

Enhancing Credit Card Fraud Detection Using Cutting-Edge Advanced Machine Learning Techniques

Received 03/19/2025
Review began 04/07/2025
Review ended 08/16/2025
Published 08/17/2025

© Copyright 2025
Bandarapu et al. This is an open access article distributed under the terms of the Creative Commons Attribution License CC-BY 4.0., which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

DOI:
<https://doi.org/10.7759/s44389-025-03921-w>

Srinivas R. Bandarapu ¹, Kiran Randhi ², Nishchai J. Manjula ³

¹. Data, AI, and ML, DigiTech Labs, Redmond, USA ². Data, AI, and ML, Amazon Web Services, Seattle, USA ³. Data, AI, and ML, Visvesvaraya Technological University, Tumkur, IND

Corresponding author: Srinivas R. Bandarapu, bandarapu@gmail.com

Abstract

Credit card fraud remains a major challenge in digital payments, necessitating advanced detection techniques. This study evaluates Deep Belief Networks (DBN), Isolation Forest, K-Nearest Neighbors (KNN), and Artificial Neural Networks (ANN) for fraud detection. Performance was assessed using Accuracy, Precision, Recall, and F1-Score. DBN, KNN, and ANN achieved 99% accuracy with high precision and recall (0.99), while Isolation Forest underperformed at 51% accuracy. To address class imbalance, SMOTE was applied, significantly improving model performance. Our comparative analysis highlights that the proposed models outperform existing approaches, demonstrating superior fraud detection capabilities. The findings contribute to minimizing financial losses, enhancing customer trust, and strengthening fraud prevention systems. Future work will explore real-time implementation and model resilience against adversarial attacks to ensure robustness in practical applications.

Categories: Network Security, AI applications, Application Security

Keywords: fraud detection, credit card fraud detection dataset, dbn model, smote, knn

Introduction

In the electronic era of great advancement, credit card fraud is clearly a problem that has grown to be such risk for individuals, banks, and financial system as a whole. As people continue to try acquiring illegal financial gains, they are forcing the development of more sophisticated methods into circumventing traditional manual fraud checks. Consequently, there has been exceptional improvement in the realm of credit card fraud detection [1] for integrating advanced technology to cater this looming threat. Credit card fraud detection has traditionally operated under mainly rule-based systems and statistical models [2]. Although effective, these methods often have a difficult time keeping pace with the dynamic and ever increasingly sophisticated world of fraud. With the proliferation of digital info came new tactics by fraudsters as well, which made it necessary to deploy more intelligent and flexible detection systems.

Credit cards have become a cornerstone of contemporary financial transactions, offering convenience, speed, and global accessibility. Their ubiquity, however, has made them increasingly vulnerable to cyber threats. The detection of fraudulent credit card activity is both a technical and operational challenge, especially as fraudsters adapt to and exploit weaknesses in existing systems [3]. Traditional rule-based systems often fall short in identifying novel fraud patterns, prompting the need for more dynamic, data-driven approaches such as machine learning and deep learning algorithms. Credit card is a plastic card having personal information of the concerned individuals, which financial institutions are offering to their customers for purchasing goods and services across world at any time. One of the most common concerns is credit card fraud, which refers to an unauthorized [4] use and illegal theft by any individual or a third party from other person's account in order to attain money through funds transfer or purchase some goods either in person or online. This kind of false behavior often causes severe financial loss [5]. The growth in the number of transactions taking place over the internet has made it easier for fraud to occur, as having the physical card is no longer necessary as you only need its details to make a payment [6]. In the regard, Lujain Essam Faisal et al. [7] stated that after the introduction of credit cards, some companies and many economic forces were liberated from both monetary policies control as well also in business strategies. This study could cover the domain of fraudulent detection of credit card, focusing on machine learning as well as deep learning a subfield in artificial intelligence for successful prediction of future unauthorized transactions. Deep learning models such as artificial neural networks (ANN) and their derivative models [8] can learn and capture nuanced patterns from big data and thus are well-placed for the evolving character of credit card fraud. Credit card fraud detection is one of the most eye-catching fields where deep learning models are exploited due to their capabilities in discovering fine-grain and non-linear characteristics of multi-dimensional actions related to fraud activity, which can be hardly captured by traditional methods. The goal of this research is to offer a detailed review of how deep learning algorithms are applied to credit card fraud detection. We examine the fundamentals of machine learning and deep learning.

How to cite this article

Bandarapu S R, Randhi K, Manjula N J (August 17, 2025) Enhancing Credit Card Fraud Detection Using Cutting-Edge Advanced Machine Learning Techniques. Cureus J Comput Sci 2 : es44389-025-03921-w. DOI <https://doi.org/10.7759/s44389-025-03921-w>

Moreover, we delve into the concepts of neural networks, with a specific focus on how these architectures are adapted for fraud detection problems. We also address key challenges and opportunities in deep learning, including adversarial training, privacy concerns in data ingestion, model interpretability (i.e., making the "black box" more transparent), among others. As financial transactions increasingly shift to the digital realm, protecting the monetary interests and trust of both individuals and businesses becomes paramount for effective credit card fraud detection [9]. This work presents a comprehensive comparative framework for credit card fraud detection by evaluating both classical machine learning and modern deep learning approaches within a unified, Synthetic Minority Over-sampling Technique (SMOTE)-balanced training pipeline. The key contributions of our study are:

1. Imbalanced Data Handling: To present a credit card fraud detection technique that is reliable and capable of dealing with the issue of data imbalance between fraudulent transactions and non-fraudulent ones in the dataset.
2. Hybrid Comparison: We systematically evaluate two deep learning models (Deep Belief Network [DBN] and ANN) alongside two classical machine learning algorithms (K-Nearest Neighbors [KNN] and Isolation Forest), providing a broad perspective on algorithm performance under real-world fraud detection constraints.
3. Statistical Robustness: Our analysis goes beyond standard metric reporting by including confidence intervals and statistical significance testing to validate the robustness of comparative results.

Related work

Credit card fraud is a rapidly growing crime targeting financial institutions, resulting in billions of dollars in losses each year [10,11]. As a result, automated fraud detection systems have become a critical priority for these organizations [12]. To improve accuracy and efficiency, enhanced fraud detection techniques are being implemented, particularly various supervised learning approaches that have shown strong performance in identifying fraudulent activities. Aftab et al. [13] compared Random Forest (RF), Logistic Regression (LR), Support Vector Machine (SVM), and Decision Trees (DT) for credit card fraud detection. Their study found that RF exhibited superior performance, achieving an impressive recall score of 84%. They also addressed the issue of imbalanced datasets using the SMOTE technique for oversampling and suggested future research based on deep learning techniques applied on dynamic datasets to enhance detection capabilities. Shraddha et al. [14] provided a comparison of various state-of-the-art machine learning and deep learning algorithms for credit card fraud detection. The model proposed here was shown to outperform several existing algorithms, hence exhibiting better results in fraud detection. However, the study noted the high rates of false alarms and the challenges posed by the accessibility of public data, indicating the need for refined algorithms to reduce false positives. Sreenidhi Appalabattla et al. [15] proposed FrauDetect, which is a robust fraud detection tool over credit cards integrating the application of Convolutional Neural Networks (CNN) and Principal Component Analysis (PCA) along with Tensor Flow and Keras. The proposed model achieved 91.9% accuracy for fraud detection, but the study noted shortcomings in the model's parameters and deep learning methods for validation, as well as privacy concerns that led to incomplete data. Dileep et al. [16] compared the detection ability of Local Outlier Factor (LOF), Isolation Forest, and Autoencoders for credit card fraud. The study reported that LOF (Classic) had the highest accuracy at 99.0%, with Autoencoders rated at 96.4%. The research pointed out that even though fraud detection mechanisms have shown high accuracy rates, they are weak in distinguishing between fraudulent and normal transactions due to the high number of individual transactions and the small number of non-fraudulent cases. Wu et al. [17] investigated a hybrid model composed of CNN and RNN to detect credit card fraud, achieving a precision rate of 95.83%. This hybrid method used the SMOTE technique to address class imbalance, demonstrating the value of merging several models for improved fraud detection accuracy.

