

Systematic Literature Survey on Credit Card Fraud Detection Techniques

Bhukya Dharma¹ and Dr.D. Latha²

¹Research Scholar, Department of Computer Science and Engineering, Adikavi Nannaya University, Rajamahendravaram, AP, India

Email: dharmaknu9@rediffmail.com

²Asst. Prof., Department of Computer Science and Engineering, Adikavi Nannaya University, Rajamahendravaram, India

Abstract—Credit cards (CC) are an important component of the present wealth. It develops into an important component of household, business, and international activities. While using credit cards responsibly and could result in major benefits, fraudulent behaviour can also significant credit and financial damage. Every year, fraudulent credit card transactions result in significant financial losses for both people and organizations. To perform these transactions, fraudsters are always looking for new technologies and methods. A significant factor affecting the increased use of electronic payments is the detection of fraudulent transactions. As a result, there is a need for effective and fast methods of identifying fraud in credit card transactions. This analysis presents Literature survey on Credit Card fraud detection techniques. In order for consumers to utilize e-banking safely and conveniently, the major goal is to protect credit card transactions. For fraud detection approaches, both positive and negative characteristics are given. It also shows its limitations and explains the state-of-the-art approaches that are currently being deployed to counter against these attacks.

Index Terms— Credit Card, fraudulent, transactions, e-banking.

I. INTRODUCTION

Credit card use is becoming common in developing countries nowadays. Users use it to make purchases, manage tabs, and conduct online exchanges. It provides some interesting information, like the ease of purchasing, the keeping of customer loan payback records, the security of purchases, and so on [1]. Still, billion-dollar losses to MasterCard extortion result from the widespread use of Visas and the lack of viable security frameworks. It is challenging to estimate the losses precisely because credit card companies sometimes refuse to disclose such reality [2]. However, some information regarding the financial losses brought on by Visa deception is publicly accessible. The primary source of information for the fraud detection procedure is credit card transactions. The detection procedure is significantly impacted by a number of important transaction data features. Data hugeness and data imbalance are the main characteristics of credit card transaction data that influence the prediction process.

Modern technologies and worldwide communication have dramatically increased fraud. The use of a credit card without authorization or with the intent to defraud is known as "fraud" and occurs when someone other than the account owner makes an unwanted usage. When someone uses another person's credit card for personal expenses without the owner's knowing or the credit card issuer's information, it is technically called card fraud. The theft

of the cardholder's personal data by the fraudsters is one of the primary elements of credit card fraud, which subsequently use it to send SMS messages or place calls. There are software programmes that thieves might use to commit credit card fraud additionally. Cardholder information need to remain confidential. Websites that use phishing techniques, credit cards that have been lost or stolen, fake cards, cards that have been intercepted, etc. are some of the ways credit card information may be stolen. To stop this abuse, necessary measures can be taken, to reduce and prevent similar fraudulent practices in the future, it is possible to study their behaviour. As a result of the problems with fraud, there is now a greater need for solutions to detect credit card fraud. Since fraudulent information is difficult, obtaining it might be challenging. Monitoring and analyzing user behaviour over a range of users helps to identify fraud by estimating the likelihood of detecting undesirable behaviours. They need to understand the many technologies, algorithms, and types are successfully used to identify credit card fraud [3]. Identification of credit card frauds, transaction records are examined. The credit card number, transaction date, recipients, and transaction amount are among the many details, make up the majority of transaction data. Automatic methods are essential because transaction datasets usually have a large number of samples, several dimensions, online updates, because it is not always possible or easy for a human analyst to spot fraudulent tendencies in data. Additionally, the cardholder is not trustworthy it comes to reporting card theft, loss, or fraudulent usage. The majority of the time now, neural networks are used to solve non-linear classification problems in the actual world. A model known as a "neural network" was created to closely resemble the organization of a biological neural network. Neural networks can also interact with information from the outside world and provide feedback [4].

Credit card transactions can be done both physically and virtually, or CNP (Card not Present). A physical swipe of the card is necessary when using one. In contrast, a virtual card has information needed to swipe a physical card, including the CVV (Card Verification Value) number, cardholder name, password, security question, and online banking. Fraudsters attempt to deceive the system in Card Not Present (CNP) fraud by pretending to be another person [5]. Genuine mail-order and online retailers are targeted by fraudsters who sell and send products through the mail and the internet. By skimming, they are acquiring private information from a credit card used in an otherwise typical transaction by another individual. A small device called a "skimmer" is used to collect and store a lot of the victim's personal information. In phishing, scammers can use a number of methods, such as websites that pretend to be banks or payment systems, to mislead customers into giving them their credit card information. When a card is lost or stolen, there is a probability that the thief may use it to make an unauthorized purchase before the cardholder blocks account.

In order to identify fraudulent transactions even before end users report fraud, automated machine learning-based systems called Fraud Detection Systems (FDS) are used by credit card providers. Such a technology attempts to identify fraudulent transactions prior to their entry into databases and so prevent fraud from occurring. False detections would be less frequent in a perfect FDS, which would disrupt legitimate transactions and inconvenience users. Every bank has a responsibility to review transaction information and determine whether or not they are fraudulent. This is a difficult and time-consuming task in a time when time is money. Data science and machine learning algorithms can help us with this important issue [6].

Using a number of algorithmic models, A kind of artificial intelligence called "machine learning" allows computers to autonomously learn new abilities and without human involvement learn from experience. Its goal is to make predictions about the future as accurate as possible. Machine learning differs significantly from conventional computing techniques, where systems are explicitly created to calculate or solve a problem. The study of the input data that is used to train a model is known as machine learning. In order to forecast unknown outcomes, the model recognizes numerous patterns in the incoming data. Supervised, unsupervised, and reinforcement learning are the three main categories of machine learning [7].

When labels are provided for both the input and the output, a machine learning technique known as supervised learning is used. Through utilizing the output and input labeled data for training, the supervised model detects patterns in the input data. The three supervised learning algorithms that are most often utilized are decision tree, logistic regression, and support vector machines. When a machine learns on its own without any supervision, it is said to be using unlabeled data. With regard to the unlabeled data, the machine looks for patterns and responds correctly. The three unsupervised learning algorithms that are most frequently utilized are K-means clustering, hierarchical clustering, and apriori algorithm. In the literature, several algorithms based on both strategies have been suggested.

