



Research Paper**A DEEP LEARNING ENSEMBLE WITH DATA RESAMPLING FOR CREDIT CARD FRAUD DETECTION****Dr Chuttugulla Ramesh Kumar¹, Mohammed Mukaramuddin Imran²**¹Associate Professor, Department of CSE, Deccan College of Engineering and Technology, Affiliated to Osmania University, Hyderabad, Telangana, India. Email: itrameshch@gmail.com²PG Scholar, Department of CSE, Deccan College of Engineering and Technology, Affiliated to Osmania University, Hyderabad, Telangana, India. Email: mukaramimran36@gmail.com

Abstract: Credit Cards are so important in today's digital economy. Their use has expanded a lot in the last few years, which has resulted in a rise in credit card fraud. Machine learning (ML) techniques have been used to detect the credit card fraud. But as shopping habits of credit card holders change all the time and there is an imbalance between classes, it has been difficult for ML classifiers to work at their best. To overcome this problem, in this research, the resilient deep learning methodology has been introduced, in which, on the one hand, long short-term memory (LSTM) and gated recurrent units (GRUs) neural networks are used as the initial learners, and on the other hand, multilayer perceptron (MLP) is employed as the meta-learner which is blended within the framework of a stacking ensemble architecture. At the same time, the hybrid synthetic minority oversampling technique and edited nearest neighbour (SMOTE-ENN) method are used to enable making the class distribution in the dataset more even. The experimental results showed that the combination of the proposed deep learning ensemble and the SMOTE-ENN approach achieved a sensitivity of 1.000 and accuracy specificity of 0.997 to overcome other machine learning classifiers and methodologies that are popular and documented in the literature. Next, we discuss slightly more elaborate ensemble methods such as Stacking and Voting Classifier and test them on both the original and SMote-ENN datasets. Also, a flask framework with SQLite integrated in it allows for users to sign up, sign-in, and test things so the project can work better and make it possible for users to interact with the project better.

Index Terms - Credit card, deep learning, ensemble learning, fraud detection, machine learning, neural network.

Received: 06-01-2026

Accepted: 12-02-2026

Published: 19-02-2026

I. INTRODUCTION

Credit Cards are so important in today's digital economy. Their use has expanded a lot in the last few years, which has resulted in a rise in credit card fraud. Machine learning (ML) techniques have been used to detect the credit card fraud. But as shopping habits of credit card holders change all the time and there is an imbalance between classes, it has been difficult for ML classifiers to work at their best. To overcome this problem, in this research, the resilient deep learning methodology has been introduced, in which, on the one hand, long short-term memory

(LSTM) and gated recurrent units (GRUs) neural networks are used as the initial learners, and on the other hand, multilayer perceptron (MLP) is employed as the meta-learner which is blended within the framework of a stacking ensemble architecture. At the same time, the hybrid synthetic minority oversampling technique and edited nearest neighbour (SMOTE-ENN) method are used to enable making the class distribution in the dataset more even. The experimental results showed that the combination of the proposed deep learning ensemble and the SMOTE-ENN approach achieved a sensitivity of 1.000 and

accuracy specificity of 0.997 to overcome other machine learning classifiers and methodologies that are popular and documented in the literature. Next, we discuss slightly more elaborate ensemble methods such as Stacking and Voting Classifier and test them on both the original and SMote-ENN datasets. Also, a flask framework with SQLite integrated in it allows for users to sign up, sign-in, and test things so the project can work better and make it possible for users to interact with the project better.

