

Enhancing Financial Security through the Integration of Machine Learning Models for Effective Fraud Detection in Transaction Systems

Jaya Krishna Modadugu¹, Ravi Teja Prabhala Venkata², Karthik Prabhala Venkata³

Abstract

Fraud detection (FD) in financial transactions involves identifying and preventing unauthorized or suspicious activities to safeguard financial systems and customer assets. However, traditional fraud detection approaches in financial transactions often fall short due to their reliance on predefined rules and static thresholds, which limits their capability to adapt to evolving fraud patterns and detect sophisticated or emerging threats effectively. To address the challenges in traditional FD methods, this paper proposes an advanced machine learning-based model, Enhancing Financial Security through the Integration of Machine Learning methods for Effective Fraud Detection in Transaction Systems (EFS-IML-EFD-TS). Initially, input data is gathered from the Financial Fraud Detection Dataset. The collected data is first pre-processed utilizing the Confidence Partitioning Sampling Filtering (CPSF) technique to handle missing values, remove duplicate records, and standardize feature scaling. The pre-processed data is further processed using the Exponential Distance Transform (EDT), which extracts discriminative features like transaction amount, time of day, and location. Then, the imbalanced data is balanced using Adaptive Support Vector-Borderline SMOTE (ASV-SMOTE), which generates high-quality synthetic samples near decision boundaries, reducing noise and improving minority class prediction. Then Interpretable Generalized Additive Neural Network (IGANN) is used to detect fraud and classify financial transactions as either genuine or fraudulent. The proposed EFS-IML-EFD-TS method achieves 98.5% precision, 98% accuracy, 97% recall, 97.5% F1-score, 0.91 MCC, a high AUC of 0.9636, low loss of 0.05, and the shortest computational time of 1.125 seconds, compared with existing methods such as Online payment fraud detection model utilizing machine learning techniques (OPFT-MLT-ANN), Financial Fraud Detection utilizing Value-at-Risk with Machine Learning in Skewed Data (FFD-MLSD-DNN), and Transparency and privacy: the role of explainable AI and federated learning in financial fraud detection (TP-AI-FFD-DNN).

Keywords: Fraud detection, Financial Security, Interpretable Generalized Additive Neural Network, Confidence Partitioning Sampling Filtering, Exponential Distance Transform, Transaction Systems.

Introduction

As online payment systems and e-commerce have grown in popularity, financial transactions have become more digital, leading to a rise in fraudulent activities [1]. Although incidents like unauthorized purchases and counterfeit cards represent a smaller portion of fraud cases, they result in a disproportionate share of financial losses [2]. In response to this growing concern, both government organizations and private businesses have significantly increased their investments in developing more robust fraud detection systems [3]. These systems are essential in identifying and preventing fraudulent transactions, thereby reducing financial losses [4]. Their efficacy is crucial for boosting security and building confidence in online financial transactions [5].

The data imbalance, where suspect transactions are far less likely than legitimate transactions, is the first of many challenges for fraud detection systems [6]. In addition to the issue of data imbalance, different misclassifications can come with different costs, with false positive instances leading to major financial fallout [7]. In addition, fraud detection systems must take into account temporal dependencies, meaning that they understand the relationships between the events based on time [8]. Temporal dynamics bring a challenge related to concept drift, meaning that the model must be updated regularly

¹ Software Engineer, Email: jayakrishna.modadugu@gmail.com, ORCID: 0009-0008-9086-6145

² Senior Manager, Software Engineer, Email: raviteja.prabhala@gmail.com, ORCID: 0009-0007-7265-212X

³ Senior Specialist, Project Management, Email: karthik030789@gmail.com, ORCID: 0009-0001-4977-9006

to remain as accurate as possible, since the types of fraud change over time [9]. Finally, the high dimensionality of data necessitates sophisticated tools to manage and analyze vast swathes of transaction data [10].

To mitigate the effects of imbalanced datasets, one solution is to modify the class weights during model training to elevate the significance of fraudulent transactions [11]. Cost sensitivity can be managed utilizing cost-sensitive learning strategies that guarantee the model penalizes misclassifications based on their financial implications [12]. Next, temporal dependencies can be accommodated by using time-series analytics that think about the sequence and timing of transactions [13]. Online learning techniques can combat concept drift, allowing the model to learn from recent data as trends replicate over time [14]. Finally, dimensionality can be lowered through feature engineering or strategies like principal component analysis (PCA), which can enable the retention of only the most significant features in fraud detection [15].

Literature Survey

Previous literature has presented a number of works that rely on the detection of fraud in financial transactions. Only a handful of them were highlighted here,

A. A. Almazroi and Nasir Ayub [16] have presented for processing financial transaction data, a novel artificial intelligence method called the ResNeXt-embedded Gated Recurrent Unit (GRU) model (RXT) was created. The growing threat of financial fraud, which presents significant risks to both consumers and financial institutions, was addressed methodically by AI technology. Data input and preprocessing were the first steps in the process, after which the SMOTE technique was used to address data imbalance. Use an ensemble AI technique for feature extraction that reveals important data patterns by combining autoencoders and ResNet (EARN). This method's drawback was that it might need a lot of time and computing power to train, particularly when dealing with big datasets.

A. U. Usmanet *al.* [17] have presented an approach that tackles the skewness and rarity of fraud cases in machine learning (ML) models in order to detect new bank account (NBA) fraud. Traditional methodologies neglect potential losses to more effectively explore fraud tendencies. The use of fraud as a worst-case scenario incorporates the use of value-at-risk (VaR) as a risk measure. VaR models potential losses as a skewed tail distribution and can be estimated using historical simulation. The BAF dataset was utilized with ML to classify risk-return characteristics which were based on VaR. A drawback to this method was it may not effectively account for new fraud trends, as it relies heavily on historical data and prescriptive risk measures.

T. Awosikaet *al.* [18] have presented a method to detect fraudulent transactions, which was a big concern for financial institutions. The approach focuses on the need for address the imbalance in transaction datasets since there were few instances of fraudulent transactions as compared to legitimate transactions, one limitation of this approach was the potential inability to detect new or changing fraud trends that may not be apparent in the training set.

Y. Cuiet *al.* [19] have presented a novel approach to adaptive and context-aware financial fraud detection that incorporates Graph Neural Networks (GNNs) and Reinforcement Learning (RL). The approach constructs a dynamic graph model for financial transactions, whereby transactions were nodes, whereas users and merchants were edges. The model introduced a novel GNN architecture, called Temporal-Spatial-Semantic Graph Convolution (TSSGC), to extract spatial relationships, temporal patterns, and semantic information from transaction data. The RL component was implemented as a Deep Q-Network (DQN), which allowed the model to minimize detection costs while having the capacity to adapt to changing patterns of fraud by adjusting the threshold for acceptable fraud detection and importance of features. This method's drawback was that it might take a lot of time and computing power to train, especially when dealing with big transaction datasets.

M. A. Talukderet *al.* [20] have presented a solution for detecting fraudulent transactions in financial institutions, specifically targeting credit card fraud. Early detection plays an important role in preventing further losses. The process involves thorough investigation of alerts; however, due to time constraints, only a limited number of warnings can be reviewed each day, which may impact the overall detection efficiency. A disadvantage of this approach was that the limited number of alerts that can be reviewed daily may result in delayed detection of some fraudulent transactions.