II. LITERATURE SURVEY

Anish Mahajan, Vivek Singh Baghel, Ramkumar Jayaraman, et. al. [8] presents Credit Card Fraud Detection using Logistic Regression with Imbalanced Dataset. A major problem for financial organizations and their clients is credit card fraud. In this analysis, they explore the application of a machine learning method called logistic regression to identify credit card fraud in an unbalanced dataset where only a small number of transactions are fraudulent. To overcome the problem of imbalanced data, they utilize undersampling of the majority class and oversampling of the minority class. The selection of an algorithm for the model is dependent on the performance of each algorithm under classification. The selection of the wrong algorithm can result in overfitting or underfitting. Logistic Regression is a universally studied algorithm and can be used in complex situations to perform binary classifications. In this paper, the algorithm needs to classify the inputs into labeled classes. It can also be represented as mapping a vector into labeled classes. This paper specifies the defects of an imbalanced dataset and provides an algorithm for binary classification that can help tackle the imbalanced classification and provide the accuracy of the aforementioned algorithm in detecting fraudulent instances in the credit card transactions dataset.

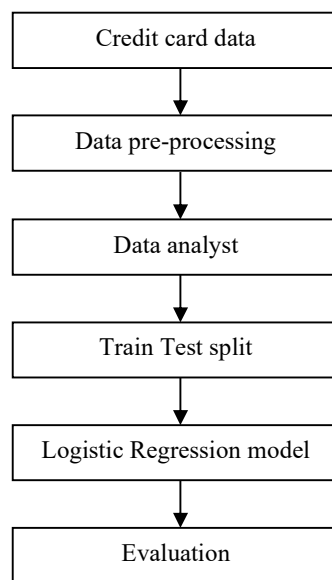


Fig. 1: Architecture Diagram

The first step of the implementation includes collecting the dataset for the model. After getting an adequate dataset, data analysis can be conducted on the dataset to further explore the dataset. This dataset, which comprises of 2,84,807 observations and 492 fraudulent transactions, this covers the transactions that took place over an average of two days. The collection of data has significant imbalances. The data pre-processing is an important step for the development of the model. Exploratory data analysis is required to understand the information in the dataset. This helps to understanding the dataset is a perfect sample of the population and is well-balanced. This step also helps us figure out the relationship between different attributes that can help us build a more advanced and accurate model.

Under this split, the dataset is divided into a training database with the goal of constructing the model (classifier) and a testing database with the goal of evaluating the model. In this model, the split is conducted through cross-validation which divides the database into k parts and iterates k times. They will split the dataset for $k = 10$ folds. The classifier then models and evaluates the classifier using the training data. Binary classifier is built and the logistic regression algorithm is applied. This is an advanced algorithm as it can classify the dataset that is widely distributed in space. Logistic regression is very efficient to train. It can also be used for multiple classifications.

The cross-validation technique gives the model merit for using the complete dataset for both training as well as for testing. The evaluation is grouped with the accuracy of the model. In general, the evaluation can be accomplished using a confusion matrix which is formed from the following: True Positive, False Positive, True Negative, False Negative. The model or the classifier built on the basis of the logistic regression gives an

accuracy of 94%. Logistic regression is a machine learning algorithm that categorizes the dataset into either Yes or No. Logistic regression serves as an advanced form of linear regression algorithm. Linear regression fails with the classification of the points that are highly distributed in the dataset. Primary ability of linear regression cannot exceed classifying the data points that overlap the line. This inability of linear regression can be easily overcome by logistic regression with its S-shaped curve. It fits an S-shaped graph line which gives the most value as 0 or 1. The graph can give the probability of a transaction is fraudulent or not.

Alavikunhu Panthakkan, Majida Appathil, Najiya Valappil, Wathiq Mansoor, Seema Verma, and Hussain Al-Ahmad, et. al. [9] compares the effectiveness of machine learning systems for detecting credit card fraud. To preserve credit card security, a suspicious transaction must be identified and stopped. Due to the availability of historical information, the use of machine learning algorithms has made it feasible to identify unusual behaviour in transactional data. The study aims to identify fraudulent activities by utilizing six different machine learning techniques, including Logistic Regression, Decision Tree, SVM (Support Vector Machine), K-Nearest Neighbour, XGBoost, and Random Forest, it's to create a balanced dataset from an existing dataset.

The first phase in the process includes developing a balanced datasets in order to provide a clearly specified outcome. Consistency and credibility are required for the procedures examined the selection of data sets. To avoid bias when the method is tested, the dataset must be filtered. Moreover, previous outcomes from which inferences may be added to the tests that have been completed on the algorithm covered in this study. 284807 legal transactions and 473 fraud transactions make up the dataset. Due to the fact that there are much more legitimate transactions than fraudulent ones, this data may provide an accuracy bias in the results. Three genuine sets of datasets are taken into consideration, and six algorithms are trained, for the purpose to look into the output accuracies are affected by the ratio of fraudulent to valid transactions. The following cases are taken into consideration:

Case 1: There are 473 fraudulent transactions out of 47300 legal transactions, or a ratio of 1:100.

Case 2: There are 473 fraudulent transactions and 4730 legitimate ones, resulting in a ratio of 1:10 between both devices.

Case 3: There is a 1:1 ratio of fraudulent to legitimate transactions when both types of transactions are represented.

The purpose of the study is to indicate that balanced datasets are preferable to unbalanced ones on which prior studies have been conducted. This paper also seeks to provide the best machine learning-based answer for credit card fraud identification. The following techniques are trained on a dataset with this as their top priority. This study utilizes the Decision tree (DT), KNN (K-Nearest Neighbour), Logistic regression (LR), Random forest (RF), Support vector machine (SVM), and XGBoost as six methods. The ratio of fraudulent transactions compared to valid transactions was used to construct three sets of datasets. The comparison matrices are depending on the F1-score of genuine and fraudulent transactions as well as precision, recall, and accuracy.

