

Research Paper

Design and Implementation of an Ensemble Learning–Based Credit Card Fraud Detection System

Sumera Naaz

PG Scholar, Department of
Computer Science and
Engineering, Shadan College of
Engineering and Technology,
Hyderabad, Telangana, India -
500086

sumeranaaz086@gmail.com

Md. Ateeq Ur Rahman

Professor, Department of
Computer Science and
Engineering, Shadan College of
Engineering and Technology,
Hyderabad, Telangana, India -
500086

ateeq@yahoo.com

Subramanian K.M

Professor, Department of
Computer Science and
Engineering, Shadan College of
Engineering and Technology,
Hyderabad, Telangana, India -
500086

kmsubbu.phd@gmail.com

Abstract: Fraudulent transactions using credit cards is a huge issue for financial institutions as hackers can pose as authorised cardholders. Several resampling techniques such as oversampling, undersampling and SMOTE are used to solve the class imbalance problem on European Data and Sparkov Data. Various approaches are adopted to improve the classification performance through ensemble learning to increase the accuracy and robustness of the classification. In this paper, a framework for ensemble is proposed using advanced resampling techniques during model training and prediction to optimise the process. Through a comprehensive evaluation of different classification models and by using Stacking Classifier with the smart combination of many base models, it is observed that it achieves a better accuracy, precision, recall and F1 score for all the methodologies. This method could be a powerful tool that helps enhance fraud detection systems by detecting fraudulent transactions accurately without raising too many false alarms. The proposed solution emphasizes the importance of incorporating ensemble methods and data balancing techniques for solving financial fraud detection problems.

Index Terms - Fintech, credit card fraud detection, ensemble learning, machine learning, simulated dataset, real-world data set.

I. INTRODUCTION

CCF detection is one of the most challenging areas in digital payments, where irregularities in credit card transactions are detected. On their own it is also quite easy to see the magnitude of this problem, as there were approximately 2.183 million cardholders of Visa and MasterCard in 2020 alone [1]. In the same year, there are over 1.4 million occurrences of identity theft, 393,207 of which are due to credit card fraud [2]. The financial losses due to fraudulent use of credit cards were huge, with \$28.6 billion in 2020 as compared to \$23.97 billion in 2017, a rise of 19.3%. These losses are estimated to be \$408 billion at the end of the decade, according to the 2022 Nilson Report [3]. With the constant accumulation of this type of data, even with the help from banks and finance

institutes, there is a clear need for efficient credit card fraud detection systems.

There are various ways established to combat credit card usage. These techniques can be roughly divided into statistical and ML based techniques [2]. Finding out fraudulent transactions is to find outliers in a dataset from statistical point of view. To this end, statistical methods have been employed, such as box-and-whisker plots, normal distribution analysis and cluster analysis. However, the complexity of credit card transaction data is increasing the use of machine learning techniques, which are more adaptable and accurate.

Fraud transaction analysis and prediction is done by ML classifiers based on previous data. DT, SVM, NNW and regression models [3, 4, 5] are popular approaches for CCF identification. These are successful methods to a greater or less extent,

but more often than not, they rely on the data quality and nature for training.

To improve detection ability, ensemble learning methods are proposed. The idea of ensemble learning is to create a more robust and stable meta-classifier that is the combination of many basic classifiers. One way is bagging, where classifiers are trained on partially different subsets of training data (with replacement) and the final predictions are made by combining the results of these classifiers, such as by voting. Another ensemble method is boosting, which incrementally builds classifiers, training them to improve on the previous classifiers' prediction errors. Some examples are AdaBoost, Gradient Boosting and XGBoost [6].

Stacking or stacked generalisation is another ensemble method that uses a number of classifiers to create a high performing meta-classifier. A potential solution to the problems of credit card fraud detection is the idea of stacking to exploit several models [7].

II. RELATED WORK

Detecting credit card fraud is an ongoing challenge for financial organisations as the amount of transactions and sophistication of criminals improve. Various research works have been carried out using statistical and machine learning techniques to decrease the chance of fraudulent transactions. The strategies aim to boost the precision and dependability of fraud identification processes, minimizing false positives and false negatives.