K. Singhet *al.* [21] have presented a safe, automated computer system for detecting financial transaction fraud. The purpose of this system was to safely process and examine transaction data, spot irregularities, and send out warnings about possible fraud. To guarantee data protection, it

integrates a number of security measures, including encryption, data obfuscation, and access controls. A disadvantage of this approach was that implementing multiple security mechanisms may increase processing time and system overhead.

Y. Tang and Z. Liu [22] have presented distributed knowledge distillation architecture for financial fraud detection. The method uses feed-forward neural networks to extract high-level relevant features after assigning weights to features using a multi-attention mechanism. Neural networks were then used to categorize financial fraud, improving detection accuracy, inference speed, and generalization ability for better decision-making in financial institutions. A disadvantage of this approach was that it may require significant computational resources to process and distill large datasets effectively. Table 1 presents the literature survey's summary.

Table 1: Literature survey's summary.

Ref	Algorithm	Advantage	Disadvantage
A. Almazroi and A. Nasir Ayub [16]	Artificial Neural Network (ANN)	Effectively detects complex fraud patterns.	High computational cost and training time.
A. U. Usman [17]	Deep Neural Network (DNN)	Enhances detection by assessing financial risk with VaR.	Limited in detecting new fraud patterns due to reliance on historical data.
T. Awosika [18]	Deep Neural Network (DNN)	Improves detection by handling data imbalance.	Struggles with new or evolving fraud patterns.
Y. Cui [19]	Graph Neural Network (GNN)	Adapts to evolving fraud with dynamic graph and RL.	High computational cost and training time.
M. Talukder [20]	Artificial Neural Network (ANN)	Enables early detection of credit card fraud to prevent further losses.	Limited alert reviews per day may delay detection of some fraud cases.
K. Singh [21]	Deep Neural Network (DNN)	Provides secure fraud detection.	Slower processing due to security overhead.
Y. Tang and Z. Liu [22]	Graph Neural Network (GNN)	Enhances accuracy, speed, and generalization in fraud detection.	Requires high computational resources.

Despite significant advancements in machine learning for financial FD, current methods still face challenges in performance, computational efficiency, and adapting to evolving fraud patterns. Techniques like ANN, GNN, and DNN are widely used, but often have very high computational complexity, and thus are less efficient when dealing with large data sets. Historical data is a key part of these models, and as such generally less effective when trying to identify fraud with a new technique against a known fraud model. Additionally, issues such as data imbalance and time to FD limit the system's overall effectiveness. All of these aspects highlight the need for a FD system with a more operational flexible, scalable, and efficient. The goal would be to make use of the data in order to identify new fraud patterns quickly and accurately without having to reload the model, and at a level of computational complexity that is not prohibitively expensive. It is with these considerations in mind that this work develops a FD system that quickly accounts for new activity, uses lower computational resources, and achieves adaptive learning to effectively detect financial fraud in large volume.

In this paper, the EFS-IML-EFD-TS method is proposed to improve fraud detection (FD) in financial transactions while overcoming the problems of traditional rule-based systems. It begins with data preprocessing using the CPSF technique to remove and handle null values, remove duplicates, and standardize features. ASV-SMOTE creates artificial samples in the vicinity of decision boundaries and provides methods to improve minority class detection and to solve data imbalance. EDT is used to

extract important features like transaction amount, time, and location. The IGANN model performs classification, correctly differentiating between legitimate and fraudulent transactions. The technique ensures effective fraud detection with few false positives by achieving high precision, recall, accuracy, F1-score, MCC, AUC, low loss, and quick computation time.

Important contribution of this research work is bridged below,

- In this research, Enhancing Financial Security through the Integration of Machine Learning Models for Effective Fraud Detection in Transaction Systems (EFS-IML-EFD-TS) is proposed.
- CPSF effectively pre-processes financial transaction data by addressing missing values, removing duplicates, and standardizing data for consistent and reliable input.
- IGANN interpretability into the classification process, making the method's decisions more transparent and suitable for financial environments that require accountability and compliance.
- The obtained results of the proposed EFS-IML-EFD-TS algorithm are compared with existing models such as OPFT-MLT, FFD-MLSD, and TP-AI-FFD, demonstrating superior performance across all metrics.

The balance paper is ordered as follows: Part 2 displays the proposed method, Part 3 displays the results and discussion, Part 5 concludes the paper.

Proposed Methodology

In this sector, the plan for EFS-IML-EFD-TS is outlined. Primary activities entail obtaining input data from a Financial FD dataset, containing a variety of transaction-related properties. The naive data is then preprocessed using a Confidence Partitioning Sampling Filter (CPSF) to treat missing values, to remove duplicates, and to normalize the factors of the input data. The data after preprocessing then goes to an Exponential Distance Transform (EDT) for feature extraction. The important features are the transaction amount, time, location, merchant category, and information about the recipient. The final dataset of features is then used as an input into an Interpretable Generalized Additive Neural Networks (IGANN) for classification of the transactions, ultimately classifying the transaction as genuine or fraudulent. IGANN uses the differentiation of features, informing back to a user understandable process of modelling complex relations present in the data, allowing for clear and accurate decision making. It aims to improve the efficiency of the process training the model, while improving the accuracy in detecting fraud, while maintaining the objective of using informative elements of the process, ideally equate to a lower overall accuracy output cost. This consolidated process is illustrated in Fig. 1. It indicates the proposed EFS-IML-EFD-TS model.

