

Dynamic Ensemble Learning for Accurate Credit Card Fraud Detection

Vishal Mahto¹, Pushpak Mahajan², Avinash Patil³, Haresh Tetgure⁴, Vinod Jagannath Kadam⁵, Rajnikant Alkunte⁶

¹ Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: mahtovishal445@gmail.com

² Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: mahajanpushpak40@gmail.com

³ Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: avinashpatil9328@gmail.com

⁴ Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: haresh.tetgure@gmail.com

⁵ Professor, Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: vjkadam@dbatu.ac.in

⁶ Assistant Professor, Dr. Babasaheb Ambedkar Technological University, Lonere, Maharashtra, India.

Email: rajnikalkunte@dbatu.ac.in

Abstract: In financial organisations, credit card fraud detection is an important problem to solve because it is possible for hackers to impersonate legitimate credit card users. We tackle the class imbalance problem in fraud detection using multiple resampling strategies like oversampling, undersampling and SMOTE techniques with various datasets like European Data and Sparkov Data. Ensemble learning is used to boost the classification performance by using multiple algorithms to improve accuracy and resistance to errors. In this paper, an ensemble based framework is proposed, which employs complex resampling methods to enhance model training and prediction. A thorough analysis of several classification models reveals that a Stacking Classifier is the optimum model; it combines many models and gives better accuracy, precision, recall and F1 scores than any other strategies. Such an approach could be a substantial benefit to fraud detection systems, allowing them to detect fraudulent transactions with a high degree of confidence and at the same time reducing the number of false positives. The suggested approach emphasises the significance of ensemble methods and data balancing techniques in dealing with the complicated nature of financial fraud detection.

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

1. INTRODUCTION

The era of digital payment brings an important challenge in detecting CCF. It's about identifying discrepancies in credit card transactions. The scale of this problem can be seen with the number of cardholders who held Visa and MasterCard alone in 2020, which was more than 2.183 million [1]. Nearly 1.4 million reports of identity theft were reported in 2013. Out of which, 393,207 cases were particularly due to credit card fraud [2]. The global cost of credit card fraud increased by 19.3% to \$28.6 billion in 2020 compared to \$23.97 billion in 2017. These losses are expected to grow to \$408 billion by the end of the next decade as reported in the Nilson Report 2022 [3]. Despite all the efforts of the banks and financial institutions, the significance of proper credit card fraud detection techniques is critical as the numbers of credit card frauds keeps on increasing.

A variety of techniques have been developed to fight against credit card misuse. These are generally classified into two types; statistical approaches and ML based approaches [2]. The task of determining fraudulent transactions is a statistical problem of finding outliers in a collection of data. Box and whisker plot, normal distribution and cluster based techniques have been used for this purpose. But the intricacies of credit card transaction data has led to the greater use of more flexible and precise ML algorithms.

ML classifiers learn from the data in the past and predict fraudulent transactions. Common methods used for the identification of CCF are choice trees, support vector machines, neural networks and regression models [3], [4], [5]. These strategies are shown to be beneficial at various levels, but their effectiveness greatly relies on the characteristics and quality of the training set.



To enhance the detecting skills, ensemble learning approaches were introduced. The purpose of ensemble learning is to combine multiple base classifiers for creating a powerful and robust ensemble classifier. One such approach is bagging where multiple samples are generated from the training data with replacement, classifiers are trained independently on each of these samples, and final predictions are made by aggregation, e.g., voting. Boosting is another method of ensembles. Boosting is a technique that consists of building classifiers sequentially, feeding the incorrect ones from the previous models into the next ones to improve the overall performance. For instance, AdaBoost, Gradient Boosting and XGBoost [6].

The third ensemble technique is called stacked generalisation, or stacking, in which several classifiers are used to form a single powerful meta-classifier. Some ideas have been proposed to deal with the complexity of credit card fraud detection: stacking, which utilizes the capability of models.

2. RELATED WORK

With the increase in the number of transactions and criminals getting more sophisticated, credit card fraud detection has been an on-going challenge for financial institutions. A lot of research has been done with application of ML and statistical techniques in reducing the risks of fraudulent transactions. These techniques attempt to increase the accuracy and robustness of fraud detection systems, while minimising false positives and false negatives.