M. Varmedja et al. [10] studied the Performance of machine learning algorithms in credit card Fraud detection system. They concentrated on the significance of imbalanced data set treatment techniques that is a common characteristic in fraud detection problem including oversampling and undersampling and SMOTE. The results showed that ensemble technique like RF and Gradient Boosting showed superior performance compared to individual models because of the better generalisation on unseen data. Raj and Portia [11] also examined various approaches towards fraud detection and concluded that machine learning classifiers are the most effective tools for detecting anomalies in transaction data. They stated that machine learning technologies were suitable for adaptive fraud trends because of how flexible they are.

The credit card fraud detection problem has been studied as a potential use of deep learning. Vimal and Vimal [12] have proposed a population based

optimised and condensed fuzzy deep belief network for better detection accuracy. The strategy was robust to extremely imbalanced datasets, by incorporating feature selection and dimensionality reduction techniques. Razooqi et al. [13] presented a hybrid system, relying on fuzzy logic and neural networks, which provided a higher accuracy and interpretability for fraud detection compared to the traditional machine learning methods. The studies show the increased interest in the application of deep learning and hybrid methods to solve the issues of detecting fraudulent transactions.

The use of visualisation techniques has been proposed as a mean to enhance fraud detection. Lokanan [14] studied the utilisation of visualisation tools in the discovery of fraudulent trends in credit card transactions and money laundering. This technique helps investigators to understand complicated data correlations and anomalies. It offers an easy-to-use approach to the analysis of very large data sets. Visualisation is not a standalone approach, but can be a supporting technique to other detection strategies.

Another new type of approach that has emerged from the credit card fraud detection is graph-based approaches. Lebichot et. al. [15] put forward a semi-supervised system for fraud detection using graph-based models. It was able to identify irregular patterns in the transaction data that would indicate fraudulent activities by modelling transactions as nodes and their relationships as edges. This solution was semi-supervised which resolved the problem of a lack of labeled data often encountered in real fraud detection.

The adaptability and scalability of ML approaches make them the preferred way to tackle the landscape of credit card fraud detection, as seen in this still. Awoyemi et al. [16] compared the accuracy of different ML classifiers such as DT, SVM, and k-NN. Their results proved that ensemble techniques such as bagging and boosting algorithms, performed better. To enhance the level of detection, many classifiers are used to produce ensemble learning, which reduces overfitting and increases generalisation.

Hybrid models have been promising to address these challenges. Fraud detection is an all-encompassing technique that utilizes statistical techniques and ML algorithms. In addition to the above, traditional outlier detection methods like clustering and density-based methods can be used as a pre-processing step to detect suspicious transactions, adding the data to ML classifiers for additional analysis. For instance, This layered strategy combines the best of both worlds so as to cover all fraudulent patterns.

Fraud detection has been a field that has been heavily adopted by machine learning and statistical techniques, but the need for XAI in this area is growing increasingly important. Development of interpretable models that can clearly explain their prediction is demanded by regulatory regulations and need for transparent decision making. This is especially true in applications in finance where false positives could lead to dissatisfied customers and false negatives could lead to large financial losses.

III. MATERIALS AND METHODS

The suggested system attempts to enhance credit card fraud detection by integrating enhanced data sampling, feature selection, and machine learning techniques. This data set has two datasets: a European Data and a Sparkov Data. To address this class imbalance problem, various data sampling strategies like OverSampling, UnderSampling and SMOTE will be employed on the dataset to experiment with. The dimensionality will be reduced using PCA and feature selection will be performed to enhance model performance. Different machine learning models will be developed and tested: RF [19], NB [18] and XGB [20] to learn various patterns in the data. In addition, the use of ensemble techniques to enhance accuracy and robustness (such as Soft Voting Classifier with RF, NB, XGBoost or Stacking Classifier with Bagging Classifier and RF, DT, LightGBM) will be explored. These techniques will be used together to construct a very efficient and effective fraud detection system.