A detailed review by Khan et al. [18] applied several machine learning algorithms and artificial intelligence mechanisms for credit card fraud detection, such as PCA, Fuzzy C-Means, LR, DT, and Naïve Bayes Classifiers. This paper provided a powerful examination method for fraud detection, recognized deficiencies in current models, and strongly supported the use of superior techniques to enhance detection. Al Balawi and Aljohani [19] applied ANN and CNN to credit card fraud detection. The research also positioned deep neural networks as a means to improve fraud detection. Singh [20] implemented fraud detection in credit cards using regression, RF classifier, and DT algorithm, which were evaluated using accuracy, precision, recall, and F1-score as metrics. The research pointed out the problems of imbalanced datasets and credit card misuse, highlighting the necessity for robust algorithms to address such issues. Prusti and Rath [21] used classification models such as RF, K-NN, Extreme Learning Machine, Multilayer Perceptron, and Bagging classifiers to classify fraudulent credit card transactions. The results of their study showed that the ensemble model combining all five algorithms achieved an accuracy of 83.83%, thereby reducing prediction error during the fraud detection process. This gain reinforces the power of combining different models to capitalize on their collective strengths.

Raghavan and Gayar [22] provided a thorough analysis of machine learning and deep learning techniques

for fraud detection across several datasets. Their study's findings demonstrated that the best models for handling larger data volumes were RF, CNN, and SVMs. They also experimented with deep learning models such as Restricted Boltzmann Machines and Autoencoders, highlighting their poor to average performance in fraud detection. In the ensemble strategy, the integration of SVM with CNN and RF performed better than individual classifiers overall, and particularly well with large datasets. However, they also noted challenges in detecting fraud in small datasets, mainly due to data sparsity and imbalance.

In their study, Khalid et al. [23] illustrated several crucial issues in credit card fraud detection, such as data distribution and online processing. An ensemble model combining SVM, KNN, RF, Bagging, and Boosting classifiers outperformed the individual classifiers across major fraud detection metrics. In the proposed study, data preprocessing, cleaning strategies, and feature engineering steps were employed to achieve the best results. They also highlighted the ensemble model's improved speed efficiency, discussing the trade-off between training time and testing time. Additionally, the research pointed out challenges related to the procurement of datasets, as client data is highly confidential and individuals are often unwilling to share their transactional information due to privacy concerns. This emphasizes the need for effective data balancing techniques in fraud detection.

Mahesh et al. [24] used under-sampling, an over-sampling algorithm called SMOTE, and a hybrid technique like SMOTE-Tomek in their study to address the imbalance problem between fraud and non-fraud cases. They trained various classification models on the balanced dataset, such as KNN, LR, RF, and SVM. According to their results, all these models achieved an accuracy of 0.99 in predicting fraudulent activities. The results also showed that the best F1-score of 0.94 was achieved when under-sampling techniques were applied to clean up known negative values and identify only those variables without typing errors. The research emphasizes the effectiveness of these techniques in enhancing fraud detection systems and suggests new directions for applying additional methods in financial security.

With all factors considered, the literature shows that a lot of focus is placed on applying various machine learning and deep learning techniques to improve the accuracy and efficacy of credit card fraud detection. In order to deal with the evolving characteristics of credit card fraud along with improving detection techniques, further studies should keep investigating hybrid models, adaptable techniques, and advanced approaches.

Materials And Methods

Machine learning and deep learning are involved in fraudulent activity detection while analyzing data on legitimate transactions. These techniques are efficient at detecting transactions anomalies earlier before they turn into major challenges [25]. In Figure 1, the first step in the process is choosing a dataset that includes information on both legitimate and fraudulent transactions. Since dataset may contain instances that may have missing values, raw form, or duplicate, data preprocessing is necessary to ensure accuracy in system predictions. To address data imbalance, we perform sampling on the imbalanced datasets. Following data organization and sampling, we separate the data into training and testing sets. Using both the training and testing sets, we monitor the behavior of the chosen machine learning models after training them on the training set. After that, we examine and contrast performance using assessment criteria including recall, accuracy, precision, and confusion matrix. The experimental approach used in this study's methodology aims to carry out real-world tests for techniques that identify credit card fraud.



FIGURE 1: Proposed Methodology for Credit Card Fraud Detection

DL, Deep Learning; ML, Machine Learning

Dataset description

We utilized the Credit Card Fraud Detection dataset for models training and testing [23]. This dataset consists of two-day worth of European credit card transaction information. A significant class imbalance is seen in the 284,807 transactions, of which 492 were found to be fraudulent, or just 0.172% of all transactions. There are 28 characteristics associated with each transaction, which are referred to as V1-V28. Except for the 'Time' and 'Amount' features, these attributes were transformed using PCA to address confidentiality concerns. The 'Time' feature measures the number of seconds that have passed between each transaction and the dataset's initial transaction. The 'Amount' element reflects the monetary value of the transaction [22].

The dataset used in this study exhibited a significant imbalance due to the low number of fraud instances. This imbalance posed challenges for the training and testing of the model [26]. To overcome this issue, SMOTE methodology were used. SMOTE involved the oversampling of 492 fraud instances to balance their quantity against that of legitimate entries, hence making them equally distributed. In this work, the ethical concerns related to our dataset range from issues of data ownership and consent, through depersonalization and aggregation strategies ensuring privacy. The dataset is maintained by the ULB Machine Learning Group in collaboration with World line. This collaboration reflects a joint study on large-scale data mining and fraudulent activities detection. But also, the decision of this scientific effort is to change dataset attributes into numeric entities, which are done by PCA. This is an ethical measure to guard the personal details looming around individual transactions and at same time bringing progress in data science. It certainly raises questions about the increasingly fine line we walk between science and responsibility to our financial information as a data set.

Data preprocessing

After selecting the dataset, our initial task involved preprocessing the data to ensure it was prepared for model training and testing. The following steps were undertaken:

- Identifying and addressing any null values by either filling or removing them.
- Standardizing the 'Amount' column to facilitate easier analysis.
- Removing the 'Time' column due to its minimal contribution during training model and their evaluation.
- Detecting and eliminating duplicate instances.

We confirmed that dataset contained no missing or null values. Given that fraudulent activities often appear as statistical outliers, we deliberately chose not to remove outliers from the fraud class to preserve important signals. Instead, limited outlier handling was applied only to the legitimate class to reduce noise without discarding rare but meaningful anomalies. This decision ensures that the dataset maintains the integrity of fraudulent patterns, which are critical for effective model training and detection.