TABLE I: COMPARISON OF ALGORITHMS WITH DATASET RATIO 1:100

		Precision	Recall	F1-score	Accuracy	Overall Accuracy
DT	Valid	100	100	100	100	93.7
	Fraud	93	82	87	87.5	
KNN	Valid	100	100	100	100	92
	Fraud	94	77	84	84	
LR	Valid	100	100	100	100	92
	Fraud	90	78	83	84	
SVM	Valid	100	100	100	100	94
	Fraud	96	80	87	88	
RF	Valid	100	100	100	100	94.75
	Fraud	95	84	89	89.5	
XGBOOST	Valid	100	100	100	100	94.75
	Fraud	95	84	89	98.5	

The tables show that an imbalanced dataset produced training results with lower accuracy and recall, which is undesirable. In contrast, an unbalanced dataset would perform significantly better in terms of accuracy because there would be more true positives. The system's overall accuracy was examined, and on a balanced dataset, the XGBoost method had the greatest accuracy. It was accurate to roughly 93.5% while maintaining good precision and recall outcomes.

Sumanth C.H, Kalyan, Bolisetti Ravi, Pokala Pavan S Balasubramani., et. al. [10] The analysis of machine learning methods for detecting credit card fraud is explained. Real-world fraud detection system deployment is

TABLE II: COMPARISON OF ALGORITHMS WITH DATASET RATIO 1:10

		Precision	Recall	F1-score	Accuracy	Overall Accuracy
DT	Valid	98	100	99	99	95.7
	Fraud	99	86	92	92.5	
KNN	Valid	99	100	99	99.5	97
	Fraud	99	90	95	94.5	
LR	Valid	99	100	99	99	96.5
	Fraud	96	92	94	94	
SVM	Valid	99	100	99	99.5	97
	Fraud	99	90	94	94.5	
RF	Valid	99	100	99	99.5	97
	Fraud	99	90	95	94.5	
XGBOOST	Valid	99	100	100	99.5	97
	Fraud	99	92	96	95.5	

TABLE III: COMPARISON OF ALGORITHMS WITH DATASET RATIO 1:10

		Precision	Recall	F1-score	Accuracy	Overall Accuracy
DT	Valid	87	97	91	92	92
	Fraud	97	87	92	92	
KNN	Valid	87	97	91	92	92
	Fraud	97	87	92	92	
LR	Valid	89	94	92	91.5	92
	Fraud	95	90	92	92.5	
SVM	Valid	88	95	92	91.5	92
	Fraud	96	89	92	92.5	
RF	Valid	88	93	91	91.5	91
	Fraud	94	89	91	91.5	
XGBOOST	Valid	89	95	92	92	92.5
	Fraud	96	90	93	93	

not without its difficulties. In real-world scenarios, a large number of payment demands is automatically scanned to identify which transactions should be authorized. All legally permitted transactions are examined by machine learning algorithms, which flag those that appear suspicious. The gathering of data is one of the most important tasks involved in creating a machine learning model. A collection of specific factors are used to collect task-related data, which is then analyzed and used to provide a beneficial outcome. The features v1 (vector 1) through v27 (vector 27) of the Kaggle data set show the number of transactions as well as other characteristics like net amount, and they classify the amount as being identifiable or not based on this property. If the value of a class attribute is zero, fraud has not been found; if it is one, fraud has been discovered to some extent. There have been 85275 legitimate transactions, 117 fraudulent transactions have been discovered, and the remaining 30 and 21 transactions may or may not be real.

Pre-processing involves the preparation of data, which might include cleaning, transforming, and selecting data. The following are some instances of data cleaning: filling in missing numbers, removing noise from noisy data, locating and removing outliers, resolving conflicts. In order to improve data quality, data transformations including smoothing, aggregating, generalization, and transformation can be used. They can select the most useful information for described system using a collection of algorithms or functions called data selection.

A method for evaluating a model's accuracy is train/test. The term "train/test" refers to the division of a data set into two sets, one for training and the other for testing. In order to detect different credit card frauds based on transactions, we analyze three machine learning algorithms in this paper: Navie Bayes, SVM, and DNN (Deep Neural Network). The support vector machine is very effective with high-dimensional data. Current layer neurons receive activations from neurons in the preceding layer and use those activations to do a simple computation. The probability-based classification technique establishes the probabilities of the target classes and the test data, that can identify credit card fraud using the Navies Bayes classifier. The ability to identify and stop fraudulent transactions. Uneven behaviours are an indication of illicit activity. Accuracy is used parameter in this study for performance evaluation. DNN algorithm achieves highest accuracy as 99% than other algorithms NB (90%) and SVM (97%).

Prabhat Singh, Shivam Singh, Vishesh Chauhan, Priya Agarwal, Shrey Agrawal, et. al. [11] offers a methodology for utilizing a machine learning algorithm to identify credit card fraud. Fraud is defined as the use of one's resources for self-improvement through fraud. Described system's main objective is to evaluate multiple

approaches and identify the best method, which enables to reduce fraud and support future system upgrades. As a defense against illegal activity, creating a way to identify credit card theft is essential. Nowadays, the great majority of transactions take place online, which means that online payment systems like credit cards are also used. The organization and the client both benefit from this strategy. Fig. 2, below represents the architectural view of the credit card fraud detection system. It is possible to implement the project concept using a number of machine learning approaches, but for now, they focus on four Python-based machine learning models.

The following are the implementation steps, which range from data import to model evaluation:

1. Importing packages: The main packages to be utilized in this project are scikit-learn, numpy, and pandas.
2. Selection of Dataset: They are utilizing a dataset from Kaggle that was created in 2013 and comprises credit card transactions from European customers.
3. Data Pre-processing: In this stage, they calculate the fraud cases (492), the non-fraud cases (284315), other cases, the proportion of fraud cases (0.17%) in all transactions. These numbers show the dataset's imbalance. Furthermore, machine learning techniques must be used to normalize information.
4. Using the train_test_split technique, they must choose features and divide dataset after normalizing information.
5. Development of machine learning models, including the Random forest (RF), KNN, Logistic regression (LR) and SVM models.
6. Evaluation: Models are assessed using their accuracy score and F1 score. These two are the most fundamental assessment techniques and are frequently used to assess models.