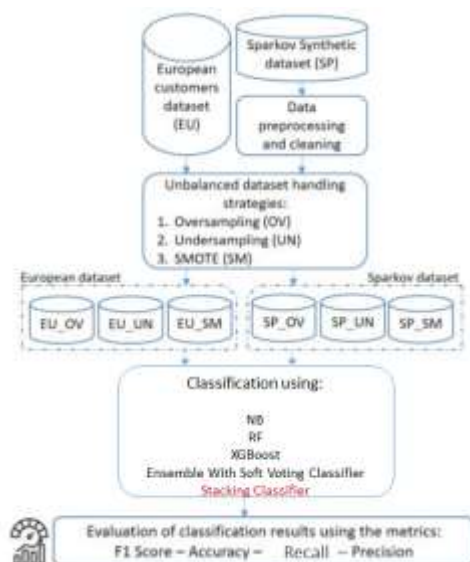


Fig.1 Proposed Architecture

The architecture is based on two datasets (European and Sparkov) to detect fraudulent

transactions. Data pre-processing and cleaning is carried out first. To overcome the problem of unbalanced dataset, the approaches used are oversampling, undersampling and SMOTE. The classification is done using NB, RF, XGB, Ensemble with Soft Voting Classifier and Stacking Classifier. The evaluation of the model is done on the basis of the following performance criteria: F1 Score, Accuracy, Recall and Precision.

i) Dataset Collection:

The following variables can be anonymised and used for identifying fraudulent transactions by ML. European and Sparkov datasets contain transaction data, including variables such as time, amount, etc. The databases are provided with marked fraud signs for study.

For credit card fraud detection dataset utilised is the European Data [8]. There are 284,807 transactions and 31 characteristics. This contains the details of the transaction, such as the time of the transaction, the transaction amounts and the 28 anonymised numerical features (V1-V28) plus the target variable "Class" indicating whether the transaction was fraudulent (1) or lawful (0). There are no missing values in the data. All the elements are continuous except for the "Class" column. The data provides a wealth of information that can help identify trends and anomalies in financial activities.

	Time	V1	V2	V3	V4	V5	V6	V7	V8
0	0.0	-1.359807	-0.072781	2.536347	1.378155	-0.338321	0.462388	0.236599	0.066698
1	0.0	1.191857	0.266151	0.166480	0.448154	0.060018	-0.082361	-0.076803	0.065102
2	1.0	-1.359354	-1.340163	1.773209	0.379780	-0.503198	1.800499	0.791461	0.247676
3	1.0	-0.968272	-0.185226	1.792963	-0.863291	-0.010309	1.247203	0.237609	0.377436
4	2.0	-1.158233	0.877737	1.548718	0.403034	-0.407193	0.095621	0.582941	-0.270533

Fig.2 Dataset Collection Table – European

The dataset shipped with Sparkov contains 1,296,675 transactions with seven columns, including transaction date and time ("trans_date_trans_time"), amount of the transaction ("amt"), and population of the city where the transaction took place ("city_pop") and the ZIP code of the city ("zip"). The data set also includes information about customers, including their date of birth ("dob") and the target variable "is_fraud", which is defined as 1 for fraudulent transactions and 0 for legitimate transactions. All the columns are non-null, with numerical and categorical data types, making it a good dataset for fraud detection research.

Unnamed: 0	trans_date_trans_time	cc_num	merchant	category	amt
0	2019-01-01 00:00:18	2703185189852095	fraud_Rippt, Kub and Mann	misc_net	4.97
1	2019-01-01 00:00:44	630423337322	fraud_Hofer, Gutmann and Ziemer	grocery_pos	107.23
2	2019-01-01 00:00:51	3889462057661	fraud_Lind-Bockelgo	entertainment	220.11
3	2019-01-01 00:01:18	3534093764340240	fraud_Kuich, Hermsdorf and Farnel	gas_transport	45.00
4	2019-01-01 00:03:08	379534208663964	fraud_Koelling-Ciel	misc_pos	41.96