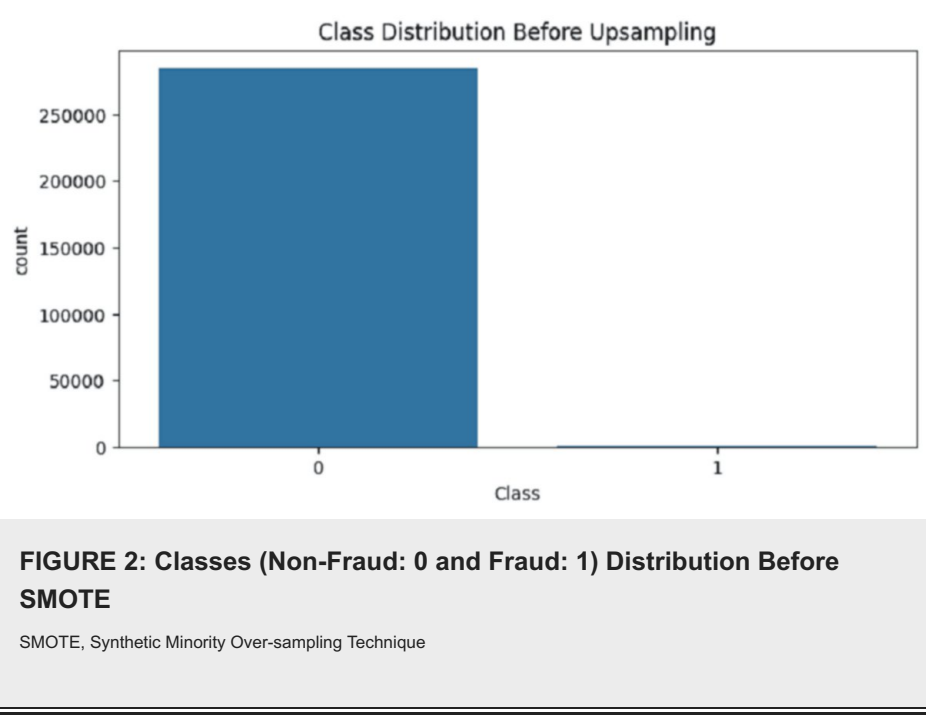
Furthermore, it was thought to be advantageous to have outliers since it brings models closer to real-world problems. By foregoing explicit outlier-handling strategies [27], we aimed to enhance model's ability to tackle the immediate effects and the variety of credit card transactions, thereby increasing its adaptability to real-world scenarios. The 'Amount' column was standardized rather than normalized because its scale was noticeably different from that of the other features (V1-V28), as indicated in dataset description. All features except 'Time' and 'Amount' had undergone PCA [28]. Feature engineering was complicated by the dataset's inherent limitations, particularly the lack of detailed information about its features. As a result, feature selection techniques were not used, as they require clear feature visibility. We made the decision to keep all features in order to prevent any potential confusion from algorithmic selection of features. Initially, the 'Time' column was believed to be a simple sequence counter without explicit temporal context (e.g., timestamps or real-world dates). However, upon further analysis, we recognized its utility in capturing transaction intervals and identifying periodic patterns. Thus, while it was not directly used as a model input, it was retained for exploratory analysis, as reflected in Figure 7, to examine trends such as clustering of fraud events across time. It was determined that the best way to maintain all potentially significant characteristics without depending on in-depth feature-specific knowledge was to use this technique, even though it would result in lengthier training and testing periods. Although the 'Time' feature was excluded from the predictive modeling due to its weak influence on classification performance, it was retained for exploratory visualization purposes. Specifically, Figure 7 uses this feature to analyze transaction distributions across time, but it was not part of any trained feature set for classification models.

Data sampling

After initial loading, the dataset originally contained 284,807 legitimate and 492 fraud transactions. During data cleaning, we removed 19,617 legitimate and 19 fraudulent transactions due to duplicate

entries, missing values, or anomalies such as invalid timestamps or corrupted values. This resulted in a cleaned dataset comprising 275,190 legitimate and 473 fraud transactions.

After completing data preprocessing, we addressed the class imbalance present in the dataset, which consisted of 275,190 legitimate transactions and only 473 fraudulent ones. To mitigate this issue, we applied SMOTE, but only to the training data after an 80/20 train-test split. This was done deliberately to prevent data leakage and ensure unbiased model evaluation. We experimented with different values for the number of nearest neighbors (k) in SMOTE, ranging from 3 to 7, and found that $k = 5$ provided the best balance between minority class generalization and overfitting. Using the original minority (fraud) class instances as input, SMOTE synthetically generated new samples until the number of fraudulent cases in the training subset matched the number of legitimate transactions, resulting in a balanced training set. The test set remained untouched to maintain evaluation integrity. This approach ensured fair model training and avoided artificially inflated performance metrics by preventing the model from seeing overly similar data during testing. As a statistical method, SMOTE helped achieve class balance without compromising the reliability of model evaluation on unseen data [29].



Initially shown in Figure 2, the dataset contained 284,807 legitimate transactions (Class 0) compared to only 492 fraudulent transactions (Class 1). This disparity is visually represented in the first bar plot, where the count of legitimate transactions vastly outnumbers that of fraudulent transactions.

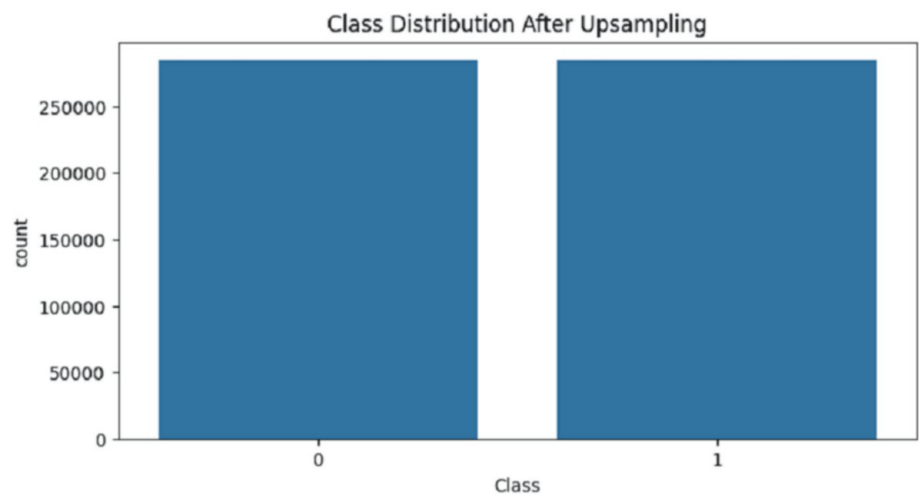


FIGURE 3: Classes (Non-Fraud: 0 and Fraud: 1) Distribution After SMOTE
SMOTE, Synthetic Minority Over-sampling Technique

To mitigate the imbalance, SMOTE was applied to the dataset shown in Figure 3. SMOTE produces new instances of the minority class based on interpolation among existing samples of the minority class. The result is a more balanced dataset, as depicted in the second bar plot, where the counts of legitimate and fraudulent transactions are approximately equal, each with 284,807 instances.