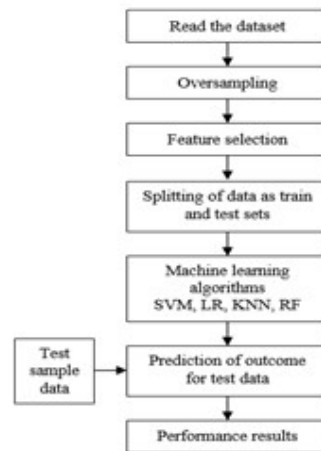


Fig. 2: Proposed System Architecture System

TABLE IV: COMPARATIVE PERFORMANCE

Algorithms	Accuracy	F1-score
SVM	99.92	77.01
KNN	99.95	85.71
LR	99.91	73.56
RF	99.92	75.70

After completing this analysis examining the accuracy score and f1 score, they discovered that KNN is effective in identifying fraudulent transactions and reducing the amount of false alarms.

Asha RB, Suresh Kumar KR, et. al. [12] the use of an artificial neural network for credit card fraud detection. Through e-commerce and a number of other websites, online payment alternatives have increased, it increases the chance of online fraud. Although there are other strategies, such as algorithmic approaches for machine learning and data mining are used to detect fraud in credit card transactions, they have not shown significant outcomes. As a result, important algorithms that are both effective and efficient must be developed. Through utilizing an artificial neural network method and comparing it with a few other machine learning algorithms, they attempt to prevent fraudsters from using our credit card before the transaction is authorized. Several machine learning techniques, such as Support vector machines (SVM), K-nearest neighbours (KNN), and artificial neural networks (ANN), are used in this study to attempt to predict the likelihood of fraud.

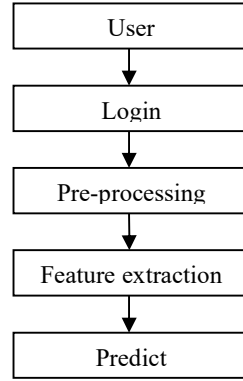


Fig. 3: Complete Architecture of Proposed System

The client transactions done at a European bank in the years 2013–2014 make up the data collection. The dataset utilized in the experiment has 31 attributes out of which 30 elements provide data on name, age, account information, and other related topics. The final attribute provides the transaction's outcome, which can be either 0 or 1. Fig. 3 demonstrates the suggested system's design. According to the design, users must provide their user name, password, email address, and phone number in order to log in to the system after registration. The user is then prompted to choose an algorithm in the selection windows that appear. Following the technique selection is data pre-processing, which indicates cleaning and normalizing the dataset before it is input into the actual model and data splitting. The model is now being evaluated and utilizing machine learning techniques, it's determine if a transaction is fraudulent or not after training.

TABLE V: COMPARATIVE PERFORMANCE

Algorithms	Accuracy	Precision	Recall
SVM	93.4	97.43	89.76
KNN	99.82	71.42	39.3
ANN	99.92	81.15	76.19

In Table 5, the performance metrics for the utilized algorithms, such as recall, accuracy, and precision, are shown. A calculation of precision, recall, and accuracy is made when the end result is evaluated using the confusion matrix. It has two classes: projected class and actual class. The following characteristics affect the confusion measurements:

False Positive (FP): This occurs when the true class is 0 while the false class is 1.

True Positive (TP): wherein the two positive values 1, are.

True Negative (TN): 0 is the result when both values are negative.

False Negative (FN): This situation arises when the true class is 0, while the actual class is 1.

$$Accuracy = \frac{(TP + TN)}{(TP + TN + FP + FN)} \dots (1)$$

$$Precision = \frac{TP}{(TP + FP)} \dots (2)$$

$$Recall = \frac{TP}{(TP + FN)} \dots (3)$$

The accuracy performance demonstrates that artificial neural networks (99.92%) forecast credit card theft with a greater degree of accuracy than support vector machines (93.4%) and k-nearest neighbour (99.82%) methods.

el Naby Aya Abd, Ezz El-Din Hemdan, El-Sayed Ayman, et. al. [13] explains the use of deep learning to identify credit card frauds. The loading of the credit card information into the suggested model is the first step in the identification of credit card frauds. On Kaggle's credit card dataset, they show a model for identifying real transactions from fraudulent ones. The files comprise two-day credit card transactions from September 2013 that

were made by European cardholders. Utilize the SMOTE technique for imbalancing credit card data after preprocessing oversampling data. 20% was used for validation, while 80% was used for training on the dataset. Apply deep learning algorithm (CNN).

The proposed method is called OSCNN (Over Sampling with Convolution Neural Network), which is based on CNN (Convolution Neural Network) with oversampling preprocessing. The minority class is initially oversampled by 0.25, resulting in a 75:25 minority class to majority class ratio in the OSCNN model. This ratio was set in order to avoid the overfitting that the SMOTE (Synthetic Minority Over-sampling Technique) causes. The completely connected layer of described model, which serves as its last layer, ends with a dense function that has a single output node for binary detection work. They used the method, changing the number of epochs, until the accuracy was optimal. The presented approach for differentiating between valid fraudulent credit card purchases is displayed in Fig. 4.