The study of Varmedja et al. [10] was aimed the effectiveness of ML algorithms to detect credit card fraud. They emphasised the necessity of using approaches such as oversampling, undersampling and SMOTE to tackle the imbalanced datasets, which is a prevalent feature of fraud detection jobs. The results showed that ensemble learning techniques, like RF and Gradient Boosting had better performance than the individual models because they have the capacity to generalise well for unseen data. Likewise, Raj and Portia [11] have investigated different techniques of fraud detection and concluded that classifiers of ML are the most suitable tools for detecting anomalies in transaction data. They said the flexibility of ML techniques suits them to changing fraud schemes.

Another potential application of DL is credit card fraud detection. In an effort to increase the accuracy of detection, Vimal et al. [12] suggested a population based optimised and condensed fuzzy deep belief network. Our strategy was robust to highly imbalanced data sets using feature selection and dimensionality reduction methods. Razooqi et al. [13] suggested a hybrid system based on fuzzy logic and neural networks that outperforms traditional ML approaches in terms of accuracy and interpretability in fraud detection. The DL and hybrid approaches used to tackle fraudulent transaction identification problems are highlighted in these works, demonstrating the rising trend of combining these methods.

There are also visualization methods which are employed for enhancing fraud detection. Lokanan [14] investigated the possibility of the use of visualisation tools to find fraudulent patterns in credit card transactions and money laundering. This way, investigators can see intricate correlations and abnormalities in data, which in turn gives them a user-friendly way to analyse large data sets. Visualisation is not a magic bullet but it provides an aid to other detection techniques.

Another new technique for credit card fraud detection is the use of graph-based algorithms. Lebichot et al. [15] proposed a semi-supervised model based on graphs for fraud detection. The nodes were used to represent the transactions, the edges the interactions between transactions and the unusual patterns of activity were discovered as being indicative of fraudulent behaviour. The scarcity of labelled data, which is a typical problem in real-world fraud detection cases, was alleviated by this semi-supervised method.

After all, ML algorithms continue to reign supreme in the field of credit card fraud detection thanks to their versatility and scalability. Awoyemi et al. [16] examined the efficiency of a number of ML classifiers including a simple decision tree, SVM and KNN. They concluded that ensemble approaches worked better, particularly bagging and boosting algorithms. The combination of the predictions of multiple classifiers in ensemble learning reduces the risk of overfitting and increases generalisation, thereby improving detection accuracy.

To address these instances, hybrid models have been applied, and solutions have been developed. Hybrid Models of ML and statistical techniques is a holistic approach to Fraud detection. For instance, typical outlier detection methods (such as clustering, density-based methods) could be applied before the data is sent to ML classifiers and examined in detail. The methods are used in a multi-layered approach, leveraging the best aspects of both methods to give full coverage of the fraudulent patterns.

ML and statistical methods have achieved a lot of success in fraud detection, but the demand for XAI in this domain is growing. For transparency in decision making and regulatory issues, there is a need to develop interpretable models that can make explicit statements regarding their predictions. This is particularly crucial in

financial applications, as false positives can result in consumer discontent and false negatives can result in substantial financial losses.

3. MATERIALS AND METHODS

The suggested method intends to enhance the credit card fraud detection using a mix of sophisticated data sampling, feature selection and ML techniques. The data set includes the European Data and Sparkov Data, which will be applied to the data sampling techniques that are used to overcome class imbalance, including OverSampling, UnderSampling and SMOTE. PCA decomposition will be used for dimensionality reduction and feature selection will be done to improve the performance of the model. To capture the various patterns in the data, we will try to build and analyse various ML algorithms such as [19] (RF), [18] (NB) and [20] (XGB). Also, for accuracy and robustness, ensemble methods will be used, such as Soft Voting Classifier aggregating RF, NB and XGBoost and Stacking Classifier which uses Bagging Classifier with RF and DT combined with LightGBM. These strategies are aimed at creating a very effective and efficient fraud detection system.