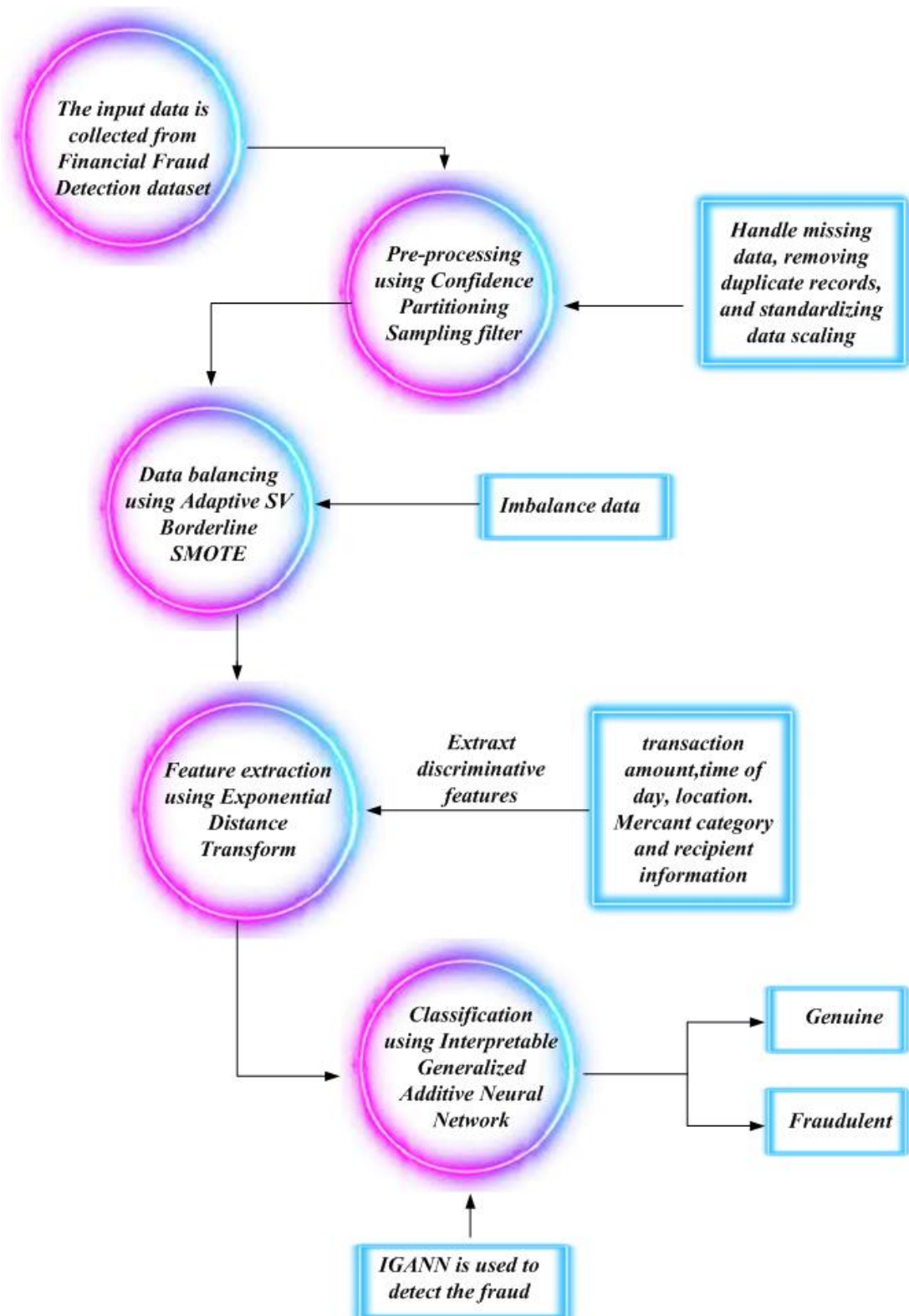


Fig 1: Block Diagram of the proposed EFS-IML-EFD-TS

Data Acquisition

The input data was derived from a dataset created for the purpose of financial fraud detection from commonly using credit card transactions, digital payment platforms and across multiple financial institutions. This dataset is prevalent for machine learning model evaluations and establishment in the detection of fraudulent financial activities. In addition to a binary label that indicates whether or not a

transaction was fraudulent, the dataset offers through its labeled records adequate labeled amounts and contains information that includes transaction amounts, transaction time, transaction location, merchant category, device id, and anonymized customer information. The original dataset consisted of over 500,000 transactions, and following pre-processing and addressing class imbalance through the use of SMOTE, addressing of missing values, removal of duplicate records, and normalization of numerical features. The following cleaning, the dataset was 100,000 well-defined records, in which were split into 15% validation set, 70% training set, and 15% test set, to maximize the effectiveness of the model while maintaining balance.

Pre-processing utilizing Confidence Partitioning Sampling Filtering (CPSF)

In this section, we discuss pre-processing with CPSF [25]. It is applied to manage missing data, standardizing the scale of data, and excluding duplicate records. The CPSF model is employed both in data selection and in model training, maximizing learning efficiency and reliability by partitioning the samples into five separate partitions based on the level of certainty, sampling the informative instances, and filtering out the more uncertain or noisy data. It maximizes model performance by only using the inputs that have been determined to be of the highest quality, reduces training time and improves robustness with respect to data imbalance or noise. In fraud detection pre-processing, a high-probability region is defined to represent normal transaction behavior, which helps handle missing data by ensuring imputations stay within typical value ranges, as expressed as equation (1).

$$a.s \int_{D_{q(y)}^v} q(y) dy = 1 - \beta \quad (1)$$

Here, $a.s$ define the cloud resources, $q(y)$ define the labelling technique, β define the edge computing system, dy define the gray matrix and $D_{q(y)}^v$ define the resource consumption. To standardize data scaling in fraud detection, normalized weights are assigned to data points based on their probability, emphasizing typical transaction behavior, as expressed as equation (2).

$$\omega_f = \frac{q(\hat{y}_f)}{\sum_{f=1}^F q(\hat{y}_f)} \quad (2)$$

Here, ω_f represents the weight probability assigned to a transaction sample in effective fraud detection systems, $q(\hat{y}_f)$ defines the test model, and F represents the state covariance matrix, which tracks subtle behavioural changes to enhance intelligent response accuracy. CPSF standardizes input dimensions, reduces computational load, maintains data consistency, and balances speed and quality efficiently. The CPSF method is used for removing duplicate records in Equation (3),

$$p(Y) \propto \lim_{\alpha, \tau \rightarrow 0} \sum_{l=1}^{+\infty} \omega_l \delta(Y - \hat{Y}_l) \quad (3)$$

Here, τ indicates the sampling interval indicates data collection frequency, $\delta(\cdot)$ represents standardizing data scaling ensures consistent feature ranges, focusing on key transaction traits, and $p(Y)$ represents the probability density function helps detect fraud anomalies. Finally, the CPSF method handles missing data, removes duplicate records, and standardizes data scaling. Next, the data balancing system receives the pre-processed data.

Data balancing using Adaptive Support Vector -Borderline SMOTE (ASV- SMOTE)

In this segment, Data balancing using Adaptive SV-Borderline SMOTE (ASV- SMOTE) [24] is discussed. Unbalanced data is balanced using the ASV-SMOTE. ASV-SMOTE was preferred over SMOTE, Borderline-SMOTE, and ADASYN due to its use of support vectors to focus on informative borderline instances, reducing noise and improving sample quality. By adaptively generating synthetic data near critical decision boundaries, it enhances minority class prediction and reduces overfitting, making it ideal for sensitive tasks like fraud detection and medical diagnosis. The kernel-based squared distance aids ASV-SMOTE in handling imbalanced fraud data, as expressed in equation (4).

$$d^\varphi(x_i, x_j)^2 = K(x_i, x_j) - 2K(x_i, x_j) + K(x_i, x_j) \quad (4)$$

Where, $d^\varphi(x_i, x_j)^2$ represents the distance between points x_i and x_j in the feature space induced by the kernel φ , (x_i, x_j) represents input data points, $K(x_i, x_j)$ indicates as Kernel function measuring similarity between x_i and x_j in the original input space. ASV-SMOTE uses feature-space interpolation in convex regions to generate synthetic samples for fraud detection, as expressed in equation (5).

$$\varphi(x^{ij}) = \varphi(x_i) + \delta_d^{ij}(\varphi(x_j) - \varphi(x_i)) \quad (5)$$

Where, δ_d^{ij} represents random number, inward generation creates the newly synthesized sample, which is then situated between x_i and x_j , and δ^{ij} is between 0,1. It combines real and synthetic data in a kernel model to enhance fraud detection on imbalanced datasets, as expressed in equation (6).

$$f(x) = \text{sgn} \left\{ \sum_{i=1}^N \alpha_i K(x, x_i) + \sum_{j=N+1}^{N+P} \alpha_j K(x, x_j^{pq}) + b \right\} \quad (6)$$

Where, $f(x)$ indicated as the output of the decision function at input x , α_i and α_j are the coefficients corresponding to the training points respectively. Finally the ASV-SMOTE performs balanced from imbalanced data and then the balanced data are given to feature extraction.