Fig.3 Dataset Collection Table - Sparkov

ii) Pre-Processing:

In this pre-processing phase we will cover the necessary procedures of data processing, data visualisation, feature selection via PCA decomposition and data sampling methods.

a) Data Processing: Data needs to be processed to be used as input to ML. It also involves data cleaning, such as handling missing values, conversion of data types and removal of data duplicates, which is required in credit card fraud detection for algorithms that are scale-dependent like neural networks. Features that are categorical may need to be encoded and features that are irrelevant should be eliminated. Models can learn trends and identify fraudulent transactions if analysed correctly.

b) Data Visualization: Visualising data can aid in understanding the distribution and relationship between features and is very important for data pre-processing. Histograms, scatter plots, heatmaps and more can be used to find patterns, correlations and anomalies in the data collection. Being able to visualise the distribution of legal versus fraud transactions and the relationship between different factors like transaction values, transaction time etc. helps to identify outlier and trends in fraud detection. Visualisation also helps to evaluate the class imbalance and to guide future preprocessing actions like as resampling or feature engineering.

c) Feature Selection: PCA is a method of dimensionality reduction of a data set. It can be used for determining the most important aspects in the data. It changes the original features to a family of linearly uncorrelated components- sorted by their variance. The PCA is employed in fraud detection to minimize the number of features by eliminating features that contribute the most to the variation and redundant or unnecessary features, thereby making the model more efficient and effective.

d) Data Sampling: Data sampling techniques are utilised to solve the problem of class imbalance existing in fraud detection datasets. To balance the data set, the number of instances of the minority class is increased through oversampling. Undersampling decreases the instances of the majority class to balance the data set. SMOTE is a technique that creates synthetic minority-class samples by interpolating between minority-class samples. These techniques can be useful in improving the model's ability to identify fraudulent transactions and ensuring that the model is not biased towards the majority class because the majority class is one of the most common fraud types.

iii) Training & Testing:

Training and testing it is a process of producing the data set for assessment of the machine learning models. Training process consists of giving the model labeled instances, and the model looks for patterns and correlations in the data, and adjusts the internal parameters of the model to reduce the errors. The trained model will be evaluated using unseen data to determine whether it generalises and forecasts correctly. This helps the model to be robust and make relevant predictions on unseen data, therefore helping to accurately detect any fraudulent transactions.

iv) Algorithms:

Random Forest: RF [19] is an ensemble learning approach which consists of a large number of DT constructed during the training process and votes for the mode class predicted by every single tree. It averages the predictions of several trees which increases the prediction accuracy and controls the overfitting. In the research, the RF is utilised because of its robustness to deal with big datasets and its capability to handle imbalanced classes. This makes it perfect for the detection of fraudulent transactions as the patterns could be complicated and varied.

Naive Bayes: A probabilistic classifier which assumes that the features are conditionally independent of each other given the class label, with the Bayes' theorem. It calculates the posterior probability for each class, and selects the class with the greatest posterior probability. The model selected for the project is Naive Bayes [18] because it is simple and efficient in the case of large dataset and where the independence of features is a reasonable assumption, which also helps in detecting frauds by efficiently classifying the transactions.

In generalized notation we write:

$$p(A,B|A) = p(A) * p(B|A) \quad (1)$$

This is read, “probability of A and B given A is equal to the probability of A times probability of B given A is known or have occurred.” This is known as conditional probability and better yet, as conditional probability given another condition or event, since the probability is being calculated under the condition of a previous event or condition.

XGBoost: XGBoost is a gradient boosting framework that sequentially develops models using DT, and optimises performance using regularisation and boosting approaches. It is noted for its exceptional precision, speed and efficiency. To enhance the accuracy of fraud identification with large amounts of data and intricate attribute interactions, XGBoost [20] is used in the study to achieve higher prediction accuracy and recall.

$$y^{\wedge}i = \sum_{k=1}^K f_k(x_i) \quad (2)$$

The last value (ith) of the array is $y^{\wedge}i$ which is the final value to be projected, and K is the number of trees in the ensemble, $f_k(x_i)$ is the prediction of the K tree for the ith data point.