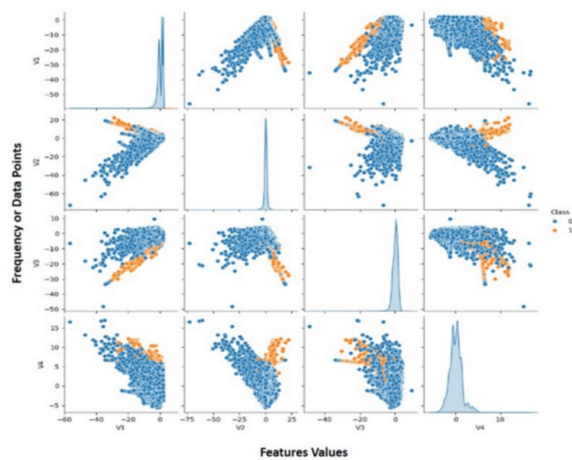


FIGURE 4: Classes Distribution Over Features V1-V4 Before SMOTE
SMOTE, Synthetic Minority Over-sampling Technique

Our dataset comprises features V1 through V28, which play a critical role in identifying patterns indicative of fraudulent activity. For illustrative purposes, we analyze the impact of SMOTE on a subset of these features: V1 to V4, as shown in Figure 4. The figure, which depicts the dataset before the application of SMOTE [30], clearly highlights the class imbalance: legitimate transactions (Class 0) significantly outnumber fraudulent ones (Class 1). This imbalance is visually evident, as the orange dots (Class 1) are sparse and scattered, making it difficult for the model to understand what traits have fraudulent transactions.

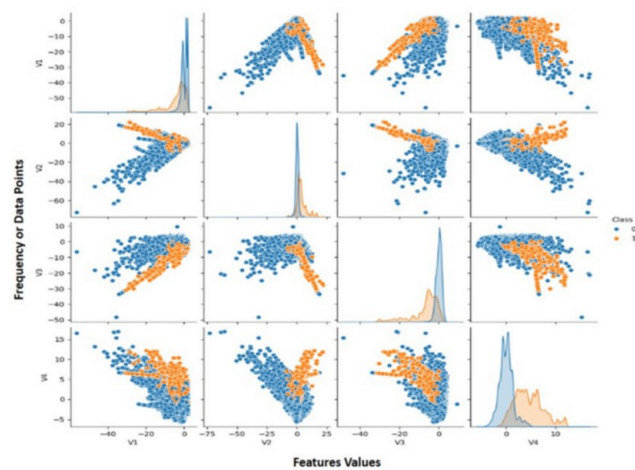


FIGURE 5: Classes Distribution Over Features V1-V4 After SMOTE

SMOTE, Synthetic Minority Over-sampling Technique

Upon applying SMOTE, Figure 5 demonstrates a transformed dataset with a balanced class distribution. SMOTE works by creating synthetic examples of the minority class (fraudulent transactions) through interpolation between existing minority instances. This results in an even representation of both classes across the feature space, as evidenced by the increased density and spread of the orange dots (Class 1).

Data statistics and visualization

In our study on credit card fraud detection, we perform a detailed analysis of various features in the dataset to understand their distributions and correlations, which are crucial for effective fraud detection. We provide an in-depth discussion of key components of the data and their relevance to our study.

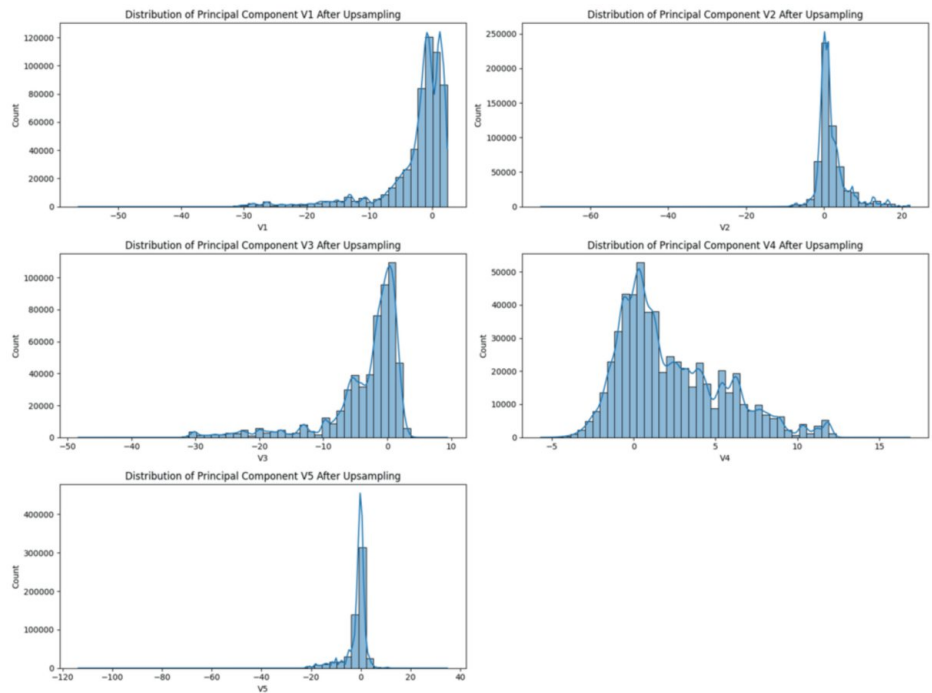


FIGURE 6: Distribution of PCA on V1-V5

PCA, Principal Component Analysis

Figure 6 illustrates the distribution of several principal components after upsampling, specifically focusing on V1 to V5. These components were selected to highlight the transformed feature space post-augmentation. The distribution of Principal Component V1 shows a skewed distribution with a significant concentration of values close to zero, accompanied by a notable tail extending towards negative values. This indicates that most data points cluster around zero, but there are significant outliers that could be important for detecting anomalies. Similarly, Principal Component V2 exhibits a more centered distribution around zero with a sharp peak and less spread compared to V1, suggesting a more compact and less variable feature. Principal Component V3 has a more strongly peaked pattern at zero that nearly resembling the negative skewness of V1 and also indicating that some important information about the data may be located near zero. Principal Component V4 generally has a normal distribution with a few provided peaks and distortions that indicating a random structure that is interlinked. Additionally, Principal Component V5 appears to be very concentrated, with few spread values to the extreme ends, as evidenced by its acute distribution and sharp peak around zero. This concentration could be essential in identifying key differentiators between fraudulent and non-fraudulent transactions.

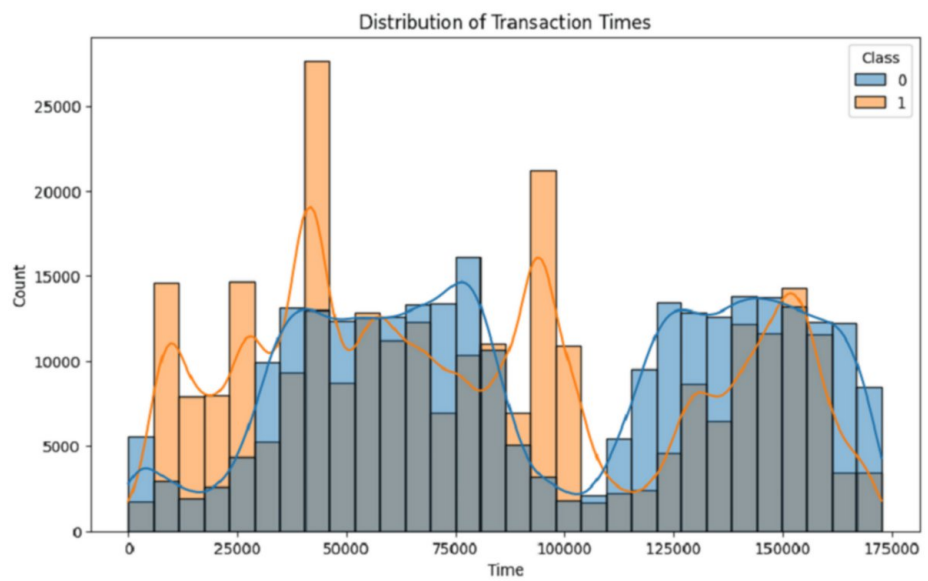


FIGURE 7: Distribution of Transaction Over Time

In Figure 7, the distribution of transaction times shows distinct patterns for non-fraudulent (Class 0) and fraudulent cases (Class 1). Non-fraudulent transactions have an even distribution across time with a few peaks, representing the times when most transactions occur. Fraudulent transactions also exhibit peaks at certain times, which could be offsets caused by common defrauders (e.g., coordinated fraud). Overlaying both classes, it shows that fraudulent transactions are not randomly distributed but rather occur during times of high non-fraudulent transaction volume. This temporal pattern underscores the need to consider the timing of transactions when detecting fraud, as it can reveal windows of vulnerability and help in the timely identification of suspicious activities.