Feature extraction using Exponential Distance Transform (EDT)

In this sector, Feature extraction utilizing Exponential Distance Transform (EDT) [26] is discussed. It is used to extract discriminative features like transaction amount, time of day, location, merchant category, and recipient information. EDT is employed in image processing and computer vision to enhance spatial awareness and feature representation by applying an exponential decay to distance values from key structures. It improves accuracy and robustness by prioritizing nearby, relevant features while reducing the influence of distant or noisy regions. EDT enables smoother transitions, supports edge-aware operations, and integrates well with machine learning models due to its differentiable nature. The minimum distance between a data point and reference points is used to extract discriminative features like transaction amount during feature exploration in fraud detection, as expressed in equation (7).

$$DT(x, y) = \min \sqrt{(x - x_i)^2 + (y - y_i)^2}, \text{ where } i \in I \quad (7)$$

Here, I indicates all irrelevant patterns in the transaction dataset, (x, y) represent the coordinates of a transaction in the system, and (x_i, y_i) represent the coordinates of transaction amount features. EDT improves fraud detection by accurately mapping transaction amounts, thereby enhancing the extraction of relevant features even in noisy or manipulated financial records. The merchant feature is extracted in equation (8),

$$IDT = \frac{1}{DT(x, y) + C} \quad (8)$$

Here C represents a constant introduced to prevent computational errors, DT represents the time difference from that transaction to a reference time point. EDT aids in handling temporal noise, improving the reliability of fraud detection by preserving essential time-based transaction patterns. The time feature is extracted in equation (9)

$$FIDT = \frac{1}{DT(x, y)^{\alpha \cdot DT(x, y) + \beta} + C} \quad (9)$$

In order to improve the assessment of recipient patterns, the distance function in the IDT is used to exponentiate the primary term. This enhances fraud detection by identifying anomalies in recipient information, which are often indicative of unauthorized or suspicious transactions. Lastly, EDT extracts discriminative features like recipient information, merchant category, location, time of day, and transaction amount. The classification model is then fed the feature extraction.

Classification Using Interpretable Generalized Additive Neural Network (IGANN)

In this segment, Classification utilizing IGANN [27] is discussed. IGANN is employed to detect fraud and classify the financial transaction as Genuine and Fraudulent. IGANN can be applied to enhance credit risk evaluation by modeling feature-wise non-linear relationships in a transparent and interpretable way. IGANN learns the individual effect of each financial indicator, such as income, credit score, and debt ratio, on the risk prediction outcome. By providing clear visualizations of how each feature contributes to the credit decision, IGANN enables financial institutions to maintain high accuracy while ensuring model transparency and regulatory compliance. The linear component and nonlinear functions in a neural network are used to classify financial transactions in fraud detection, as expressed in equation (10).

$$\hat{y} = \langle a, x \rangle + b + \sum_{l=1}^L S_l f_l(x) \quad (10)$$

Where, $\hat{y} = \langle a, x \rangle + b$ represents the overall trend of the connection between the input x and the output y , while the summation term $\sum_{l=1}^L S_l f_l(x)$ represents a set of nonlinear adjustments to this relationship. The neural network function models the classification of financial transactions as genuine in fraud detection, as expressed in equation (11).

$$f_{\theta}(X) = \sum_{K=1}^d \sum_{j=1}^N \alpha_j^k \sigma(X_k W_j^k) \quad (11)$$

Here $f_{\theta}(X)$ is denoted as the output function, parameterized by θ , d refers to the count of input features or dimensions in the transaction data, N represents the number of basic functions or neurons in the model; α_j^k are the weights that scale the output of each basis function, $\sigma(X_k W_j^k)$ are the weights applied to each basis function, $X_k W_j^k$ which modulate the strength of the non-linearity applied to the inputs. The neural network function is used to classify financial transactions as fraudulent in fraud detection, as expressed in equation (12).

$$f_{\theta}(X) = \sum \bar{\alpha}_j \sigma(\langle X, \bar{W}^j \rangle) \quad (12)$$

Where $f_{\theta}(X)$ the output function is parameterized by θ , which maps the input X transaction data to a classification outcome; $\bar{\alpha}_j$ are the weights associated with each basis function, determining their contribution to the final decision. Finally IGANN method has classified the financial transaction as Genuine and Fraudulent.

Result and Discussion

The results of proposed method are discussed in this sector. The proposed EFS-IML-EFD-TS method is implemented and simulated in Python, compiled using Jupiter Notebook, and executed on a system with 64 GB RAM, Intel Core i9-13900K CPU, and 500 GB SSD storage. The process begins by splitting the dataset into training (70%) and testing (15%) sets, followed by performance evaluation using various classifiers. The obtained result of the proposed EFS-IML-EFD- approach is analyzed with existing systems like OPFT-MLT, FFD-MLSD, and TP-AI-FFD respectively.

Performance Measure

This is an important step in choosing the best classifier. Performance metrics that are assessed include detection rate, F1-score, recall, accuracy, and precision. The performance metric is used to scale the performance metrics. To scale the performance metric, the True Negative (TN), True Positive (TP), False Negative (FN) and False Positive (FP) samples are needed.

Accuracy

The accuracy of a method evaluates its overall correctness based on the percentage true negative and of true positive predictions among all forecasts. It gives an indication of how well the method identifies instances that are positive and negative over the whole dataset.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (13)$$

Precision

One measure of machine learning method's efficiency is precision, or how well method creates positive forecasts. It is measured using the equation (14) that follows.

$$Precision = \frac{TP}{(TP + FP)} \quad (14) \quad 3.1.3$$

Recall

Recall measures a method's capacity to correctly identify all relevant instances, focusing on minimizing false negatives. It is crucial in situations where capturing all true positives is more important than avoiding false positives.

$$Recall = \frac{TP}{TP + FN} \quad (15)$$

Performance Analysis

Fig 2–8 displays the simulation outcomes of the proposed EFS-IML-EFD-TS method. Then the proposed EFS-IML-EFD-TS method is compared with the existing OPFT-MLT, FFD-MLSD-DNN, and TP-AI-FFD-DNN methods respectively.

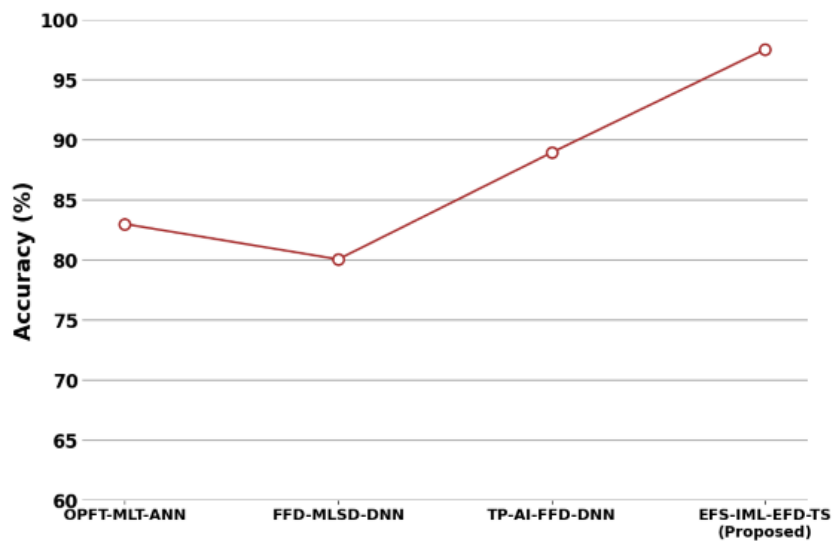


Fig 2: Performance Analysis of Accuracy

Fig 2 illustrates the performance analysis of accuracy. OPFT-MLT achieves an accuracy of 83%, FFD-MLSD-DNN scores 80%, and TP-AI-FFD-DNN reaches 89%. The proposed EFS-IML-EFD-TS

method outperforms the others with an accuracy of 98%. The proposed method demonstrates the highest accuracy, highlighting its superior effectiveness for the given task.