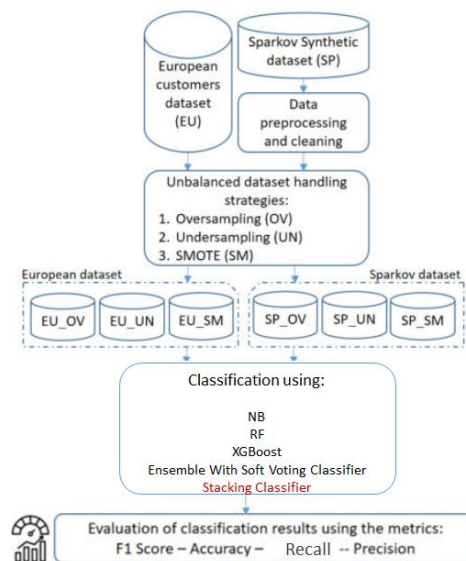


Fig.1 Proposed Architecture

This design uses two datasets (European and Sparkov) to detect fraud transactions. The data is first preprocessed and cleansed. To overcome unbalanced dataset problems, the approaches of oversampling and undersampling and SMOTE are employed. The classification is done using NB, RF, XGB, Ensemble with Soft Voting Classifier and Stacking Classifier. The performance of the model is evaluated using F1 Score, Accuracy, Recall and Precision measures.

i) Dataset Collection:

The European and Sparkov datasets are transaction data with features like as time, quantity, and anonymised details used to identify fraudulent transactions using ML. Labelled fraud indicators are also included in the databases for study.

The dataset that has been used for credit card fraud detection is the European Data [8] that contains 31 anonymised numerical features (V1 to V28), a target variable “Class” indicating the normal transactions (0) or fraudulent transactions (1), and transaction data like time of transaction, amount of transaction etc., and has been totalised to 284,807 transactions. All features are continuous (apart from the Class column) and there are no missing values in the dataset. The data offers a rich source for finding 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	0.239599	0.098698
1	0.0	1.191857	0.266151	0.166480	0.448154	0.060018	-0.082361	-0.078803	0.085102
2	1.0	-1.358354	-1.340163	1.773209	0.379780	-0.503198	1.800499	0.791461	0.247676
3	1.0	-0.966272	-0.185226	1.792993	-0.863291	-0.010309	1.247203	0.237609	0.377436
4	2.0	-1.158233	0.877737	1.548718	0.403034	-0.407193	0.095921	0.592941	-0.270533

Fig.2 Dataset Collection Table – European

The Sparkov data is made up of 1296675 samples of which 7 columns contain transaction details such as trans_date_trans_time, transaction amount (amt), and location details such as city population (city_pop), and ZIP code (zip). It also contains details of the customers such as their date of birth (“dob”) and a target variable “is_fraud” that specifies whether a transaction corresponds to a fraudulent or a normal transaction (1 or 0). It's a suitable dataset for fraud analysis because there aren't any nulls in any of the columns and the data types are all numerical and categorical.

	Unnamed: 0	trans_date_trans_time	cc_num	merchant	category	amt
0	0	2019-01-01 00:00:18	2703186189652095	fraud_Rippin, Kub and Mann	misc_net	4.97
1	1	2019-01-01 00:00:44	630423337322	fraud_Heller, Gutmann and Zieme	grocery_pos	107.23
2	2	2019-01-01 00:00:51	38859492057661	fraud_Lind-Buckridge	entertainment	220.11
3	3	2019-01-01 00:01:16	3534093764340240	fraud_Kutch, Hermiston and Farrell	gas_transport	45.00
4	4	2019-01-01 00:03:06	375534208663984	fraud_Keeling-Crist	misc_pos	41.96

Fig.3 Dataset Collection Table - Sparkov

ii) Pre-Processing:

We will discuss some of the pre-processing steps like data processing, visualization of data and feature selection techniques using PCA decomposition and data sampling method.

a) Data Processing: When preparing the raw data for ML, data processing is required. This includes cleansing the dataset by dealing with missing values, data kinds, and duplicates. In fraud detection on credit card transactions, transaction data needs to be normalised or standardised especially when using algorithms that rely on scale like neural networks. Then, some encoding of category features might be necessary as well as the removal of irrelevant features. By handling the data properly, the models will be able to discover trends and recognise fraudulent transactions.