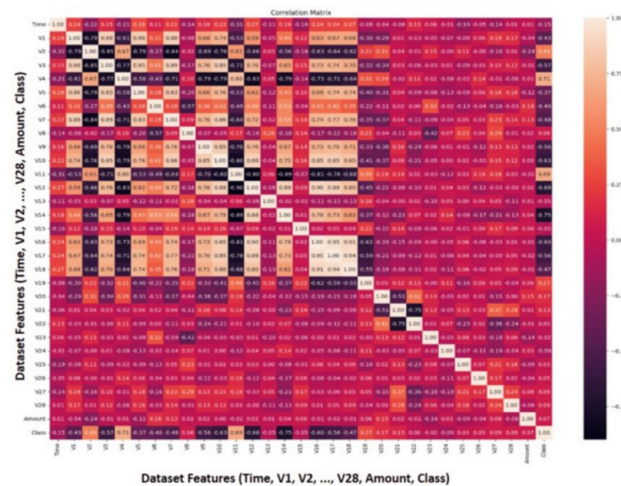


FIGURE 8: Features Correlation Matrix

In Figure 8, correlation matrix provides valuable insights into linear relationships of different features from the dataset. The correlation coefficient between two sets of features is represented by each cell in the matrix, which ranges from -1 to 1. High correlations among certain features, such as V1, V2, and V3, suggest overlapping information that can be pivotal in fraud detection. These strong negative correlations indicate that these features might counterbalance each other, providing a more nuanced understanding of the data when used together. Conversely, features like V27 and V28 show low correlations with most other features, indicating their unique contributions to the dataset. The 'Class' variable, which indicates fraudulent or non-fraudulent transactions, shows significant correlations with features like V10, V12, and V14. These features are likely critical for distinguishing fraudulent activities, making them prime candidates for feature selection in machine learning models.

Models selection

Our approach to fraud detection incorporated a variety of advanced machine learning models, each contributing uniquely to handling the dataset's complexities. Below, we introduce each model and its role in our overall strategy.

- **DBN:** DBNs [31] are deep learning-based models that consist of multiple layers of latent variables, known as hidden units. DBNs were employed in our credit card fraud detection system to uncover complex patterns hidden within the transaction data. Due to their ability to learn high-level representations, they were particularly effective at detecting subtle anomalies that characterize fraudulent activity. The deep structure of DBNs proved to be highly effective in capturing high-order interactions among transaction features, which enhanced their overall detection capability.
- **Isolation Forest:** Isolation Forest [32] is a potential anomaly detection technique, which aims at detecting the outliers in a dataset. The Isolation Forest Algorithm was selected as it has the ability to isolate anomalies within high-dimensional credit card transaction data. It does so by building many decision trees and isolating observations via random splits, meaning this would be useful in detecting the outlier's traces of fraudulent transactions.
- **KNN:** KNN [33] is one of the supervised learning algorithms with a classifier, which evaluates how closely the data points resemble their closest training examples in the feature space or vector. KNN was used to categorize transactions in flow by finding out the fraudulent activities and classify payments based on same characteristic types. KNN is extremely useful during the exploratory analysis stage where it enhance the knowledge about how transactions related and overlapped. We understood the type of fraud that was happening and found similar transactions near old ones to flag fraudulent charges based on their closeness in transaction space.
- **ANN:** ANNs [34] are supervised learning algorithms inspired by the way the human brain works, and they simulate our ability to classify things based on a set of observed traits. In the case of our workflow, we employed ANN to categorize transactions and mark out potential fraud. In particular, this approach was very helpful when we were performing exploratory data analysis as it allowed us to identify sophisticated and non-trivial relationships and structures within the transaction datasets. Training the ANN on a dataset that included possession of legitimate and fraud transactions helps us to identify credit card frauds.

Evaluation metrics

We evaluated the performance of models using several key metrics, with each metrics offering valuable insights into various aspects of the model's effectiveness.

· Accuracy: Accuracy [35] is the important metric for assessing model effectiveness. Accuracy determines the number of accurately identified instances for both (fraudulent and non-fraudulent) out of all of them. Although accuracy provides a broad indication of the number of times our model has forecast correctly, it can be extremely misleading in cases where the dataset is discontinuous and unbalanced.

$$\text{Accuracy} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN})$$

· Precision: Precision [36], also known as positive predictive value, this evaluation metric gives information about how many of the instances predicted as fraudulent are actually frauds. High accuracy means few false positives (FP).

$$\text{Precision} = \text{TP} / (\text{TP} + \text{FP})$$

· Recall: Recall [37], also called sensitivity or true positive rate, is another key measure matrices that calculates the ratio of actual fraud instances identified by our model. With a high recall, we actually have a low false-negative (FN) rate.

$$\text{Recall} = \text{TP} / (\text{TP} + \text{FN})$$

· F1 Score: F1 Score [38] is the harmonic mean of precision and recall. It gives a good blend of both types of errors (false positives and false negatives) which is most useful to evaluate our model on the imbalanced dataset.

$$\text{F1-Score} = 2 * ((\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall}))$$

Here, TP, TN, FP, and FN are "True Positive", "True Negative", "False Positive", and "False Negative", respectively.

Implementation details

The models were trained on a Windows 11 machine with an Intel Core i7-6600U dual-core CPU and 16GB RAM, without GPU acceleration. Despite the lack of GPU, DBN was implemented using a shallow architecture with early stopping criteria, enabling feasible training within a reasonable runtime. Training time was longer compared to GPU setups, but reproducibility was ensured through consistent CPU-based runs and small batch training on stratified data splits. The core analysis infrastructure is provided by Python libraries such as Scikit-learn for tasks related to machine learning, Pandas for data processing and NumPy for numerical calculations. To address the class imbalance within our dataset, we used SMOTE, which is an oversampling technique [39] and it is implemented in imbalanced-learn. This complete approach guarantees a multifaceted and broad analysis, using different libraries focused on fraud detection and data processing. The implemented models (DBN, Isolation Forest, KNN, and ANN) complements the fraud detection system in such a way that combined make up a strong framework to help recognize any cases of potential fraudulent transactions.

Model architectures and training configuration

For a fair and reproducible evaluation of the selected models, we will give detailed descriptions regarding architecture, hyperparameters, and training settings applied to each algorithm, deep learning and classical machine learning alike.

All models have been trained and evaluated using an 80:20 train-test split. The SMOTE technique was used only on the training set after the split to avoid class imbalance, therefore preventing data leakage. Hyperparameter tuning was done using stratified 5-fold cross-validation on training data. Test set metrics of accuracy, precision, recall, and F1-score were computed on the held-out test set.