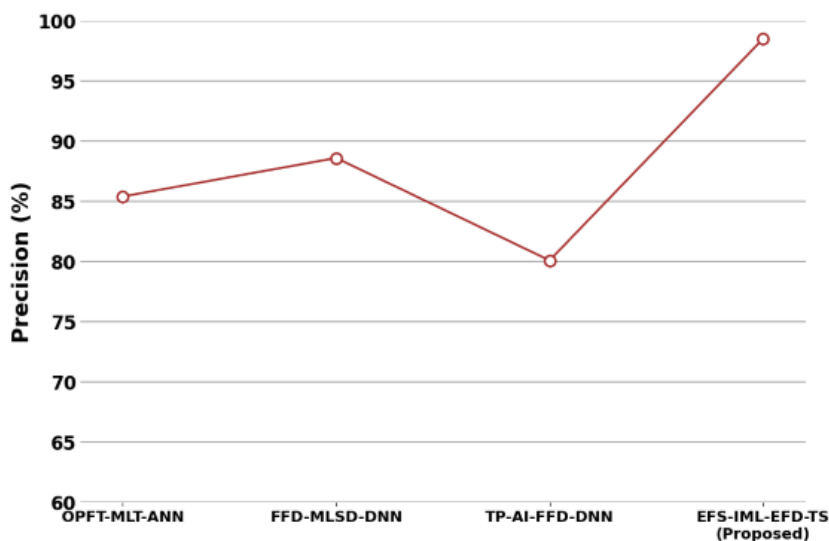


Fig 3: Performance Analysis of Precision

Fig 3 illustrates the performance analysis of precision. OPFT-MLT-ANN achieves 85.5% precision, FFD-MLSD-DNN around 88.5%, TP-AI-FFD-DNN drops to about 80.5%, and the proposed EFS-IML-EFD-TS significantly outperforms the others with 98.5%. The proposed method demonstrates the highest precision, highlighting its superior effectiveness in minimizing false positives and enhancing accuracy.

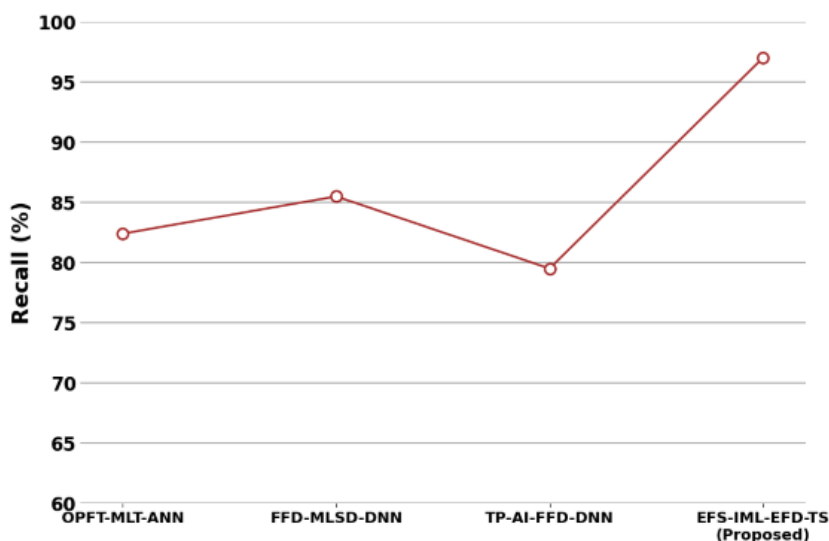


Fig 4: Performance Analysis of recall

Fig 4 illustrates the performance analysis of recall. OPFT-MLT-ANN achieves 82% recall, FFD-MLSD-DNN improves to around 85.5%, TP-AI-FFD-DNN drops to about 79.5%, and the proposed EFS-IML-EFD-TS significantly outperforms the others with 97%. The proposed method demonstrates the highest recall, showcasing its superior capability in correctly identifying actual positive cases and minimizing false negatives.

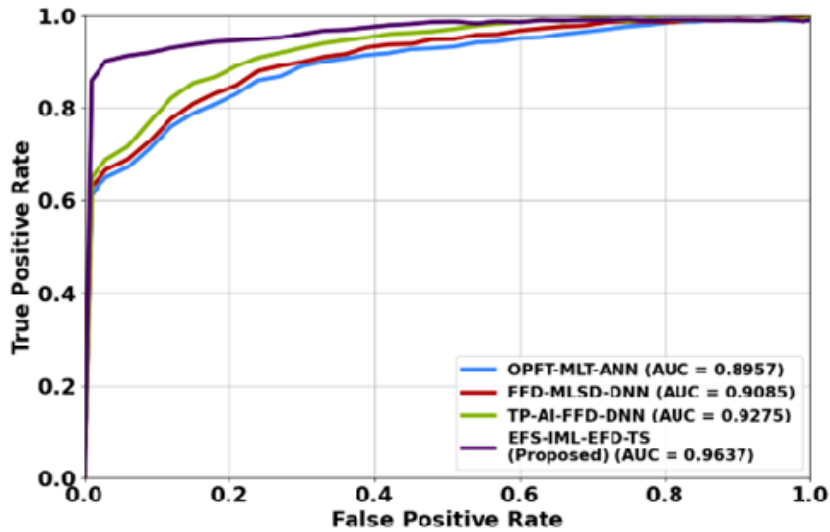


Fig 5: Performance Analysis of AUC

Fig 5 illustrates the performance analysis of the AUC curve. Each technique's false positive rate (FPR) and true positive rate (TPR) are contrasted using the AUC curve. The proposed EFS-IML-EFD-TS method achieves a TPR of about 0.97 while maintaining a low FPR, consistently outperforming the others. In comparison, OPFT-MLT-ANN reaches a TPR of 0.89, FFD-MLSD-DNN about 0.91, and TP-AI-FFD achieves around 0.93. With the highest AUC of 0.9636, the proposed method demonstrates superior discriminatory power, effectively distinguishing between negative and positive instances across all thresholds.

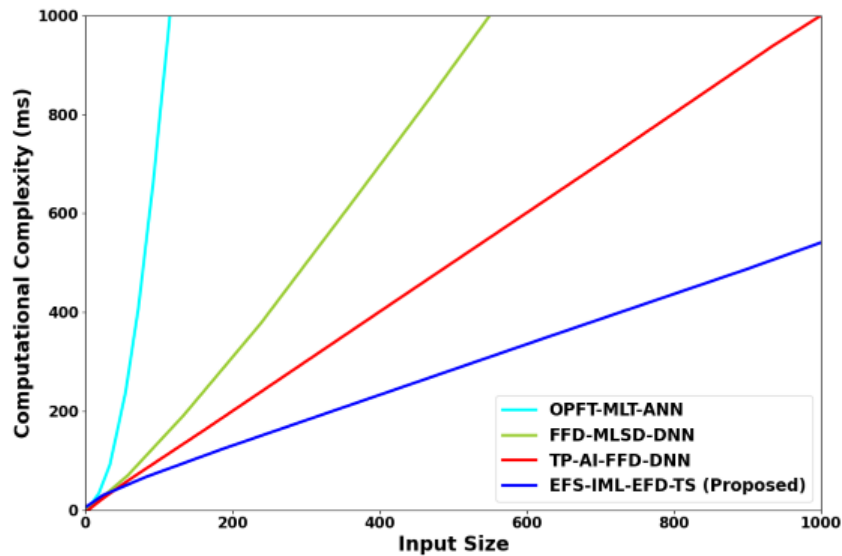


Fig 6: Performance Analysis of Computational Complexity

Fig 6 illustrates the performance analysis of the computational complexity. The computational complexity graph compares the processing time against input size for each method. The proposed EFS-IML-EFD-TS model demonstrates the lowest and most consistent growth, starting at around 2 ms for small inputs and reaching only about 35 ms at an input size of 1000. In contrast, TP-AI-FFD-DNN starts at 5 ms and grows linearly to about 120 ms. FFD-MLSD-DNN begins near 10 ms and climbs steeply to 420 ms, while OPFT-MLT-ANN shows the highest complexity, rising rapidly from 12 ms to over 600 ms. These values highlight the superior scalability and computational efficiency of the proposed method, especially as input size increases.