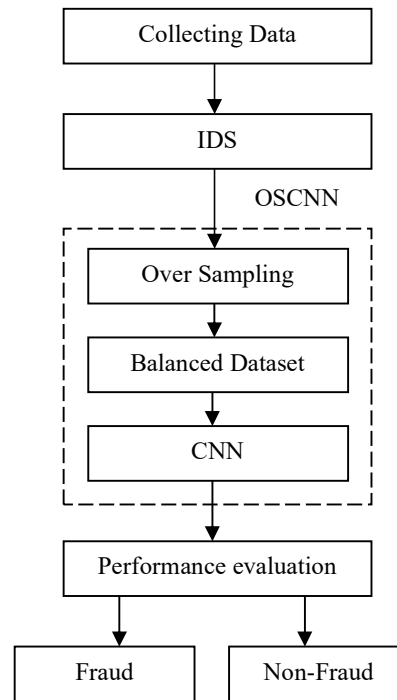


Fig 4: Proposed System (Oscnn)

The original dataset was modelled using Multi-layer perceptron (MLP) and CNN. SMOTE was then applied to the dataset, and MLP and CNN were used to evaluate the results. The results are then compared, and the best accurate model is selected. Precision, recall, and accuracy are computed, and the outcome is assessed using the confusion matrix. Since 50 epochs provide the best accuracy, the epoch's number has been changed. They were able to demonstrate through experiments that the proposed approach OSCNN was more accurate than the MLP and MLP+SMOTE models.

TABLE VI: COMPARATIVE PERFORMANCE

Parameters		MLP	MLP+SMOTE	OSCNN
Accuracy	30 Epochs	88	96	97
	50 Epochs	88	97	98.7
Precision	30 Epochs	40	97	97
	50 Epochs	40	97	97
Recall	30 Epochs	40	89	91
	50 Epochs	40	89	91

When compared to the MLP and the MLP with SMOTE, accuracy from 0.88 to 0.989, they advise using the supplied OSCNN technique and applying it to the dataset to ensure model performance.

Taha Altyeb Altaher, Jameel Sharaf Malebary et. al. [14] the analysis is an Intelligent Approach to Credit Card Fraud Detection Using an Optimized Light Gradient Boosting Machine (OLightGBM). A Light gradient boosting machine's (LightGBM) parameters are tuned using the suggested technique, which effectively incorporates a Bayesian-based hyperparameter optimization algorithm. The main contribution of described research is a intelligent approach for identifying fraud in credit card transactions using an optimized light gradient boosting machine. This method uses a Bayesian-based hyperparameter optimization algorithm to optimize the parameters of the light gradient boosting machine. Using two real-world public credit card transaction data sets that included both fraudulent and valid transactions, experiments were carried out to show the efficacy of proposed OLightGBM for identifying fraud in credit card transactions.

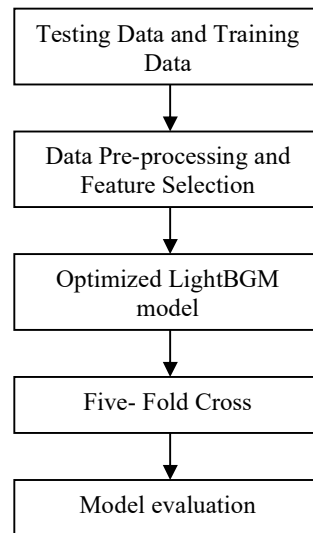


Fig. 5: Overall Framework of the Proposed Approach

284,807 credit card transactions made by cardholders in Europe in September 2013 are included in the initial batch of data. The second data set consists of an actual data collection of e-commerce transactions called the UCSD-FICO (University of California, San Diego- Financial and Controlling) Data Mining Contest 2009 Dataset. In this study, the model is trained and tested using a cross validation technique to get more accurate results. In Fig. 5, the general design of the intelligent technique for credit card fraud detection that is described.

Although there are a lot of elements, it is essential to choose the most crucial and those that are crucial for the effective detection of credit card fraud. The most crucial characteristics are chosen by LightGBM using the Information gain (IG) approach, it also makes the training data decrease the dimensionality. The parameters of the LightGBM algorithm are intelligently tuned using the suggested approach's Bayesian based hyper parameter optimization technique. Massive data volumes and dispersed data processing may be efficiently managed by the quick LightGBM approach. To maintain the precision of the information gain estimation, LightGBM utilizes the Gradient-based one side sampling (GOSS) approach. In order to evaluate the effectiveness of the proposed intelligent technique, there are two real-world data sets used, and performance evaluation measures are used to compare to other machine learning methods.

To provide more accurate estimates, the model is trained and tested in each subset of the two data sets using cross validation. After this, the data set's average for all the specified metrics is determined. The performance of the proposed approach is evaluated in this study using a 5-fold CV test. Each data set is partitioned into five equal-sized subgroups at random. The training data set is made up of the last four subsets, accounting for 80% of the total data set, and one subset (20% of the total data set) is put aside for evaluating the usefulness of the recommended approach at each stage of validation. Once each subgroup has been utilized, this procedure is performed five more times. Through, the averaging the outcomes of the five test subsections, effective estimation of the proposed technique's overall performance on a 5-fold CV test is possible.

In this approach, the following parameters are used: Confusion Matrix, Recall, Precision, Accuracy (ACC), AUC (Area Under Curve), and F1-score. Table 3 displays an evaluation of performance of the suggested strategy

utilizing the two real-world data sets and the 5-fold CV methodology. The terms used by the Confusion Matrix to evaluate the effectiveness of credit card fraud detection are as follows:

The number of fraudulent credit card transactions that were accurately recognized is referred to as "TP" (True Positive).

It demonstrates that many genuine instances of credit card fraud have been labeled as False positives, or FPs.

The term FN, which stands for false negative, indicates that numerous fraudulent credit card transactions have been recognized as regular.

True Negative, often known as TN (True negative), is the proportion of common credit card transactions that were really recognized. The following definitions apply to the measurements that were utilized.

$$Accuracy = \frac{(TP + TN)}{(TP + TN + FP + FN)} \dots (4)$$

$$Precision = \frac{TP}{(TP + FP)} \dots (5)$$

$$Recall = \frac{TP}{(TP + FN)} \dots (6)$$

$$F1 - Score = 2 * \frac{Precision \times Recall}{(Precision + Recall)} \dots (7)$$

TABLE VII: PERFORMANCE EVALUATION OF THE PROPOSED APPROACH

Parameters	Dataset 1 (for average 5 folds)	Dataset 2 (for average 5 folds)
AUC	90.2	92.9
Accuracy	98.4	98.3
Recall	40.54	28.33
Precision	97.3	91.7
F1-Score	56.9	43.27

TABLE VIII: ACCURACY COMPARISON

Approach	Accuracy (%)
Isolation Forest	95
Random Forest	95.5
ANN	92
Proposed Approach	98.4

Based on the accuracy attained for the same data set, Table 4 presents a performance comparison of the proposed approach with other results of the study. Compared to other ways, the suggested strategy had the greatest Accuracy (98.40%). The results of the experiments show that the suggested technique performed better than other machine learning algorithms in terms of Precision, Accuracy, F1-score and AUC.