b) Data Visualization: To achieve this, visualisation of data can be used to understand the distribution and the link between the features which is important to pre-processing. Some methods that reveal patterns, correlations and anomalies in the data are histograms, scatter plots and heatmaps. Visualising a distribution of legal and fraudulent transactions as well as how much the data was worth and how long the transactions took helps detect patterns and outliers in fraud detection. Visualisation can be used to understand the class imbalance and to determine what next steps are required in the preprocessing such as resampling or feature engineering.

c) Feature Selection: PCA is a method that reduces dimensionality of a dataset by selecting the most important features. It converts the original features to a new set of features that are linearly uncorrelated and ordered by variance. PCA reduces the complexity of the data, maintaining most of the information. For example, when it comes to fraud detection, PCA can be used to determine which features are most significant for variance, and to remove redundant or irrelevant features, resulting in more efficient, effective models.

d) Data Sampling: Data sampling techniques are used to address the issue of imbalance class in the fraud detection data sets. Increase the minority class instances refers to oversampling and reduce the majority class instances refers to undersampling. SMOTE (Synthetic Minority Over-sampling Technique) generates synthetic instances for the minority class thru interpolation between existing instances. The strategies aim to provide well-rounded datasets, reducing the likelihood that the ML algorithms are skewed toward the dominant class and hence are prone to fraud.

iii) Training & Testing:

Preparing the data set that will be used for testing the ML model. During training, the model uses labelled examples to discover correlations and patterns in the data. It optimises internal parameters for the purpose of reducing errors. Once trained, the model is tested with new data to determine its ability to generalise and make predictions. This method guarantees that the model is solid and can predict with certainty the transactions that have not yet occurred, thereby enabling a successful identification of fraudulent transactions.

iv) Algorithms:

Random Forest: RF [19] is an ensemble learning approach, which creates several decision tree models during training and outputs the class that is the majority vote of the trees. It averages out the projections from numerous trees and therefore it enhances accuracy and decreases overfitting. The research uses Random Forest due to its robustness that it can be applied for huge datasets and its capability of handling imbalanced classes, which is suitable in identifying fraudulent transactions since the pattern might be complicated and varied.

Naive Bayes: A classifier that uses Bayes' Theorem and assumes that, for each class, the features are conditionally independent. It estimates the posterior distribution of each class and selects the class that has the maximum probability. In the project the Naive Bayes [18] is employed because of its simplicity and efficiency particularly for large data sets and when it can be assumed that the features are independent. This makes it easier to detect fraud in the process of efficiently categorizing transactions.

In generalized notation we write:

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

This translates to "probability of A and B is equal to probability of A times the probability of B, given A. This is called conditional probability or, better, a joint conditional probability because the probability is calculated under a condition or event.

XGBoost: XGBoost is a gradient boosting framework which constructs models sequentially using decision trees. It is a fusion of optimisation, regularisation and boosting techniques. It is renowned for its unparalleled accuracy, speed and efficiency. In the research, XGBoost [20] is applied to handle the large dataset and complex feature relationship, thus improving the ability to detect fraud and the recall rates of the model.

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

Where $y^{\wedge}i$ is the expected value for the i th data point, K is the number of trees in the ensemble and $f_k(x_i)$ is the forecast of the K th tree for the i th data point.

Ensemble With Soft Voting Classifier (RF + NB + XGB): Ensemble With Soft Voting Classifier (RF + NB + XGB) Soft Voting Classifier is an averaging of the predicted probabilities of different classifiers (RF, NB and XGB) for final classification. It takes advantage of the single models to increase the overall forecast accuracy. This ensemble approach is used in the project to create a stronger fraud detection system which can detect various patterns and reduces the risk of misclassification errors.

Stacking Classifier (Bagging Classifier with RF and DT with LightGBM): Combine multiple models such as RF, DT and LightGBM to boost the predictive performance by aggregating the prediction result of multiple models with a meta-learner. This technique takes advantage of the strengths of each model and minimises biases. It's being utilized in the project to boost the effectiveness of fraud detection. It fuses the forecasts of different algorithms and produces more effective and generalised fraud detection.