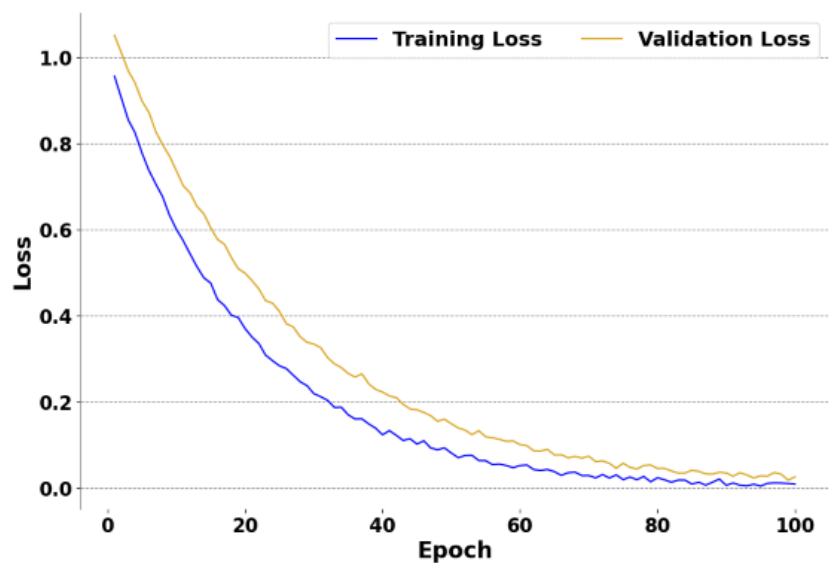


Fig 7: Performance Analysis of Loss

Fig 7 illustrates the performance analysis of the loss. A machine learning method's training progress over 100 epochs is depicted by the loss curve graph, where the Y-axis indicates loss and the X-axis represents epochs. Both the validation loss and training loss decrease steadily from around 1.0 to near 0.05, indicating effective learning. The close alignment of the two curves throughout training suggests good generalization and minimal over fitting. The curves begin to flatten around epochs 80 and 100, signalling convergence and a well-optimized model.

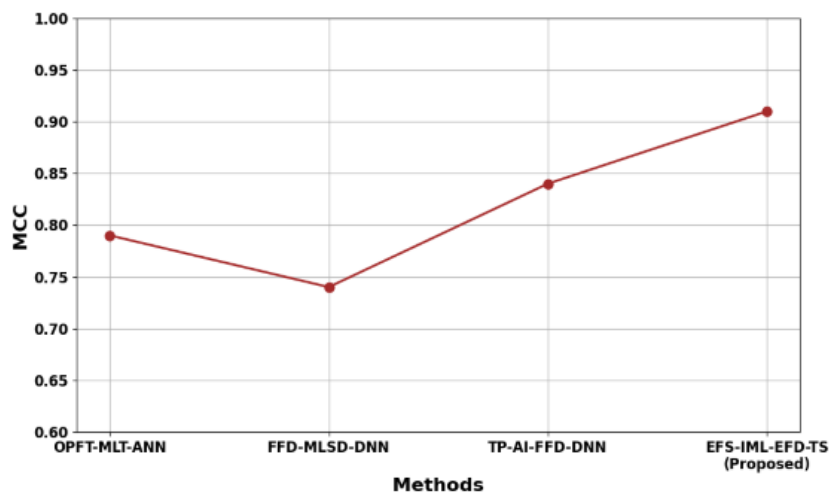


Fig 8: Performance Analysis of MCC

Fig 8 illustrates the performance analysis of the MCC. The MCC line plot compares the classification performance of four methods, with the Y-axis representing the Matthews Correlation Coefficient. OPFT-MLT-ANN achieves an MCC of 0.79, FFD-MLSD-DNN slightly lower at 0.74, and TP-AI-FFD-DNN improves significantly to 0.84. The proposed EFS-IML-EFD-TS technique outperforms all others with the highest MCC of 0.91, representing superior classification accuracy. MCC reflects the balance between true and false predictions; the results highlight the proposed technique's effectiveness and reliability in delivering high-quality predictions.

Table 2: Comparison Results of the Performance Analysis

Methods	F1-Score	Computational Time
OPFT-MLT-ANN	92.5%	1.159

FFD-MLSD-DNN	89.6%	1.136
TP-AI-FFD-DNN	87.2%	1.147
EFS-IML-EFD-TS (proposed)	97.5%	1.125

Table 2 displays the comparison results of the performance analysis. In this analysis, the F1-score performance of the methods is as follows: the proposed EFS-IML-EFD-TS achieved the highest F1-score at 97.5%, followed by OPFT-MLT at 92.5%, FFD-MLSD at 89.6%, and TP-AI-FFD at 87.2%. In terms of computational time, the proposed EFS-IML-EFD-TS also demonstrated the shortest processing time of 1.125 seconds, while FFD-MLSD required 1.136 seconds, TP-AI-FFD took 1.147 seconds, and OPFT-MLT recorded the longest time of 1.159 seconds.

Conclusion

In conclusion, the method EFS-IML-EFD-TS presented in this manuscript provides a solid approach to solution fraud detection in financial transaction systems. When you combine state-of-the-art pre-processing with the interpretable model used in this work, it generates the best detection capability by integrating advances in detection while ensuring the reliability to detect fraud and not providing faulty fraud detection in financial transactions. The application of this model can invariably strengthen fraud prevention onboard in the fintech space while generating trust and ensures innovation in the financial services industry. The EFS-IML-EFD-TS method is implemented in Python. The proposed EFS-IML-EFD-TS experiment achieves precision at 98.5%, accuracy of 98%, recall of 97%, an F1-score of 97.5%, MCC of 0.91, AUC of 0.9636, with loss of 0.05, and computational time of only 1.125 seconds, demonstrating to outperformed all methods concerning performance, accuracy and efficiency. The proposed EFS-IML-EFD-TS frameworks of financial fraud detection have an excellent opportunity to increase the accuracy of detection and computational efficiency across a variety of financial datasets. The experiments in this work have also revealed the ongoing challenges of addressing class imbalance with financial transactions and the challenge of enabling generalization across different transaction patterns, thus the need for future work to consider hybrid learning solutions and improved feature engineering strategies to develop the model performance while keeping computational time to a possible minimum.