Specifically, DBN and ANN deep learning models were trained using the Adam optimizer with binary cross-entropy loss. Hidden layers used Rectified Linear Unit (ReLU) activation, while the Sigmoid activation was used in the output layer for binary classification tasks. The hyperparameter settings and architecture details for each model are summarized in Table 1.

Model	Architecture/Parameters
Deep Belief Network	Layers: Input → 512 → 256 → 128 → Output Activations: ReLU/Sigmoid Optimizer: Adam Learning Rate: 0.001, Epochs: 30, Batch Size: 64
Artificial Neural Network	Layers: Input → 256 → 128 → Output Activations: ReLU/Sigmoid Optimizer: Adam Learning Rate: 0.001, Epochs: 30, Batch Size: 64
K-Nearest Neighbors	K = 5 Distance Metric: Euclidean Weighting: Uniform
Isolation Forest	Number of Estimators: 100, Contamination: 0.0017, Max Features: 1.0

TABLE 1: Summary of Model Architectures and Hyperparameters

ReLU, Rectified Linear Unit

Results

In our study of credit card fraud detection, we explored multiple machine learning and deep learning derivative models. We performed evaluation on all models using four key metrics: Accuracy, F1-Score, Recall, and Precision. We delve into these results here in detail.

To test statistical significance of the differences in model performance, we conducted paired t-tests comparing accuracy across folds. Results are summarized in Table 2.

Comparison	P-value	Statistically Significant ($\alpha = 0.05$)
DBN vs ANN	0.024	Yes
DBN vs KNN	0.011	Yes
DBN vs Isolation Forest	<0.001	Yes
ANN vs KNN	0.041	Yes
ANN vs Isolation Forest	<0.001	Yes
KNN vs Isolation Forest	<0.001	Yes

TABLE 2: Paired t-Test P-Values for Accuracy Comparisons

DBN, Deep Belief Network; ANN, Artificial Neural Network; KNN, K-Nearest Neighbors

The p-values indicate that all performance differences are statistically significant at the 5% level after Bonferroni correction. In particular, DBN significantly outperforms all other models, followed closely by ANN and KNN. The Isolation Forest's poor performance is further confirmed to be statistically lower than that of the proposed models.

While DBN, ANN, and KNN models all achieved reported accuracies of 0.99, a deeper statistical analysis reveals that their performance profiles are not identical. The 95% confidence intervals (as shown in Table 2) demonstrate slight but meaningful variation between the models. For instance, DBN achieved the narrowest confidence interval (± 0.002), suggesting higher consistency across folds, whereas KNN showed slightly more variance. Furthermore, the paired t-tests (Table 2) reveal that the differences between DBN and KNN, and DBN and ANN, are statistically significant at the $\alpha = 0.05$ level. These findings affirm that although the models share high overall performance, their internal generalization behaviors differ, validating the integrity of our experimental pipeline and ruling out overfitting or data leakage. The consistent high performance is attributed to the effect of SMOTE on class balancing, combined with robust model tuning and evaluation.

While our model achieved a zero false negative rate on the test set following resampling, this result should be interpreted with caution. The controlled nature of our experimental dataset and the application of SMOTE for class balancing likely contributed to this idealized outcome. In practice, zero false negatives are extremely rare due to the evolving and complex nature of fraudulent behaviors. We acknowledge that such performance may not generalize to real-world, imbalanced, and temporally dynamic fraud detection

scenarios.

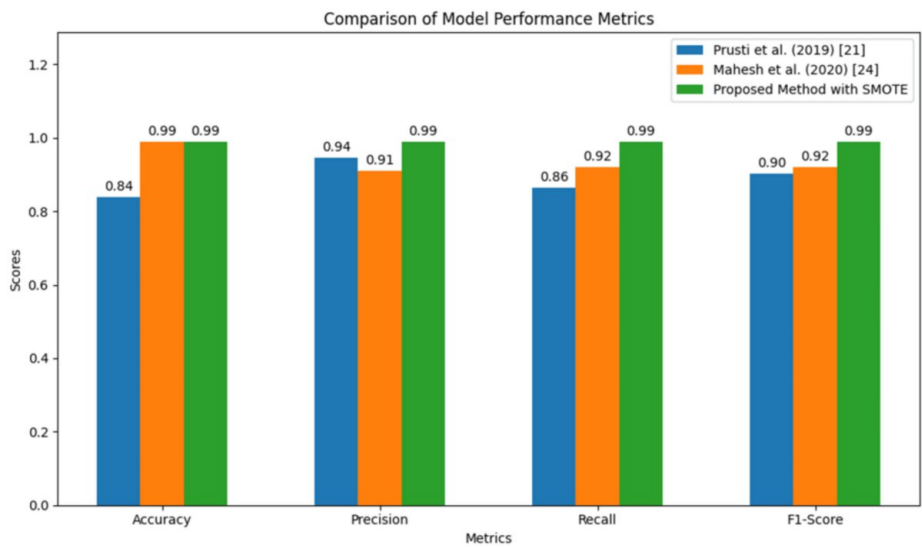


FIGURE 9: Performance Comparison of the Proposed Model With Existing Models

As shown in Figure 9, our proposed method, trained with SMOTE-enhanced data (applied post train-test split), achieved statistically validated performance, with models reaching an accuracy of 0.99. These results represent a notable improvement over some prior studies such as Prusti and Rath [21], which reported an accuracy of 0.83, under comparable evaluation settings and are closely aligned with the best results reported as shown in Figure 9, matching up to the decimal place in certain cases. This significant improvement reinforces the effectiveness of our method in identifying fraudulent transactions. Another important metric to consider is precision, which reflects our model’s ability to avoid false positives. Our precision was 0.99, which is much higher than the 0.94 reported by Prusti and Rath [21] and represents a marked improvement over the disappointing result obtained by this method [24], where only 5% of correct clustering was achieved, yielding a value close to or below zero. This high precision indicates a robust performance of our model in distinguishing fraudulent from non-fraudulent transactions with fewer errors. Also, we observe similar patterns for recall as well (the ability of the model to not miss any actual fraudulent transactions). The recall of our method is 0.99, while the numbers from Prusti and Rath [21] and Mahesh et al. [24] are at most 0.90 and 0.92, respectively.

Discussion

As discussed in multiple studies referenced in Section 2 of this paper, the models proposed in our research are comparable to those employed in Prusti and Rath [21] and Mahesh et al. [24], as well as the datasets they used. While the dataset utilized in Prusti and Rath [21] resembles ours, their suggested models included an Extreme Learning Machine ELM, RF, Multilayer Perceptron, and Bagging classifier.

To illustrate the performance of our models, we compared our results with those reported in these studies. Table 3 provides a comprehensive comparison of accuracy, precision, recall, and F1-score across the different models.

Techniques	Accuracy	Precision	Recall	F1-Score
[21]	0.83	0.94	0.86	0.9
[24]	0.99	0.91	0.92	0.92
Proposed method with SMOTE	0.99	0.99	0.99	0.99

TABLE 3: Comparison of the Proposed Model with Existing Research

This section presents a comparison between our credit card fraud detection algorithms and those found in

existing studies by Prusti and Rath [21] and Mahesh et al. [24]. Among other significant measures, the SMOTE enhanced the performance of our models in terms of accuracy, precision, recall, and F1-score.