Ensemble With Soft Voting Classifier (RF + NB + XGB): Soft Voting Classifier takes the prediction of various classifiers (RF, NB and XGB) and calculates the average of each predicted probabilities to get the final prediction. It takes advantage of the best qualities of each model to improve the total forecast accuracy. In the project, this ensemble approach is used to build a more powerful fraud detection system as it collects different patterns and minimises the chances of misclassification.

Stacking Classifier (Bagging Classifier with RF and DT with LightGBM): Ensemble Machine Learning Model of stacking different models like RF, DT with LightGBM using a meta-learner to improve prediction ability. This technique takes advantage of the capabilities of each model and reduces biases. It is utilised in the project to improve the accuracy of fraud detection, by aggregating the prediction of different algorithms,

resulting in a more reliable and generalised detection system.

IV. RESULTS AND DISCUSSION

Accuracy: Accuracy of a test is the ability to distinguish patient from healthy cases. To check for accuracy, determine the percent of true positive and true negative cases out of the total number of cases tested. Mathematically it is:

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (3)$$

Precision: Precision is the proportion of truly positive cases or samples in the presence of all cases or samples that are classified as positive. Precision is determined using the following formula:

$$Precision = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}} \quad (4)$$

Recall: Recall is a measure of the model’s ability to find all relevant instances of a class in machine learning. It quantifies the completeness of a model, based on the number of true positives and all positives, i.e., true positives/total positives.

$$Recall = \frac{TP}{TP + FN} \quad (5)$$

F1-Score: F1 score is a measure of the accuracy of a ML model. Combined model precision and recall scores (Accuracy metric) indicates the frequency at which the model predicted correctly within the dataset.

$$F1 \text{ Score} = 2 * \frac{Recall * Precision}{Recall + Precision} * 100 \quad (6)$$

The performance metrics-F1-score, precision, recall and accuracy for each method are shown in Table (1 to 6). All metrics indicate that the Stacking Classifier always outperforms all other algorithms. A comparison of the metrics for various methods is also included in the tables.

Table.1 Performance Evaluation Metrics for OverSampling – European dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.925	0.925	0.925	0.928
Naive Bayes	0.793	0.801	0.793	0.867
XGBoost	0.937	0.937	0.937	0.939
Ensemble	0.910	0.911	0.910	0.921
Stacking Classifier	1.000	1.000	1.000	1.000

Graph.1 Comparison Graphs for OverSampling – European dataset

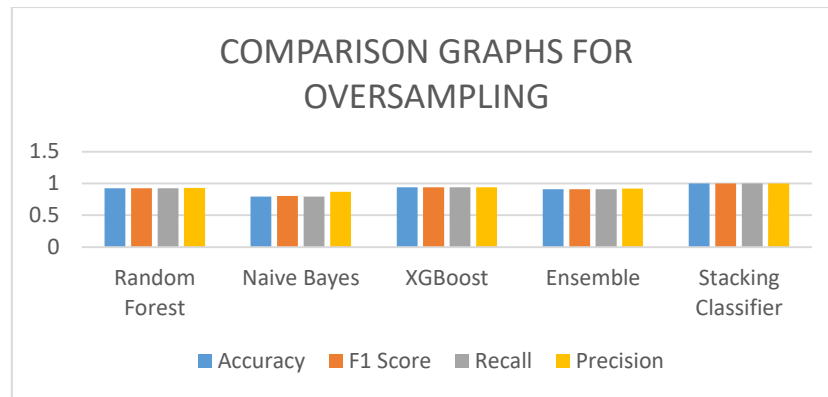


Table.2 Performance Evaluation Metrics for UnderSampling – European dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.878	0.878	0.878	0.878
Naive Bayes	0.787	0.794	0.787	0.854
XGBoost	0.898	0.899	0.898	0.900
Ensemble	0.883	0.884	0.883	0.898
Stacking Classifier	0.954	0.954	0.954	0.959