REFERENCES

- [1] Al-dahasi, E. M., Alsheikh, R. K., Khan, F. A., & Jeon, G. (2025). Optimizing fraud detection in financial transactions with machine learning and imbalance mitigation. *Expert Systems*, 42(2), e13682.
- [2] Gandhar, A., Gupta, K., Pandey, A. K., & Raj, D. (2024). Fraud detection using machine learning and deep learning. *SN Computer Science*, 5(5), 453.
- [3] Reddy, S. R. B., Kanagala, P., Ravichandran, P., Pulimamidi, R., Sivarambabu, P. V., & Polireddi, N. S. A. (2024). Effective fraud detection in e-commerce: Leveraging machine learning and big data analytics. *Measurement: Sensors*, 33, 101138.
- [4] Zioviris, G., Kolomvatsos, K., & Stamoulis, G. (2024). An intelligent sequential fraud detection model based on deep learning. *The Journal of Supercomputing*, 80(10), 14824-14847.
- [5] Zeng, Q., Lin, L., Jiang, R., Huang, W., & Lin, D. (2025). NNEnsLeG: A novel approach for e-commerce payment fraud detection using ensemble learning and neural networks. *Information Processing & Management*, 62(1), 103916.
- [6] Zhu, H., Wang, C., & Chai, S. (2024). Detecting evolving fraudulent behavior in online payment services: Open-category and concept-drift. *IEEE Transactions on Services Computing*.
- [7] Alghamdi, S., Daim, T., & Alzahrani, S. (2024). Organizational Readiness Assessment for Fraud Detection and Prevention: Case of Airlines Sector and Electronic Payment. *IEEE Transactions on Engineering Management*.
- [8] Khattri, V., Kumar Nayak, S., Kumar Singh, D., & Bhateja, V. (2024). Design and Implementation of Fraud Detection-Decision Support System Framework. In *Identification and Mitigation of Fraudulent Online Transactions Using Authentication and Fraud Detection System* (pp. 91-107). Singapore: Springer Nature Singapore.

- [9] Charizanos, G., Demirhan, H., &İçen, D. (2024). An online fuzzy fraud detection framework for credit card transactions. *Expert Systems with Applications*, 252, 124127.
- [10] Zhao, C., Sun, X., Wu, M., & Kang, L. (2024). Advancing financial fraud detection: Self-attention generative adversarial networks for precise and effective identification. *Finance Research Letters*, 60, 104843.
- [11] Surianarayanan, C., Kunasekaran, S., &Chelliah, P. R. (2024). A high-throughput architecture for anomaly detection in streaming data using machine learning algorithms. *International Journal of Information Technology*, 16(1), 493-506.
- [12] Cai, S., & Xie, Z. (2024). Explainable fraud detection of financial statement data driven by two-layer knowledge graph. *Expert Systems with Applications*, 246, 123126.
- [13] Innan, N., Sawaika, A., Dhor, A., Dutta, S., Thota, S., Gokal, H., ...&Bennai, M. (2024). Financial fraud detection using quantum graph neural networks. *Quantum Machine Intelligence*, 6(1), 7.
- [14] [Tian, Y., & Liu, G. (2024). Spatial-temporal-aware graph transformer for transaction fraud detection. *IEEE Transactions on Industrial Informatics*.
- [15] [Kumar, S., Vical, U., Sinha, S., &Upadhyay, K. K. (2025). Regression model for credit card fraud detection. In *Advances in Electronics, Computer, Physical and Chemical Sciences* (pp. 99-105). CRC Press.
- [16] Almazroi, A. A., &Ayub, N. (2023). Online payment fraud detection model using machine learning techniques. *Ieee Access*, 11, 137188-137203.
- [17] Usman, A. U., Abdullahi, S. B., Liping, Y., Alghofaily, B., Almasoud, A. S., & Rehman, A. (2024). Financial Fraud Detection Using Value-at-Risk with Machine Learning in Skewed Data. *Ieee Access*.
- [18] Awosika, T., Shukla, R. M., &Pranggono, B. (2024). Transparency and privacy: the role of explainable ai and federated learning in financial fraud detection. *IEEE Access*.
- [19] Cui, Y., Han, X., Chen, J., Zhang, X., Yang, J., & Zhang, X. (2025). FraudGNN-RL: A Graph Neural Network With Reinforcement Learning for Adaptive Financial Fraud Detection. *IEEE Open Journal of the Computer Society*.
- [20] Talukder, M. A., Khalid, M., & Uddin, M. A. (2024). An integrated multistage ensemble machine learning model for fraudulent transaction detection. *Journal of Big Data*, 11(1), 168.
- [21] Singh, K., Kolar, P., Abraham, R., Seetharam, V., Nanduri, S., & Kumar, D. (2024). Automated Secure Computing for Fraud Detection in Financial Transactions. *Automated Secure Computing for Next-Generation Systems*, 177-189.
- [22] Tang, Y., & Liu, Z. (2024). A distributed knowledge distillation framework for financial fraud detection based on transformer. *Ieee Access*.
- [23] <https://www.kaggle.com/datasets/sriharshaeedala/financial-fraud-detection-dataset/data>
- [24] Guo, J., Wu, H., Chen, X., & Lin, W. (2024). Adaptive SV-Borderline SMOTE-SVM algorithm for imbalanced data classification. *Applied Soft Computing*, 150, 110986.
- [25] Qiang, X., Xue, R., & Zhu, Y. (2024). Confidence partitioning sampling filtering. *EURASIP Journal on Advances in Signal Processing*, 2024(1), 24.
- [26] Li, B., Chen, J., Yi, H., Feng, M., Yang, Y., Zhu, Q., & Bu, H. (2024). Exponential distance transform maps for cell localization. *Engineering Applications of Artificial Intelligence*, 132, 107948.
- [27] Kraus, M., Tschernutter, D., Weinzierl, S., & Zschech, P. (2024). Interpretable generalized additive neural networks. *European Journal of Operational Research*, 317(2), 303-316.

Appendix

		Frequency	Percentage
Gender	Male	72	48.0
	Female	78	52.0
Age	Under 18 years old	23	15.3
	18-30 years old	22	14.7
	31-40 years old	19	12.7
	41-50 years old	35	23.3
	51-60 years old	28	18.7
	60 years old and over	23	15.3
Education	High school and under	46	30.7
	Undergraduate /Post-secondary	58	38.7
	Master's degree and over	46	30.7
	Total	150	100.0

Appendix 1, Table of Statistical Characteristics of the Sample Population

Item	Factor load		Commonality	CR value
	PA	NA		
lively				
keen				
cheerful				
keen				
thrilled				
proud				
glad				
Energetic				
Grateful				
Shameful				
Awful				
Scared				
Tight				
Frightened				
Guilty				
Irritable				
Shaking with fright				
Dread				
Note: PA = positive affect, NA = negative affect. All CR values are significant at the $p < 0.001$ level.				

Appendix 2, PANAS scale

Latent variable	Measurement term	Factor loading	Cronbach's Alpha	AVE	Combined Reliability (CR)
Functional value	SQ1				
	SQ2				
	SQ3				
	SQ4				
	SQ5				
Sentimental value	EMV (1)				
	EMV (2)				
	EMV (3)				
	EMV (4)				
Novelty value	NV (1)				
	NV (2)				
	NV (3)				
	NV (4)				
	NV (5)				
Advertising value	XD (1)				
	XD (2)				
	XD (3)				
Social value	SV (1)				
	SV (2)				
	SV (3)				
	SV (4)				
	SV (5)				
	SV (6)				
Satisfaction	SAT (1)				
	SAT (2)				
	SAT (3)				
	SAT (4)				
Stance	ATT (1)				
	ATT (2)				

Appendix 3, PERVAL scale

Anderson and Fornell;Oliver Scale			
Variant	Serial number	Measurement item	Sources of indicators
Customer satisfaction	SA1	It's a wise choice to buy cultural and creative products in this art museum.	Anderson and Fornell(1994); Oliver (1997) ;
	SA2	SA2I was delighted by the interactive experience of buying cultural and creative products from this art museum.	
	SA3	SA3Overall, I was satisfied with the interactive experience of purchasing cultural and creative products from this art museum.	
	SA4	SA4I think the interactive experience of purchasing cultural and creative products with this art museum meets my expectations.	
	SA5	SA5I feel good about using this art museum to buy cultural and creative products.	