Ruttala Sailusha, V. Gnaneswar, R. Ramesh, G. Ramakoteswara Rao, et. al. [15] It is discussed that to use machine learning to detect credit card fraud. The development of the credit card fraud detection system made it possible to identify fraudulent behaviour. This paper's primary goal is to use algorithms to categorize transactions that are both fraud and non-fraud transactions in the dataset. The major goal of this research is to concentrate on machine learning techniques. Both the Adaboost method and the random forest algorithm are used. In order to determine that algorithm is the identifies credit card fraud transactions, these two algorithms are compared. The partitioning of the data, model training, model deployment, and evaluation criteria are all included in the process flow for the credit fraud detection problem. A European credit card supplier provided information on credit card theft. The dataset was downloaded from Kaggle. The dataset contains the transactions that cardholders made in September 2013 using their cards. The information includes the transactions completed during the previous two days. The data collection includes 284,807 transactions of which 492 are fraudulent. Out of all the transactions, fraud transactions made up of 0.172%. The Principal component analysis (PCA) transformation turns the dataset containing the input variable into numerical values. Due to concerns about

privacy, this is completed. To create the model, they must now separate the data into training and testing sets. To create the Random Forest and Adaboost models, they use the training data. The two models are then developed.

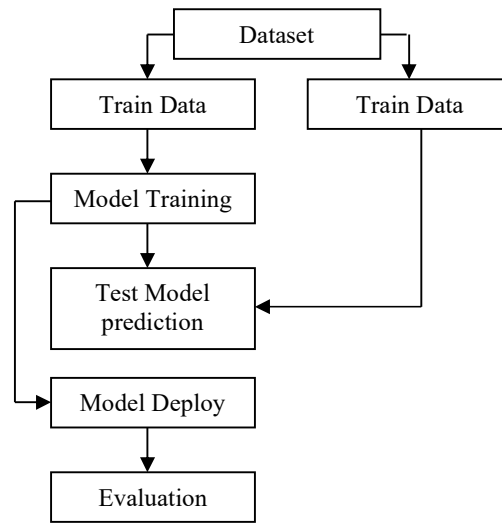


Fig. 6: Process Flow

One of the common supervised learning methods is called Random Forest. This may be applied for classification as well as regression. But categorization issues are this approach is most frequently utilized. Although a forest is frequently made up of trees, the Random Forest technique develops decision trees using the sample data and derives predictions from each sample in a similar way. Therefore, the Random Forest algorithm is an ensemble method. Since it averages the results, this method performs better than single decision trees at preventing over-fitting.

One of the ensemble methods is boosting. By using weaker classifiers as a starting point, this method creates stronger classifiers. This might be achieved by developing a strong model out of a weak model in the series. A model is initially created using the training data. The first model is then used to produce the second model, with the mistakes in the first model being fixed. This approach is repeated iteratively until the maximum number of models are added or all of the training dataset's predictions are accurate. One of the most effective boosting algorithms created for the binary classification was called Adaboost.

The outputs of the two algorithms are based on the F1-score, as well as accuracy, precision, recall, and other factors. The confusion matrix is used to plot the ROC (*Receiver Operating Characteristic*) curve. The most efficient method for detecting fraud is the when the algorithms from Random Forest and Adaboost are compared has the best accuracy, precision, recall, and F1-score. Finally, a comparison of credit card fraud transactions that is more accurate. TNR, TPR, FPR, and FNR are the four outputs of the matrix.

1. The quantity of fraudulent transactions that the system even identifies as as the "True Positive Rate (TPR)". 2. The fraction of legitimate transactions that the system that acknowledges as actual, often known as the True Negative Rate (TNR). 3. The percentage of real transactions that are wrongly classified as fraudulent, or False Positive Rate (FPR). 4. False Negative Rates (FNR) refer to fraudulent transactions that are incorrectly classified as authorized.

According to the confusion matrix, the true positives and false positives for the train data are 190490 and 0, respectively. The true negatives and false negatives are both 0. In the test results, 93818 real positives, 37 fake positives, 7 real negatives, and 125 fake negatives were detected. The train data included 190464 true positives and 120 false positives, whereas there 26 true negatives and 201 false negatives, according to the confusion matrix. In the test results, 93811 real positives, 65 false positives, 14 real negatives, and 97 false negatives were found. The two algorithms have the same accuracy, however they differ in terms of precision, recall, and F1 score. The greatest accuracy, recall, and F1-score are achieved by algorithms that use random forest. They get to the conclusion that the Random Forest method detects credit card fraud better than the Adaboost method.

Mohammed Azhan, Shazli Meraj, et. al. [16] demonstrates that to use machine learning and deep learning techniques to detect credit card fraud. A more thorough examination of the various credit card fraud schemes is included in the proposed study. The present study addresses how Machine Learning and Neural Networks can be

used to identify potential fraudsters by referring to their previous mistakes and details of previous fraudsters, even though all of them cannot be dealt with at the same time. Naive Bayes classifiers, logistic regression, K nearest neighbours, multiple linear regression, random forest classifiers, and a basic neural network classifier are a few examples of machine learning techniques.

The Machine Learning Group at ULB (Universite Libre de Bruxelles) assembled and made the dataset available. The dataset has no missing values and solely numerical inputs, but it is significantly unbalanced, with a smaller number of data series belonging to the positive class than the negative class. The information on whether or not the transaction was fraudulent is contained in the column Class, where the value 1 denotes fraud and 0 denotes. The mentioned Github repository contains the source codes for each model as well as an accumulation of the results.