Graph.2 Comparison Graphs for UnderSampling – European dataset

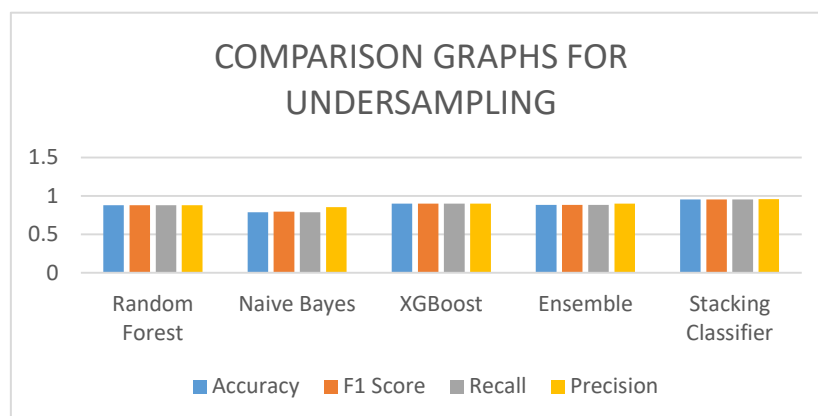


Table.3 Performance Evaluation Metrics for SMOTE – European dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.941	0.941	0.941	0.941
Naive Bayes	0.807	0.813	0.807	0.871
XGBoost	0.955	0.955	0.955	0.956
Ensemble	0.923	0.923	0.923	0.930
Stacking Classifier	1.000	1.000	1.000	1.000

Graph.3 Comparison Graphs for SMOTE – European dataset

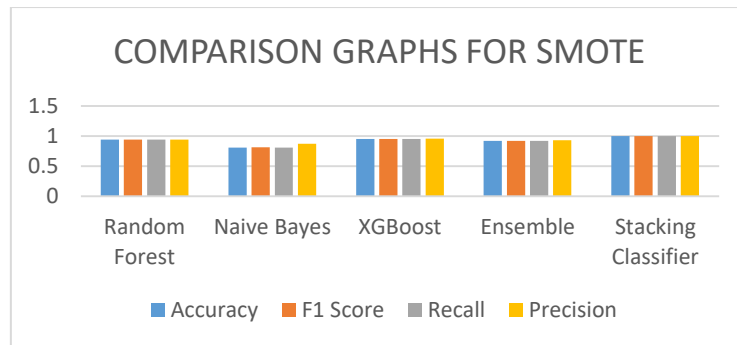


Table.4 Performance Evaluation Metrics for OverSampling – Sparkov dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.861	0.862	0.861	0.881
Naive Bayes	0.806	0.811	0.806	0.860
XGBoost	0.913	0.914	0.913	0.921
Ensemble	0.864	0.866	0.864	0.888
Stacking Classifier	1.000	1.000	1.000	1.000

Graph.4 Comparison Graphs for OverSampling – Sparkov dataset

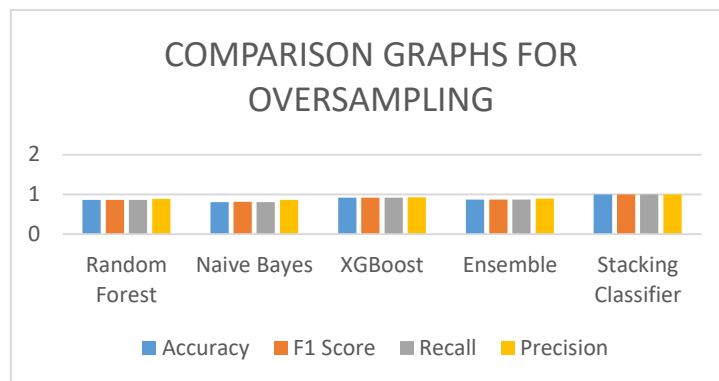


Table.5 Performance Evaluation Metrics for UnderSampling – Sparkov dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.861	0.862	0.861	0.884
Naive Bayes	0.843	0.845	0.843	0.867
XGBoost	0.911	0.911	0.911	0.919
Ensemble	0.865	0.866	0.865	0.888
Stacking Classifier	0.995	0.995	0.995	0.995

