies.

Federated learning has a lot of promise for use in actual fraud situations in the future. To better describe intricate transaction patterns, future studies can investigate the integration of sophisticated neural network designs like CNNs, DNNs, and transformers. FedProx and FedAvgM are examples of customized federated strategies that can be used to handle client data variability. Furthermore, security can be improved by implementing strong privacy- preserving technologies as secure multi-party computation (SMPC), homomorphic encryption, and differential privacy. The creation of a standardized, scalable, and privacy- conscious framework for fraud detection may also be aided by real-time data handling and cooperation between financial institutions.

REFERENCES

- [1] R. K. Chanda, P. K. Pagadala, C. K. Edukulla, S. S. Archana, S. Gurram, and S. R. Maram, "Enhancing credit card fraud prediction using decision trees, SMOTE, and hyper-tuned random forests: A comprehensive approach," in Proceedings of IEEE, 2023.
- [2] H. P. N., P. D. Rathika, and P. A., "Privacy preservation using federated learning for credit card transactions," in Proceedings of IEEE, 2023.
- [3] K. D'souza, S. Puthusseri, and A. G. Samuel, "Scalable federated learning for privacy-preserving credit card fraud detection," in Proceedings of IEEE International Carnahan Conference on Security Technology (ICCST), 2023.
- [4] S. Lynch, A. M. Abdelmoniem, and S. S. Gill, "Centralised and decentralised fraud detection approaches in federated learning: A performance analysis," in Applications of AI for Interdisciplinary Research, 1st ed., CRC Press, Jul. 2024, pp. 1–20, doi: 10.1201/9781003467199-18.
- [5] H. Zheng, "Federated learning-based credit card fraud detection: A comparative analysis of advanced machine learning models," in Proceedings of International Conference on Data Science, Advanced Algorithm and Intelligent Computing (DAI), vol. 70, 2025.

- [6] R. K. Chaudhary, R. Kumar, and N. Saxena, "A systematic review on federated learning system: A new paradigm to machine learning," *An International Journal* (Springer), vol. 67, Nov. 2024.
- [7] T. C. Tran and T. K. Dang, "Machine learning for prediction of imbalanced data: Credit fraud detection," in *Proceedings of IEEE*, 2021.
- [8] N. Prabhakaran and R. Nedunchelian, "Oppositional cat swarm optimization-based feature selection approach for credit card fraud detection," *Computational Intelligence and Neuroscience*, vol. 2023, Article ID 2693022, 2023, doi: 10.1155/2023/2693022.
- [9] S. R. Dammavalam and M. Mukheed, "Credit card fraud detection using machine learning," *International Journal of Advanced Engineering and Management (IJAEM)*, vol. 5, no. 1, pp. 1–6, Jan. 2023.
- [10] M. A. Mim, N. Majadi, and P. Mazamder, "A soft voting ensemble learning approach for credit card fraud detection," *Heliyon*, vol. 10, no. 3, p. e25466, Feb. 2024, doi: 10.1016/j.heliyon.2024.e25466.
- [11] L. Ali and A. Kasem, "Enhancement model to detect credit card fraud based on processing data," *European Modeling Studies Journal*, vol. 6, no. 5, pp. 1–8, 2022.
- [12] V. T. Gowda, "Credit card fraud detection using supervised and unsupervised learning," in *Proceedings of CMC, NCO, SOFT, CDKP, MLT, ICAITA*, 2021, doi: 10.5121/csit.2021.111107.
- [13] K. S. Srinivas, "Credit card fraud detection using supervised machine learning algorithms," *International Journal of Creative Research Thoughts (IJCRT)*, vol. 10, no. 9, Sep. 2022. [Online]. Available: <https://ijcrt.org/>
- [14] A. Mniai and K. Jebbari, "Credit card fraud detection by improved SVDD," in *Proceedings of World Congress on Engineering*, Jul. 2022.
- [15] S. Bagga, A. Goyal, N. Gupta, and A. Goyal, "Credit card fraud detection using pipelining and ensemble learning," Elsevier, 2020.
- [16] Y. Xie, A. Li, L. Gao, and Z. Liu, "A heterogeneous ensemble learning model based on data distribution for credit card fraud detection," *Wireless Communications and Mobile Computing*, vol. 2021, Article ID 2531210, 2021, doi: 10.1155/2021/2531210.
- [17] F. Carcillo, Y.-A. Le Borgne, O. Caelen, Y. Kessaci, F. Oble, and G. Bontempi, "Combining unsupervised and supervised learning in credit card fraud detection," *Information Sciences*, vol. 2019, pp. 245–257, doi: 10.1016/j.ins.2019.05.042.
- [18] S. R. Adapa, M. A. S. Nirob, S. Bhatt, M. Yerram, and A. P. Nivas, "Enhancing credit card fraud detection: A novel approach with random forest and behavioral biometrics," *International Journal for Research in Applied Science and Engineering Technology (IJRASET)*, vol. 12, no. 3, Mar. 2024.
- [19] J. Gao, Z. Zhou, J. Ai, B. Xia, and S. Coggeshall, "Predicting credit card transaction fraud using machine learning algorithms," *Journal of Intelligent Learning Systems and Applications*, vol. 11, no. 3, 2019, doi: 10.4236/jilsa.2019.113003.
- [20] Y. Xie, A. Li, B. Hu, L. Gao, and H. Tu, "A credit card fraud detection model based on multi-feature fusion and generative adversarial network," *Computer Materials & Continua*, vol. 2023, pp. 123–136, doi: 10.32604/cmc.2023.037039.
- [21] A. Dal Pozzolo, "Credit Card Fraud Detection Dataset," Kaggle. [Online]. Available: <https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud>
- [22] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002.
- [23] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Proceedings of the International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2017, pp. 1273–1282.
- [24] D. J. Beutel, T. Topal, A. Mathur, X. Qiu, T. Parcollet, and N. D. Lane, "Flower: A friendly federated learning research framework," *arXiv preprint, arXiv:2007.14390*, 2020.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)