Using classification reports from the Sklearn library and a confusion matrix, the models were assessed. The f1-score for the negative class in the multiple linear regression model was 1 indicating that there were no false positives or false negatives in the predicted values, whereas the f1-score for the positive class was 0.77. The positive class had a precision value of 0.75, making it unable to identify 25% of the positive frauds. The model, which had a negative class accuracy score of 1, accurately predicted all negative classes. The logistic regression's positive class f1-score in this case was 0.70, the negative class f1-score is similarly 1. For the negative and positive classes, the model's accuracy was found to be 1 and 0.66, respectively. The negative class's f1-score in this example for K Nearest Neighbours is similarly 1, but the positive class's f1-score improved from the previous two classifiers to 0.78. The classifier for multiple linear regression with a threshold of 0.3 maintained the same accuracy.

According to the results, the Naive Bayes Classifier was a total failure. It has a very poor f1-score for predicting the positive class. It was the poorest approach for separating positive and negative classes because of its precision for the positive class, which was 0.02. The Random Forest Classifier, an ensemble approach, is the next step. To prevent over-fitting on the training data, the maximum depth and number of decision trees were both limited to 3. The accuracy and f1-score of the classifier were stated the negative and positive classes, 1.0 and 0.69. Differential weights were used to a neural network classifier in order to correct the class imbalance. The weight ratio of 1:1.04 produced the desired results. It is clear that balancing the network's weights can have certain advantages. It is important to note that the classifier's accuracy was 1.0 and 0.75 for the negative and positive classes, respectively, and that the f1-score was 1.0 and 0.78 additionally. The train and test splits were kept same for each of these weight distribution models to guarantee the accuracy of the results. Each model was also run for an equal number of epochs while maintaining all other hyper-parameters at their original values. The accuracy and f1-score of the various methods used in the paper are compared in Table 9.

TABLE IX: F1-SCORE AND PRECISION OF VARIOUS ALGORITHMS USED

Model	F1-Score		Precision	
	0	1	0	1
Multiple Linear Regression	1	0.75	1	0.77
Logistic Regression	1	0.66	1	0.70
K Nearest Neighbors	1	0.75	1	0.78
Naive Bayes	1	0.02	1	0.05
Random Forest	1	0.69	1	0.71
Neural Network	1	0.75	1	0.78

They may determine from the results that K Nearest Neighbours performs better than the other classifiers. It is important to point out that additional methods may have been used to enhance the neural network's performance. Additionally, weight balancing has shown to be useful.

Sangeeta Mittal, Shivani Tyagi, et. al. [17] examines the effectiveness of machine learning algorithms for detecting credit card fraud. It is impossible for credit card providers and businesses to find these fraudulent transactions during the thousands of legitimate ones. Machine learning techniques can be used to resolve this issue if enough data is gathered and made available. Popular supervised and unsupervised machine learning methods have been used in this study to identify credit card thefts in a significantly imbalanced dataset. The range of supervised learning algorithms, has been taken into consideration, ranging from the earliest to the most recent. These consist of hybrid algorithms, neural networks for both deep learning and classical learning, Bayesian techniques, and tree-based algorithms. As supervised learning algorithms, Random Forest, Neural Networks, Deep Learning, Support Vector Machine, Naive Bayes, Logistic Regression, Extended Gradient

Boosted Tree, Quadratic Discriminant Analysis, K-Nearest Neighbour, and Hybrid Supervised Approaches are selected.

Unsupervised approaches group identical data rows together and categorize them into the same class. The ability to identify outliers rows that don't fit into any of the clusters makes these extremely valuable. Because there are more fraud cases than usual, which might be regarded as outliers, unsupervised approaches are therefore especially helpful in this situation. Self organizing maps (SOM), K-means, isolation forests (IF), and Local outlier factors (LOF) are a few unsupervised detection techniques that were employed in this work. Andrea Dal Pozzolo and peers conducted research on large data mining and fraud detection with the dataset, which was gathered and examined as part of a Worldline and ULB (University Libre de Bruxelles) Machine Learning Group collaboration. A total of 284,807 transactions done by cards across Europe in September 2013 are included in the report. The data set is quite unbalanced since it contains 492 fraud transactions. A total of 28 features that were during principal components analysis of real features are supplied due to the confidentiality issue. The transaction amount is the attribute "Amount." The transaction type label is the property "Class", has a value of 1 in cases of fraud and 0 in all other cases.

Positive Predictive Value (PPV) or Precision or Sensitivity, Negative Predictive Value (NPV), Balanced Accuracy, Specificity, Prevalence and Diagnostic Odd Ratio (DOR) are used parameters in this study. Two measures are available to identify classification efficiency in unbalanced datasets: balanced accuracy and diagnostic odds ratio.

$$\text{Balanced Accuracy} = \frac{(\text{Sensitivity} + \text{Specificity})}{2} \dots (8)$$

$$\text{Specificity} = \frac{TN}{(TN + FP)} \dots (9)$$

$$\text{Sensitivity/ PPV} = \frac{TP}{(TP + FP)} \dots (10)$$

$$\text{NPV} = \frac{TN}{(TN + FN)} \dots (11)$$

$$\text{Prevalence} = \frac{(TP + FN)}{(TN + FN + TP + FP)} \dots (12)$$

$$\text{DOR} = \frac{(PPV * NPV)}{(1 - PPV)(1 - NPV)} \dots (13)$$

The selected methods (ANN, NN, ensemble model, XGBT LR, LOF, and IF) all generated results that were almost desirable. Higher DOR levels are regarded as better outcomes; this statistic is also balanced. While other hybrid models produced comparable results, XGBT provided the best DOR. Unsupervised Learners, IF LOF ranked the highest in both the among unsupervised groups. As a result, unsupervised approaches, particularly IF and LOF, remain successful when managing extremely unbalanced datasets. Compared to previous strategies, unsupervised algorithms perform well across all metrics because they manage dataset skewness more effectively.