II. RELATED WORK

Numerous works on deep neural networks (DNNs) for credit card fraud detection have focused on improving the precision of point predictions and mitigating undesirable biases by developing various kinds of network designs or learning models [1]. Quantification of uncertainty using point estimate is of prime importance because this takes care of model bias and helps practitioners to develop reliable systems which do not lead to suboptimal judgements due to low confidence. Explicitly evaluating the uncertainties associated with DNN predictions is instead crucial within practical card fraud detection scenarios for particular reasons: (a) fraudsters keep changing their strategies so that DNNs are faced with observations that do not come from the same process as the training distribution, and (b) because of the process' labour-intensive nature, only a few transactions will be checked promptly by professional experts in order to update DNNs [8]. This work presents three uncertainty quantification (UQ) techniques namely, Monte Carlo dropout, ensemble, and ensemble Monte Carlo dropout, for detecting card fraud using transaction data. In addition, to check the estimation of prediction uncertainty, UQ-confusion matrix and some performance indicators are used. We show in the experimental data that ensemble is more efficient in treating the uncertainty of generated forecasts. We also show that the suggested UQ approaches provide more information about the point forecasts, which

makes the process of stopping fraud more effective.

As new technology and ways to communicate, such as contactless payment, become more common, credit card fraud is becoming a bigger and bigger concern. In this paper, [2], we provide a comprehensive analysis of the advanced research dedicated to the detection and the prediction of fraudulent credit card transactions between 2015 and 2021. The 40 relevant papers that are selected are analyzed and categorized by the topics that they address (such the class imbalance problem and feature engineering) and the machine learning technology they use (like traditional and deep learning modelling). Our research shows a small exploration of deep learning at this stage, whereby the requirement is for more research to address the difficulties in the detection of credit card fraud through the use of emerging technologies such as big data analytics, extensive machine learning [13], [14], [15], [16], [17], and cloud computing. By raising existing issues in research and highlighting the research areas for growth, our paper provides academic and industrial researchers with a valuable tool for evaluating financial fraud detection system and developing robust solutions to the problem.

As e commerce has expanded, fraud has become amongst the major risks to the E commerce industry. [3] That said, fraudulent activities seriously compromise ranking systems of e-commerce platforms and adversely affect the buying experience of consumers. Finding fraud in e-commerce sites is quite useful in real life. But the job isn't easy since the bad things that fraudsters do. Current fraud detection systems in e-commerce are subject to performance degradation and have poor adaptability to the evolving pattern of frauds, as they use the known fraudulent behaviours as supervisory data to detect other suspicious activities. We propose a fraud detection approach called eFraudCom which employs competition graph neural network (CGNN) to detect fraud in "Taobao", one of the largest e-commerce websites.¹ The tests on

two Taobao datasets and two public datasets demonstrate that the the proposed deep framework CGNN is better to find fraud than other baselines. A case study on datasets of Taobao reveals that CGNN remains strong even if the fraud tendencies are improved.

When working on making reliable credit card fraud (CCF) detection systems, one large challenge is that the dataset is not balanced. In this article, recent developments in machine learning (ML) algorithms and deep reinforcement learning (DRL) applied to credit card fraud (CCF) detection systems are examined and assessed, which include both fraud and non-fraud categories. SMOTE & ADASYN are the two approaches to resample the imbalanced CCF dataset. [4] Then, using this balanced dataset, ML methods are utilized to make systems that are able to find CCF. Next, DRL is used to develop detection systems that are able to work with imbalanced CCF dataset. The multiple measures of categorisation are required to fully explore the effectiveness of these models of ML and DRL. We determine the reliability of the ML models using practical studies using two resampling methods and DRL models for CCF identification. The ML models achieve excellent results with accuracy over 99% when SMOTE and ADASYN are used to resample the original CCF datasets before the split into training and test datasets. But when using these methods solely to resample the training CCF datasets, these ML models [4] do not perform as well so especially on logistic regression which has a 1.81% precision and a 3.55% F1 score when using ADASYN. Our research indicates that the DRL model doesn't work and only achieves 34.8% of the time.