4. RESULTS AND DISCUSSION

Accuracy: The ability of a test to discriminate between the patient and healthy cases. The percentage of accurate test results can be determined by the ratio of true positive and true negative cases of all cases tested. This is in a mathematical way:

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

Precision: Precision is the accuracy of positive instances or samples that are correctly classified. Precision is defined as:

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

Recall: When the model is able to find all relevant examples of a class, it has a high rate of recall. It is based on how well a model is able to capture the observations of a class, and compares the number of correct positive observations to the number of total positive observations.

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

F1-Score: F1 score is the precisions of a ML model. The overall accuracy achieved by the model. The accuracy statistic is the percentage of correct predictions the model made on the entire data set.

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

The performance measures (F1-score, precision, recall and accuracy) for each algorithm are evaluated in tables (1-6). The Stacking Classifier outperforms all other methods in all the metrics. Also, the tables contain a comparison of the other algorithms' metrics

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

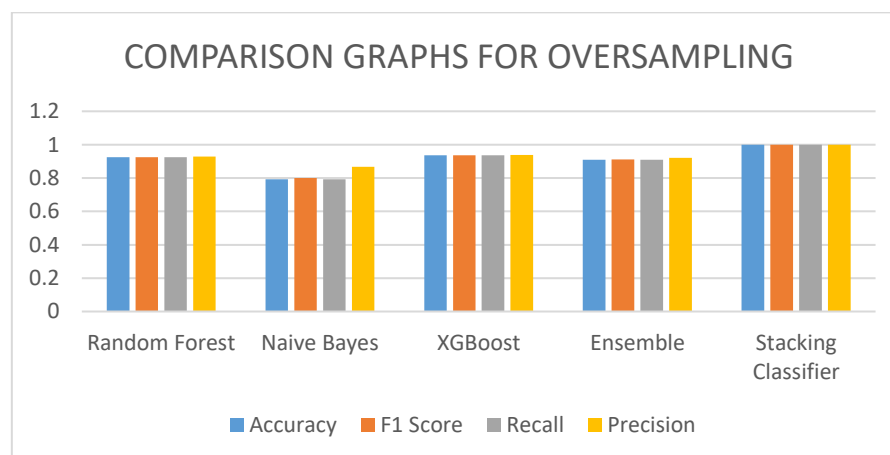


Table.2 Performance Evaluation Metrics for UnderSampling – European dataset