Graph.5 Comparison Graphs for UnderSampling – Sparkov dataset

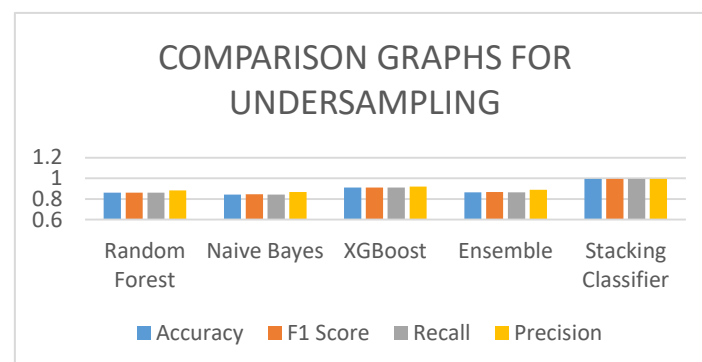
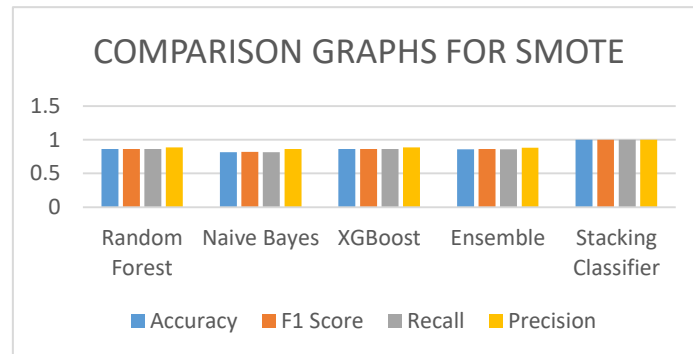


Table.6 Performance Evaluation Metrics for SMOTE – Sparkov dataset

Model	Accuracy	F1 Score	Recall	Precision
Random Forest	0.861	0.863	0.861	0.885
Naive Bayes	0.814	0.819	0.814	0.863
XGBoost	0.863	0.865	0.863	0.885
Ensemble	0.860	0.861	0.860	0.883
Stacking Classifier	0.999	0.999	0.999	0.999

Graph.6 Comparison Graphs for SMOTE – Sparkov dataset



Graphs (1 to 6). The accuracy, F1-score, recall and precision. The Stacking Classifier performs better than the other models, getting the highest values in all the measures. These are graphically illustrated above.

V. CONCLUSION

Malicious transaction on credit cards can be identified which helps in preventing massive amount of financial loss to the institutions. These losses continue to rise, despite the efforts of financial stakeholders, and there is a need to improve detection systems. Ensemble models in particular, and machine learning methods in general, have proven to be very promising for overcoming this challenge. Among all the models mentioned, the ensemble techniques have proven to be very effective in detecting fraudulent transactions, such as Soft Voting Classifier and Stacking Classifier, which comprises of RF, NB and XGBoost. Combining these models, they yield the superior accuracy, precision, f1-score and recall for detecting anomalies. In addition, the Stacking Classifier obtained great accuracy and is one of the most successful methods in the identification of credit card fraud. These powerful algorithms can significantly reduce the risk of fraud and improve the safety of digital transactions for financial institutions, preventing significant financial losses.

For future research, the authenticity for elements which are multiplied, changed and randomised in the authentication step should be taken into consideration to make fraud detection more robust. Another important aspect will be the focus on the explainability and interpretability of algorithms. Knowing and interpreting the decision making

process of such models can provide deeper understandings of the model behaviour, increase transparency and build trust in the model predictions. This will result in the creation of more powerful fraud detecting systems with greater interpretation.