TABLE X: PERFORMANCE EVALUATION OF SUPERVISED AND UNSUPERVISED MACHINE LEARNING ALGORITHMS

Techniques	PPV	NPV	Prevalence	Specificity	DOR
NB	0.06	0.99	0.001	0.97	6.319
RF	0.99	0.16	0.998	0.91	18.857
KNN	NaN	0.99	0.001	1.0	0
LR	0.99	0.63	0.99	0.87	168.56
XGBT	0.99	0.92	0.99	0.81	1138.5
SVM	NaN	NaN	NaN	0.92	0
ANN	0.99	0.84	0.99	0.77	462
DL	0.98	0.93	0.86	0.98	651
QDL	0.97	0.89	0.97	0.47	261.6
NN	0.99	NaN	0.99	0.0	0
SOM Hybrid	0.92	0.84	0.99	0.83	60.375
IF	0.99	0.99	0.99	1.0	9801
LOF	0.99	0.99	0.998	1.0	9801
K-Means	0.99	NaN	0.998	1.0	0

III. CONCLUSION

A Literature survey on credit card fraud detection techniques are presented in this analysis. Through the extreme imbalance between legitimate and fraudulent transactions, a unusual categorization problem is identifying credit card fraud. Some researchers are beginning to question the design and application of new technology, despite the fact that it is readily available and extensively supported by banks and merchants worldwide to decrease or possibly eliminate the effects of credit card fraud. The primary data used in the fraud detection process are credit card transactions. This paper discusses and explains various methods for detecting credit card fraud. Performance metrics used frequently in credit card fraud detection include precision, accuracy, F1-score, recall. Furthermore, organizations are interested in finding techniques that reduces the costs and improve profit; they can locate and select the approach from the above studies.

REFERENCES

- [1] Steven Lufthanza Bong, Albertus Baskara Yunandito Adriawan, Ghina Kamilah, Evaristus Didik Madyatmadja, "Analysis of Consumer Behavior in online shopping on Social Media", 2023 International Conference On Cyber Management And Engineering (CyMaEn), Year: 2023
- [2] Yathartha Singh, Kiran Singh, Vivek Singh Chauhan, "Fraud Detection Techniques for credit card Transactions", 2022 3rd International Conference on Intelligent Engineering and Management (ICIEM), Year: 2022
- [3] Sahil Negi, Sudipta Kumar Das, Rigzen Bodh, "Credit Card Fraud Detection using Deep and Machine Learning", 2022 International Conference on Applied Artificial Intelligence and Computing (ICAAIC), Year: 2022
- [4] N. Prabha, S. Manimekalai, "Imbalanced data Classification in Credit Card fraudulent activities Detection using Multi-Class Neural Network", 2022 Second International Conference on Artificial Intelligence and Smart Energy (ICAIS), Year: 2022
- [5] Madhushika Delgolla, Thilina Halloluwa, Anuradha Rathnayake, "A rule based approach to minimize false-positive declines in Electronic card not present financial transactions using feature engineering techniques", 2021 21st International Conference on Advances in ICT for Emerging Regions (ICTer), Year: 2021
- [6] Jyoti Singh Kirar, Dhiraj Kumar, Diptirtha Chatterjee, Prasoon Singh Patel, Shailendra Nath Yadav, "Exploratory Data analysis for Credit Card Fraud Detection", 2021 International Conference on Computational performance Evaluation (ComPE), Year: 2021
- [7] Jain V, Agrawal M, Kumar A. "Performance analysis of machine learning algorithms in credit cards fraud detection", In 2020 8th International Conference on Reliability, Infocom Technologies and Optimization (Trends and Future Directions)(ICRITO) 2020 Jun 4 (pp. 86-88). IEEE.
- [8] Anish Mahajan, Vivek Singh Baghel, Ramkumar Jayaraman, "Credit Card Fraud Detection using Logistic Regression with Imbalanced Dataset", IEEE Xplore, 2023
- [9] Alavikunhu Panthakkan, Najiya Valappil, Majida Appathil, Seema Verma, Wathiq Mansoor, and Hussain Al-Ahmad, "Performance comparison of Credit Card Fraud Detection system Machine Learning", 2022 5th International Conference on Signal Processing and Information Security (ICSPI), 2022
- [10] C.H Sumanth, Pokala Pavan Kalyan, Bolisetti Ravi, S Balasubramani., "Analysis of credit card fraud Detection using Machine Learning Techniques", 2022 7th International Conference on Communication and Electronics Systems (ICCES), Year: 2022
- [11] Prabhat Singh, Vishesh Chauhan, Shivam Singh, Priya Agarwal, Shrey Agrawal, "Model for Credit card fraud Detection using Machine Learning Algorithm", 2021 International Conference on Technological Advancements and Innovations (ICTAI), Year: 2021
- [12] Asha RB, Suresh Kumar KR, "Credit card fraud detection using artificial neural network", Global Transitions Proceedings 2 (2021)
- [13] Aya Abd el Naby, Ezz El-Din Hemdan, Ayman El-Sayed, "Deep Learning Approach for Credit Card Fraud Detection", 2 nd IEEE International Conference on Electronic Engineering, Menoufia University, Egypt ICEEM2021
- [14] Altyeb Altaher Taha, Sharaf Jameel Malebary, "An Intelligent Approach to Credit Card Fraud Detection Using an Optimized Light Gradient Boosting Machine", IEEE Access, VOLUME 8, 2020
- [15] Ruttala Sailusha, V. Gnaneswar, R. Ramesh, G. Ramakoteswara Rao, "Credit Card Fraud Detection Using Machine Learning", Proceedings of the International Conference on Intelligent Computing and Control Systems (ICICCS 2020)
- [16] Mohammed Azhan, Shazli Meraj, "Credit Card Fraud Detection using Machine Learning and Deep Learning Techniques", Proceedings of the Third International Conference on Intelligent Sustainable Systems [ICISS 2020], 2020
- [17] Sangeeta Mittal, Shivani Tyagi, "Performance Evaluation of Machine Learning Algorithms for Credit Card Fraud Detection", 9th International Conference on Cloud Computing, Data Science & Engineering (Confluence), 2019.