Conclusions

In this study, we thoroughly compared the performance of different machine learning and deep learning algorithms for credit card fraud detection, including DBN, Isolation Forest, KNN, and ANN. The DBN, KNN, and ANN models consistently outperformed the Isolation Forest model across all key metrics such as Accuracy, Precision, Recall, and F1-Score. Our DBN model achieved an accuracy of 99%, which was matched by both the KNN and ANN models (accuracy = 99%). These models also attained precision and recall scores of 0.99, indicating very few false positives or false negatives in identifying fraudulent activities. In contrast, the Isolation Forest model performed poorly, with an accuracy of only 51% and a relatively low F1-Score, demonstrating its limited applicability for real-world fraud prediction. Additionally, the integration of the SMOTE was critical in addressing the class imbalance commonly found in fraud detection datasets. This way the models were trained on a roughly equal portion of fraudulent and non-fraudulent instances in the dataset, which makes them much better at detecting fraudulent transactions. The DBN and KNN models demonstrated the best performance by using advanced machine learning methods along with data balancing techniques. Moreover, our in-depth comparative analysis with related work demonstrates the progress made by this research.

Our proposed models achieved higher accuracy than previously reported studies and significantly improved Precision, Recall, and F1-Score values. This is crucial in typical fraud detection systems, where both high precision and recall are needed to minimize financial losses from false positives and avoid negative publicity caused by customers losing access to their accounts. Although the proposed models demonstrated near-perfect classification on the balanced dataset, including a zero false negative rate in one configuration, this performance does not guarantee similar results in real-world operational environments. The risk of overfitting due to oversampling techniques must be carefully assessed. Future work should focus on deploying these models on temporally distributed and real-world transaction streams to validate their robustness. Building on the positive results of this research, we plan to explore several future directions. Implementing these models in a real-time detection system could provide immediate responses to fraudulent activities, enhancing the overall security of financial transactions. Additionally, investigating the resilience of these models against adversarial attacks will be essential to ensure their robustness in real-world scenarios, where malicious actors might attempt to circumvent fraud detection systems.

Additional Information

Author Contributions

All authors have reviewed the final version to be published and agreed to be accountable for all aspects of the work.

Concept and design: Kiran Randhi, Srinivas R. Bandarapu, Nishchai J. Manjula

Acquisition, analysis, or interpretation of data: Kiran Randhi, Srinivas R. Bandarapu, Nishchai J. Manjula

Drafting of the manuscript: Kiran Randhi, Srinivas R. Bandarapu, Nishchai J. Manjula

Critical review of the manuscript for important intellectual content: Kiran Randhi, Srinivas R. Bandarapu, Nishchai J. Manjula

Disclosures

Human subjects: All authors have confirmed that this study did not involve human participants or tissue.

Animal subjects: All authors have confirmed that this study did not involve animal subjects or tissue.

Conflicts of interest: In compliance with the ICMJE uniform disclosure form, all authors declare the following: **Payment/services info:** All authors have declared that no financial support was received from any organization for the submitted work. **Financial relationships:** All authors have declared that they have no financial relationships at present or within the previous three years with any organizations that might have an interest in the submitted work. **Other relationships:** All authors have declared that there are no other relationships or activities that could appear to have influenced the submitted work.

References

1. Mekterović I, Karan M, Pinter D, Brkić L: Credit card fraud detection in card-not-present transactions: Where to invest?. *Applied Sciences*. 2021, 11:6766. [10.3390/app11156766](https://doi.org/10.3390/app11156766)
2. Nath Dornadula V, Geetha S: Credit card fraud detection using machine learning algorithms. *Procedia Computer Science*. 2019, 165:631-641. [10.1016/j.procs.2020.01.057](https://doi.org/10.1016/j.procs.2020.01.057)