REFERENCES

- [1] Z. Faraji, "A review of machine learning applications for credit card fraud detection with a case study," SEISENSE J. Manage., vol. 5, no. 1, pp. 49–59, Feb. 2022.
- [2] F. K. Alarfaj, I. Malik, H. U. Khan, N. Almusallam, M. Ramzan, and M. Ahmed, "Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms," IEEE Access, vol. 10, pp. 39700–39715, 2022.
- [3] Nilson Report. Card Fraud Worldwide. Accessed: May 2023. [Online]. Available: <https://nilsonreport.com/>
- [4] N. S. Alfaiz and S. M. Fati, "Enhanced credit card fraud detection model using machine learning," Electronics, vol. 11, no. 4, p. 662, Feb. 2022.
- [5] B. Arora and Sourabh, "A review of credit card fraud detection techniques," Recent Innov. Comput., pp. 485–496, 2022.
- [6] S.Srinidhi, K.Sowmya, and S.Karthika, "Automatic credit fraud detection using ensemble model," in ICT Analysis and Applications. Springer, 2022, pp. 211–224.
- [7] A. Alharbi, M. Alshammari, O. D. Okon, A. Alabrah, H. T. Rauf, H. Alyami, and T. Meraj, "A novel text2IMG mechanism of credit card fraud detection: A deep learning approach," Electronics, vol. 11, no. 5, p. 756, Mar. 2022.
- [8] Kaggle. (2022). European Cardholders Dataset. Accessed: May 2023. [Online]. Available: <https://www.kaggle.com/datasets/mlgulg/creditcardfraud>
- [9] Sparkov Data Generation on Github, Sparkv simulator, 2020.
- [10] D. Varmedja, M. Karanovic, S. Sladojevic, M. Arsenovic, and A. Anderla, "Credit card fraud detection—machine learning methods," in Proc. 18th Int. Symp. INFOTEH-JAHORINA (INFOTEH), Mar. 2019, pp. 1–5.
- [11] S B. E. Raj and A A. Portia, "Analysis on credit card fraud detection methods," in Proc. Int. Conf. Comput., Commun. Electr. Technol. (ICCCET), 2011, pp. 152–156.
- [12] J. M. V and D. Vimal, "Population based optimized and condensed fuzzy deep belief network for credit card fraudulent detection," Int. J. Adv. Comput. Sci. Appl., vol. 11, no. 9, 2020.

- [13] T. Razooqi, P. Khurana, K. Raahemifar, and A. Abhari, "Credit card fraud detection using fuzzy logic and neural network," in Proc. 19th Commun. Netw. Symp., 2016, pp. 1–5.
- [14] M. E. Lokanan, "Financial fraud detection: The use of visualization techniques in credit card fraud and money laundering domains," J. Money Laundering Control, vol. 26, no. 3, pp. 436–444, Apr. 2023.
- [15] B. Lebichot, F. Braun, O. Caelen, and M. Saerens, "A graph-based, semi supervised, credit card fraud detection system," in Proc. 5th Int. Workshop Complex Netw. Appl. (COMPLEX Network). Springer, 2017, pp. 721–733.
- [16] J.O.Awoyemi, A.O.Adetunmbi, and S.A.Oluwadare, "Credit card fraud detection using machine learning techniques: A comparative analysis," in Proc. Int. Conf. Comput. Netw. Informat. (ICCNi), Oct. 2017, pp. 1–9.
- [17] I. Sohony, R. Pratap, and U. Nambiar, "Ensemble learning for credit card fraud detection," in Proc. ACM India Joint Int. Conf. Data Sci. Manage. Data, Jan. 2018, pp. 289–294.
- [18] K. M. Leung, "Naive Bayesian classifier," Polytech. Univ. Dept. Comput. Sci./Finance Risk Eng., vol. 2007, pp. 123–156, Nov. 2007.
- [19] A. Prinzie and D. Van den Poel, "Random forests for multiclass classification: Random MultiNomial logit," Expert Syst. Appl., vol. 34, no. 3, pp. 1721–1732, Apr. 2008.
- [20] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in Proc. 22nd ACM SIGKDD Int. Conf. Knowl. Discovery Data Mining, Aug. 2016, pp. 785–794.