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

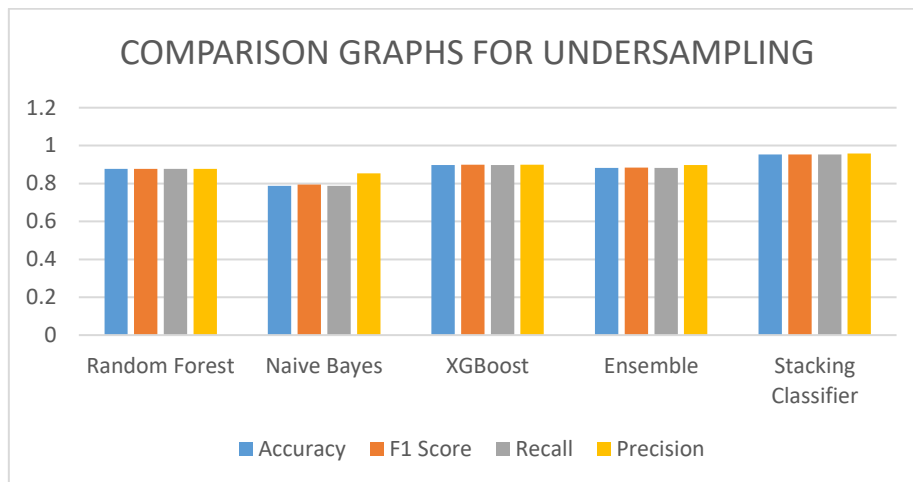


Table.3 Performance Evaluation Metrics for SMOTE – European dataset

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

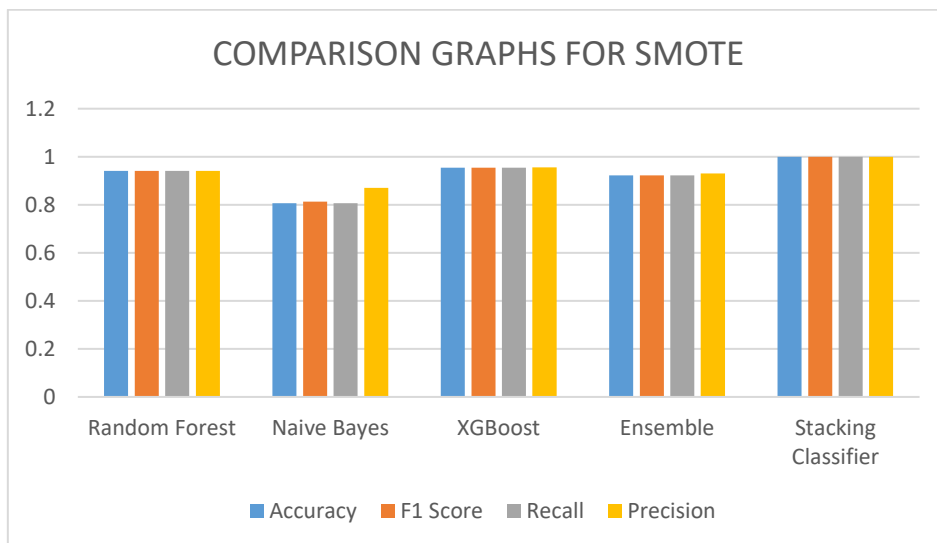


Table.4 Performance Evaluation Metrics for OverSampling – Sparkov dataset

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

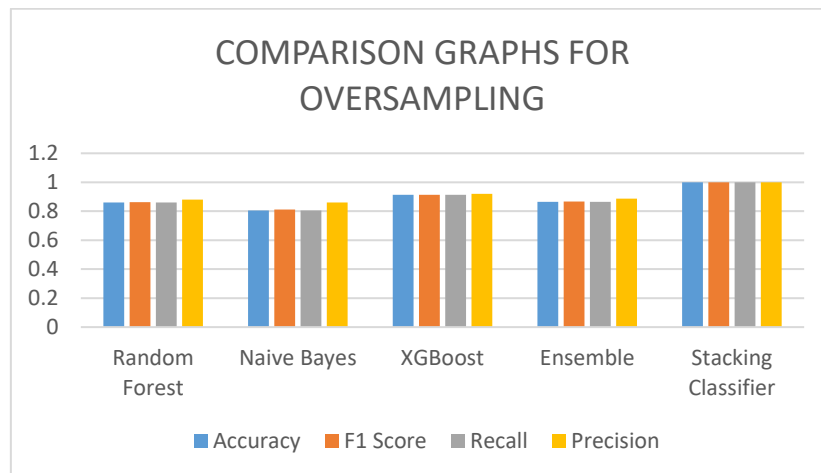


Table.5 Performance Evaluation Metrics for UnderSampling – Sparkov dataset

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

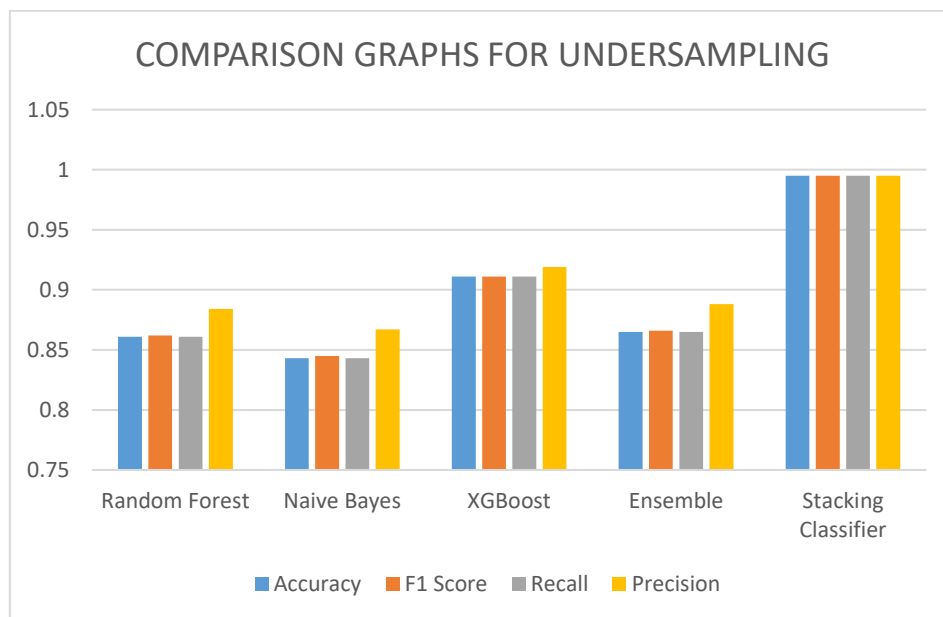
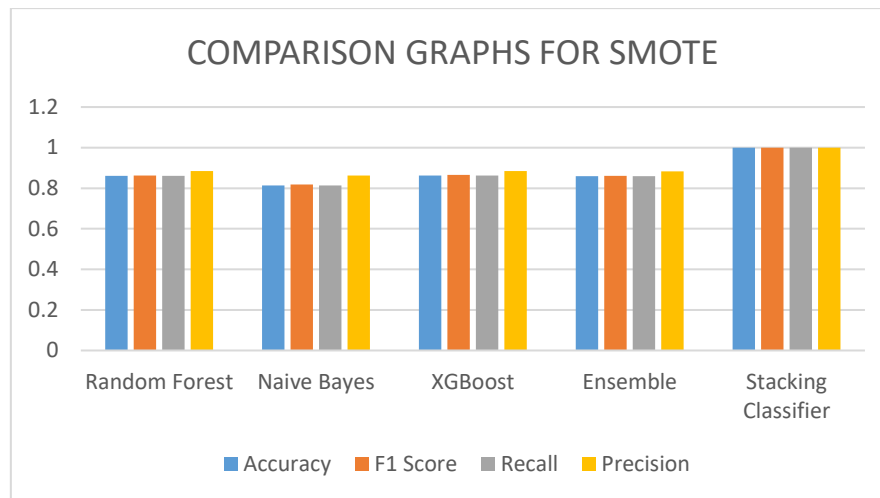


Table.6 Performance Evaluation Metrics for SMOTE – Sparkov dataset

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) show Accuracy (blue), F1-score (orange), recall (grey) and precision (light yellow). The Stacking Classifier outperforms the other models in comparison, attaining the best values for all the measures. The unit conclusions are clearly illustrated on the above graphs.

5. CONCLUSION

Fraudulent transactions on credit cards identification helps to prevent huge financial losses for organisations. However, such losses are still on the rise despite the financial stake-holders' attempts to eradicate them, highlighting the need for more efficient methods of detection. In this respect, ML techniques, particularly ensemble models, show great promise. Ensemble models like Soft Voting Classifier or Stacking Classifier have been proving to be very successful in detecting fraudulent transactions, employing a combination of Random Forest, Naive Bayes, and XGBoost. These models use the strength of each individual classifier to enhance the accuracy, precision, f1-score and recall in anomaly detection. In instance, the Stacking Classifier has achieved excellent accuracy being one of the most successful algorithms for credit card fraud detection. The application of these high-performance algorithms allows financial institutions to greatly minimise the risk of fraud and increase the security of digital transactions, therefore helping to prevent huge monetary losses.

The impact of elements multiplication, change and randomisation during authentication process in enhancing the fraud detection resilience should be considered for future studies. It will be equally important to stress the explainability as well as the interpretability of algorithms. Knowledge and interpretation of the decision making mechanisms of these models can help gain a better understanding of their behaviour, ensure transparency and build trust in their forecasts. These advances will help to build more reliable and interpretable fraud detection systems.

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, "European Cardholders Dataset," 2022. Accessed: May 2023. [Online]. Available: <https://www.kaggle.com/datasets/mlgulg/creditcardfraud>
9. Sparkov Data Generation on GitHub, "Sparkov 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. (ICCN)*, 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.