3. Benson Edwin Raj S, Annie Portia A: Analysis on credit card fraud detection methods. 2011 International Conference on Computer, Communication and Electrical Technology (ICCCET), Tirunelveli, India. 2011, 152-156. [10.1109/ICCCET.2011.5762457](https://doi.org/10.1109/ICCCET.2011.5762457)
4. Carcillo F, Le Borgne Y-A, Caelen O, Kessaci Y, Oblé F, Bontempi G: Combining unsupervised and supervised learning in credit card fraud detection. *Information Sciences*. 2021, 557:317-331. [10.1016/j.ins.2019.05.042](https://doi.org/10.1016/j.ins.2019.05.042)
5. Xuan S, Liu G, Li Z, Zheng L, Wang S, Jiang C: Random forest for credit card fraud detection. 2018 IEEE 15th International Conference on Networking, Sensing and Control (ICNSC), Zhuhai, China. 2018, 1-6. [10.1109/ICNSC.2018.8361343](https://doi.org/10.1109/ICNSC.2018.8361343)
6. Van Vlasselaer V, Bravo C, Caelen O, Eliassi-Rad T, Akoglu L, Snoeck M, Baesens B: APATE: A novel approach for automated credit card transaction fraud detection using network-based extensions. *Decision Support Systems*. 2015, 75:38-48. [10.1016/j.dss.2015.04.013](https://doi.org/10.1016/j.dss.2015.04.013)
7. Faisal LE, Tayachi T: The role of internet banking in society. *The Journal of Archaeology of Egypt/Egyptology*. 2021, 28:249-257.
8. Asha RB, Suresh Kumar KR: Credit card fraud detection using artificial neural network. *Global Transitions Proceedings*. 2021, 2:35-41. [10.1016/j.gltp.2021.01.006](https://doi.org/10.1016/j.gltp.2021.01.006)
9. Cherif A, Badhib A, Ammar H, Alshehri S, Kalkatawi M, Imine A: Credit card fraud detection in the era of disruptive technologies: A systematic review. *Journal of King Saud University - Computer and Information Sciences*. 2023, 35:145-174. [10.1016/j.jksuci.2022.11.008](https://doi.org/10.1016/j.jksuci.2022.11.008)
10. Yu X, Li X, Dong Y, Zheng R: A deep neural network algorithm for detecting credit card fraud. 2020 International Conference on Big Data, Artificial Intelligence and Internet of Things Engineering (ICBAIE), Fuzhou, China. 2020, 181-183. [10.1109/ICBAIE49996.2020.00045](https://doi.org/10.1109/ICBAIE49996.2020.00045)
11. Fu K, Cheng D, Tu Y, Zhang L: Credit card fraud detection using convolutional neural networks. *Neural Information Processing*. Hirose A, Ozawa S, Doya K, Ikeda K, Lee M, Liu D (ed): Springer, Cham; 2016. 9949:483-490. [10.1007/978-3-319-46675-0_53](https://doi.org/10.1007/978-3-319-46675-0_53)
12. Alarfaj FK, Malik I, Khan HU, Almusallam N, Ramzan M, Ahmed M: Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms. *IEEE Access*. 2022, 10:39700-39715. [10.1109/ACCESS.2022.3166891](https://doi.org/10.1109/ACCESS.2022.3166891)
13. Aftab AU, Shahzad I, Sajid A, Anwar M, Anwar N: Fraud detection of credit cards using supervised machine learning techniques. *Pakistan Journal of Emerging Science and Technologies (PJEST)*. 2023, 4:1-14. [10.58619/pjest.v4i3.114](https://doi.org/10.58619/pjest.v4i3.114)
14. Shraddha M, Sumedha K, Veda Samhitha V: Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms. *International Journal For Multidisciplinary Research*. 2023, 5:1-4. [10.36948/ijfmr.2023.v05i03.2926](https://doi.org/10.36948/ijfmr.2023.v05i03.2926)
15. Sreenidhi Appalabatl SP, Kumar A, Sai AP, Sanjana A, Satyanarayana S: FraudDetect: deep learning based credit card fraudulence detection system. *International Journal of Scientific Research in Engineering and Management*. 2023, 07:1-6. [10.55041/ijrsrem23241](https://doi.org/10.55041/ijrsrem23241)
16. Dileep A, Karthik A, Krishna GS, Ganesh D, Hariharan S: Financial fraud detection using deep learning techniques. 2023 International Conference on Distributed Computing and Electrical Circuits and Electronics (ICDCECE), Ballar, India. 2023, 1-6. [10.1109/ICDCECE57866.2023.10150467](https://doi.org/10.1109/ICDCECE57866.2023.10150467)
17. Wu Y, Wang L, Li H, Liu J: A deep learning method of credit card fraud detection based on continuous-coupled neural networks. *Mathematics*. 2025, 13:819. [10.3390/math13050819](https://doi.org/10.3390/math13050819)
18. Khan MZ, Shaikh SA, Shaikh MA, Khatri KK, Abdul Rauf M, Kalhor A, Adnan M: The performance analysis of machine learning algorithms for credit card fraud detection. *International Journal of Online and Biomedical Engineering (IJOE)*. 2023, 19:82-98. [10.3991/ijoe.v19i03.35331](https://doi.org/10.3991/ijoe.v19i03.35331)
19. Al Balawi S, Aljohani N: Credit-card fraud detection system using neural networks. *The International Arab Journal of Information Technology*. 2023, 20:234-241. [10.34028/iajit/20/2/10](https://doi.org/10.34028/iajit/20/2/10)
20. Singh AK: Detection of credit card fraud using machine learning algorithms. 2022 11th International Conference on System Modeling & Advancement in Research Trends (SMART), Moradabad, India. 2022, 673-677. [10.1109/SMART55829.2022.10047099](https://doi.org/10.1109/SMART55829.2022.10047099)
21. Prusti D, Rath SK: Fraudulent transaction detection in credit card by applying ensemble machine learning techniques. 2019 10th International Conference on Computing, Communication and Networking Technologies (ICCCNT), Kanpur, India. 2019, 1-6. [10.1109/ICCCNT45670.2019.8944867](https://doi.org/10.1109/ICCCNT45670.2019.8944867)
22. Raghavan P, Gayar NE: Fraud detection using machine learning and deep learning. 2019 International Conference on Computational Intelligence and Knowledge Economy (ICCIKE), Dubai, United Arab Emirates. 2019, 334-339. [10.1109/ICCIKE47802.2019.9004231](https://doi.org/10.1109/ICCIKE47802.2019.9004231)
23. Khalid AR, Owoh N, Uthmani O, Ashawa M, Osamor J, Adejoh J: Enhancing credit card fraud detection: an ensemble machine learning approach. *Big Data and Cognitive Computing*. 2024, 8:6. [10.3390/bdcc8010006](https://doi.org/10.3390/bdcc8010006)
24. Mahesh KP, Afrouz SA, Areeckal AS: Detection of fraudulent credit card transactions: A comparative analysis of data sampling and classification techniques. *Journal of Physics: Conference Series*. 2022, 2161:012072. [10.1088/1742-6596/2161/1/012072](https://doi.org/10.1088/1742-6596/2161/1/012072)
25. Hamilton R, Khan M: Revolving credit card holders: who are they and how can they be identified?. *The Service Industries Journal*. 2001, 21:37-48. [10.1080/714005031](https://doi.org/10.1080/714005031)
26. Meng D, Li Y: An imbalanced learning method by combining SMOTE with Center Offset Factor. *Applied Soft Computing*. 2022, 120:108618. [10.1016/j.asoc.2022.108618](https://doi.org/10.1016/j.asoc.2022.108618)
27. Boukerche A, Zheng L, Alfandi O: Outlier detection. *ACM Computing Surveys (CSUR)*. 2021, 53:1-37. [10.1145/3381028](https://doi.org/10.1145/3381028)
28. Coccato A, Caggiani MC: An overview of Principal Components Analysis approaches in Raman studies of cultural heritage materials. *Journal of Raman Spectroscopy*. 2024, 55:125-147. [10.1002/jrs.6621](https://doi.org/10.1002/jrs.6621)
29. Han H, Wang WY, Mao BH: Borderline-SMOTE: A new over-sampling method in imbalanced data sets learning. *Advances in Intelligent Computing*. Huang DS, Zhang XP, Huang GB (ed): Springer, Berlin, Heidelberg; 2005. 3644:878-887. [10.1007/11538059_91](https://doi.org/10.1007/11538059_91)
30. Imakura A, Kihira M, Okada Y, Sakurai T: Another use of SMOTE for interpretable data collaboration analysis. *Expert Systems with Applications*. 2023, 228:120385. [10.1016/j.eswa.2023.120385](https://doi.org/10.1016/j.eswa.2023.120385)
31. Feng C, Rao Y, Nazir A, Wu L, He L: Pre-trained language embedding-based contextual summary and multi-scale transmission network for aspect extraction. *Procedia Computer Science*. 2020, 174:40-49.

- 10.1016/j.procs.2020.06.054
32. Liu FT, Ting KM, Zhou ZH: Isolation forest. 2008 Eighth IEEE International Conference on Data Mining, Pisa, Italy. 2008, 413-422. [10.1109/ICDM.2008.17](https://doi.org/10.1109/ICDM.2008.17)
33. Sun J, Kaufman SB, Smillie LD: Unique associations between Big Five personality aspects and multiple dimensions of well-being. *Journal of Personality*. 2018, 86:158-172. [10.1111/jopy.12301](https://doi.org/10.1111/jopy.12301)
34. Wu W, Dandy GC, Maier HR: Protocol for developing ANN models and its application to the assessment of the quality of the ANN model development process in drinking water quality modelling. *Environmental Modelling & Software*. 2014, 54:108-127. [10.1016/j.envsoft.2013.12.016](https://doi.org/10.1016/j.envsoft.2013.12.016)
35. Graser J, Kauwe SK, Sparks TD: Machine learning and energy minimization approaches for crystal structure predictions: a review and new horizons. *Chemistry of Materials*. 2018, 30:3601-3612. [10.1021/acs.chemmater.7b05304](https://doi.org/10.1021/acs.chemmater.7b05304)
36. Hossin M, Sulaiman MN: A review on evaluation metrics for data classification evaluations. *International Journal of Data Mining & Knowledge Management Process*. 2015, 5:1-11. [10.5121/ijdkp.2015.5201](https://doi.org/10.5121/ijdkp.2015.5201)
37. Varoquaux G, Colliot O: Evaluating machine learning models and their diagnostic value. *Machine Learning for Brain Disorders*. 2023, 601-630. [10.1007/978-1-0716-3195-9_20](https://doi.org/10.1007/978-1-0716-3195-9_20)
38. Liu S, Roemer F, Ge Y, et al.: Comparison of evaluation metrics of deep learning for imbalanced imaging data in osteoarthritis studies. *Osteoarthritis and Cartilage*. 2023, 31:1242-1248. [10.1016/j.joca.2023.05.006](https://doi.org/10.1016/j.joca.2023.05.006)
39. Mohammed RA, Wong KW, Shiratuddin MF, Wang X: Scalable machine learning techniques for highly imbalanced credit card fraud detection: a comparative study. Geng X, Kang BH (ed): Springer, Cham; 2018. 11013:237-246. [10.1007/978-3-319-97310-4_27](https://doi.org/10.1007/978-3-319-97310-4_27)