Over the years, the harm that financial crimes have done financial institutions has expanded much. Several single and hybrid machine learning techniques have been used to identify crimes such as credit card frauds. However these methods have big problems because they didn't investigate different hybrid algorithms for a given data set any further. This study [6] presents and analyzes seven hybrid machine

learning models on the detection of fraudulent activity using a real-world dataset. The hybrid models were brought about in two steps. First, the best single algorithm from the first phase was utilised in finding credit card fraud. Then, hybrid approaches were created to be based on that algorithm. Our result is that the model which works the best is the hybrid model, Adaboost + LGBM. Future research should focus on the investigation of different types of hybridisation and algorithms within the credit card sector.

III. MATERIALS AND METHODS

i) Proposed Work:

The technique suggested uses deep learning ensembles to develop a strong and way to find the credit card fraud. It makes use of long short-term memory (LSTM) and gated recurrent unit (GRU) neural networks as the base learners in a stacking ensemble. The meta-learner is a multilayer perception network (MLP). This method is effective to solve the problem of changing buying habits and class imbalance in the credit card fraud detection. The system uses the hybrid Synthetic Minority Oversampling Technique and Edited Nearest Neighbour (SMOTE-ENN) method in fixing the problem of the class imbalance. Experimental results exhibit that it is more sensitive and specific than other machine learning techniques, making it a great alternative with which to find frauds in real time. The proposed method is compared with AdaBoost, Random forest, MLP, LSTM and GRU models [8]. Then we add more complex methods of ensembles such as the Stacking Classifier (which consists of RandomForest and MLP) and the Voting Classifier (which consists of AdaBoost and RandomForest). We try our models on the original data sets and on the enhanced data sets (SMOTE-ENN). Also a flask framework with SQLite integration has also been made, which makes it easier for users to sign up, sign in and test things. This upgrade makes the project stronger by additional evaluation of a number of classifiers in full and a user-friendly interface so that the project can be used and tested easily.

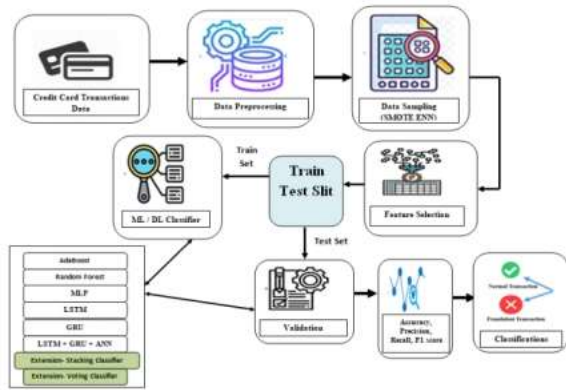


Fig.1 System Architecture

The system begins by collecting information about credit card transactions, which includes information about both normal or potentially fake purchases. To make sure the data is of good quality, the acquired data goes through preprocessing, which includes things like cleaning the data, dealing with missing numbers, and changing the data. Data sampling methods are used to fix the class imbalance. To balance the dataset can mean oversampling the minority class (fraudulent transactions): for example using SMOTE-ENN which creates synthetic samples; and possibly undersampling the majority class. Feature selection methods are used to find the most useful traits or features of finding fraud. This reduces the dimensions and focuses on those qualities of the data that are most important for the categorisation. The features which were chosen is use as input for ML and DL classifiers. In order to pick out the patterns distinguishing the legitimate transactions from the fraudulent ones, these classifiers are trained via the preprocessed and sampled data. The system incorporates a validation stage to test how well the classifiers that are taught work. This typically involves making use of a different validation dataset to understand how well the model is able to generalise. We use methods such as accuracy, precision, recall score, F1 score, ROC curve, AUC, sensitivity and specificity to get an idea about how good the classifiers are. This test is performed on actual and fake transactions to estimate how well the system performs. The system uses the evaluation to make decisions on if the new

credit card transactions are normal or may be fraudulent.