Appendix 4, Anderson and Fornell;Oliver scale

KMO and Bartlett's test			
KMO	Sample	Suitability	.822
Quantity			
Bartlett's	test	of	Approxim
sphericity			ate
			chi-square
		Degrees	300
		of	
		freedom	
		Significan	.000
		ce	

Variant	Alpha、Cronbach Alpha	Item count
Pre-interactive	.884	5
Mid-interaction	.876	4
Post-interaction	.777	3
Emotional experience	.839	6
Perceived value	.778	4
Overall satisfaction	.819	3

Appendix 5、Reliability test form, validity test (KMO & Bartlett's test)

Total Variance Explained									
Compon ents	Initial eigenvalue			Extract the sum of the squares of the loads			Rotating load sum of squares		
	Total	Perce ntage of varian ce	Cumulati ve %	Tot al	Percent age of varianc e	Cumulati ve %	Tot al	Percent age of varianc e	Cumulati ve %
1	6.034	24.136	24.136	6.034	24.136	24.136	3.559	14.236	14.236
2	3.999	15.997	40.133	3.999	15.997	40.133	3.349	13.396	27.632
3	2.525	10.101	50.234	2.525	10.101	50.234	3.021	12.084	39.716
4	1.672	6.689	56.923	1.672	6.689	56.923	2.498	9.992	49.708
5	1.325	5.299	62.222	1.325	5.299	62.222	2.183	8.733	58.441
6	1.190	4.761	66.983	1.190	4.761	66.983	2.135	8.542	66.983
7	.833	3.333	70.316						
8	.771	3.083	73.400						
9	.655	2.620	76.020						
10	.622	2.489	78.508						
11	.555	2.219	80.728						
12	.525	2.101	82.828						
13	.475	1.901	84.729						
14	.460	1.839	86.568						
15	.429	1.715	88.283						
16	.394	1.575	89.859						
17	.371	1.483	91.341						
18	.357	1.427	92.768						
19	.313	1.254	94.022						
20	.296	1.186	95.208						
21	.286	1.144	96.352						
22	.275	1.099	97.450						
23	.251	1.003	98.453						
24	.229	.916	99.369						
25	.158	.631	100.000						
Extraction method: principal component analysis									

Appendix 6、Total Variance Interpretation Table

	CMIN/DF	GFI	RMR	RMSEA	NFI	IFI	TLI	CFI
Standard	≤3.00	≥0.90	≤0.08	≤0.08	≥0.90	≥0.90	≥0.90	≥0.90
Measured values	1.119	0.928	0.052	0.028	0.914	0.990	0.988	0.990

Appendix 7、Main Fit Indicators for Study 1 (Structural Equation of Emotional Experience) Model

Hopeless		Standardised path coefficients	S.E.	C.R.	P	Conclusion
H	Pre-interaction-Emotional experience	.458	.097	4.532	***	Set up
	Emotional experience-Overall satisfaction	.318	.111	2.929	.003	Set up
	Pre-interactive-Overall satisfaction	.245	.102	2.358	.018	Set up

Appendix 8、Coefficients and Significance of Major Paths in Study 1 (Structural Equation of Emotional Experience)

Hop eles s	path	Effect	SE	S.E.	P	LB (95%CI)	UB (95%CI)	Conclusi on
	Pre-interaction-Emotional experience-Ove rall satisfaction	Total	.391	.087	.001	.221	.549	Establish ed
		Direct	.245	.107	.028	.029	.441	Establish ed
		Indirect	.146	.057	.003	.049	.269	Establish ed

Appendix 9、Results of mediation analyses for Study 1 (Structural Equations of Emotional Experience)

	CMIN/DF	GFI	RMR	RMSEA	NFI	IFI	TLI	CFI
Standard	≤3.00	≥0.90	≤0.08	≤0.08	≥0.90	≥0.90	≥0.90	≥0.90
Measured values	1.132	0.909	0.071	0.030	0.844	0.985	0.982	0.985

Appendix 10、Main Fit Indicators for the Study II (Perceived Value Structural Equation) Model

Hopeless	path	SE	Estimate	S.E.	C.R.	P	Conclusion
	Mid-interaction→ Perceived value	.483	.356	.076	4.665	***	Established
	Emotional experience→ Overall satisfaction	.362	.377	.081	4.655	***	Established
	Perceived value→ Overall satisfaction	.332	.380	.103	3.697	***	Established

Appendix 11、Study 2 (Perceived Value Structure Equation) Main Path Coefficients and Significance

Parameter	Estimate	Lower	Upper	P	Conclusion
Mid-interaction→Emotional experience→Overall satisfaction	.168	.072	Established	.000	Established
Mid-interaction→Perceived value →Overall satisfaction	.160	.058	Established	.001	Established

Appendix 12、Results of mediation analyses for Study 2 (Perceived Value Structure Equation)

	CMIN/DF	GFI	RMR	RMSEA	NFI	IFI	TLI	CFI
Standard	≤3.00	≥0.90	≤0.08	≤0.08	≥0.90	≥0.90	≥0.90	≥0.90
Measured values	1.172	0.934	0.063	0.069	0.896	0.954	0.934	0.953

Appendix 13、Main fit metrics of the model for Study 3 (Interactive Late Validation Model)

Hopeless	Path	SE	S.E.	C.R.	P	Conclusion
	Post-interaction→ Perceived value	.252	.093	2.373	.018	Established
	Post-interaction→Overall satisfaction	.230	.088	2.275	.023	Established
	Perceived value→Overall satisfaction	.350	.105	3.309	***	Established

Appendix 14、Study 3 primary path coefficients and significance

Hop eles s	path	Effect	SE	S.E.	P	LB (95%CI)	UB (95%CI)	Conclusi on
	Post-interaction →Perceived value→Overall satisfaction	Total	.318	.105	.005	.144	.488	Establish ed
		Direct	.230	.109	.035	.048	.414	Establish ed
		Indirect	.088	.044	.022	.022	.164	Establish ed

Appendix 15、Results of Study III intermediation analyses

Rotated component matrix ^a						
	Ingredient					
	1	2	3	4	5	6
Pre-interactiveA1	.848					
Pre-interactiveA2	.823					
Pre-interactiveA3	.758					
Pre-interactiveA4	.805					
Pre-interactiveA5	.744					
Mid-interactionB6			.803			
Mid-interactionB7			.806			
Mid-interactionB8			.799			
Mid-interactionB9			.800			
Post-interactionC10					.831	
Post-interactionC11					.767	
Post-interactionC12					.844	
Emotional experienceA113		.699				
Emotional experienceA114		.626				
Emotional experienceA115		.777				
Cultural resonanceA216		.667				
Cultural resonanceA217		.671				
Cultural resonanceA218		.762				
Cognitive valueB119				.730		
Cognitive valueB120				.724		
Sentimental valueB221				.712		
Sentimental valueB222				.793		
Overall satisfactionC123						.821
Overall satisfactionC124						.694
Interactive engagementC225						.825
Extraction method: principal component analysis						
Rotation method: Kaiser normalised maximum variance method						
a. Rotation has converged after 7 iterations						

Appendix 16、Rotated component matrix