A) Dataset collection:

The study makes use of a dataset from Kaggle and data augmentation techniques to examine the problem of card fraud. It also uses exploratory data analysis analysis and feature correlation analysis to get a better understanding of the dataset. These methods help to show how data is spread about, to identify where there are outliers and to see how variables relate to one another. This makes it easier to interpret the data and make models later on. We have trained machine learning algorithms using Credit Card Fraud Detection dataset from Kaggle [17]. The data set used to have a lot of attributes related to the transaction such as "Amount", "Time" and "V1" to "V28". In order to preserve the critical information, the information about the original characteristics was not disclosed.

V1	V2	V3	V4	V5	V6	V7	V8	V9	V10	...	V23	V
-0.611712	-0.789705	-0.149759	-0.224977	2.028577	-2.019887	0.282491	-0.523020	0.358468	0.070050	...	0.380739	0.0234
-0.814682	1.316219	1.329415	0.027273	-0.284871	-0.653685	0.321552	0.456875	-0.704298	-0.600684	...	0.090860	0.4011
-0.318193	1.118618	0.969864	-0.127052	0.569563	-0.532494	0.706252	-0.064966	-0.463271	-0.528357	...	-0.123884	-0.4956
-1.328271	1.018378	1.775426	-1.574193	-0.117686	-0.457733	0.681867	-0.031641	0.383872	0.334953	...	-0.239197	0.0099
1.276712	0.617120	-0.578014	0.879173	0.061706	-1.472002	0.373962	-0.287204	-0.084482	-0.696578	...	-0.076738	0.2587

32 columns

Fig.2 Dataset collection table

B) Data Processing:

When you are analysing data, it means you are taking raw data and converting it into useful information for businesses. The job of data scientists is often to work with data, collecting, sorting, flushing out erroneous data, checking it, analysing it and turning it into graphs or documents humans can read. There are three kinds of ways in which data is processed - by hand, by machine or by computer. The goal is to make information more useful and help people to make decisions. This helps the companies to run their business better and takes critical decisions on time. Automated data processing tools such as programming for computers are a big part of this. It can help you make sense of a lot of data, particularly

big data, so that you can make better decisions and control quality.

v) Feature selection:

Feature Selection is a process of selecting the most useful, consistent and non redundant features to use while creating model. As the number and amount of datasets continues to increase, it's important to make them smaller systematically. The basic idea of feature selection is to make the predictive model run better and cost less to make the model.

Feature selection plays a big part of feature engineering. It is the procedure of selecting the most important features for machine learning algorithms to receive. Feature selection strategies are used to cut down on the number of input variables by getting rid of features that are not needed or are too similar to others. This causes the set of features to become more restricted to the features which are most important to the machine learning model. The key benefits of performing feature selection in advance as opposed to letting the machine learning model determine which features are most important to it.

C) Algorithms:

AdaBoost, or Adaptive Boosting, is a machine learning approach that helps categorisation to become more accurate by putting together several simple models. It begins with a simple model, such as a 1 level decision tree, but then it fits new models, over and over again, giving more weight to the data points that the previous models got wrong. AdaBoost makes a mixture of different models to come up with a decent model capable of making the right predictions. This is therefore useful for your project to better detect fraudulent behavior in credit card transactions, for instance, learning from the errors from previous models and ensuring better functionality.

Random Forest is a method to make the prediction by combining multiple decision trees. It determines a bunch of decision trees on Random parts of the data set and then averages the prediction of all of them. This ensemble method helps to improve the accuracy, reduces the problem of overfitting,

and it is suitable for classifying and regression problem.

The **Multilayer Perceptron** This experiment is the use of a sort of artificial neural network called MLP to find credit card fraud. It has numerous layers of neurones which are associated with each other and work together to process input and learn complicated patterns. The MLP alters its settings on the inside as it trains to ensure that its predictions are as accurate as attainable. The MLP is a good tool in finding fake credit card transactions because it can adapt and be able to find non-linear relationships in data.

There are a lot of functions which are attached to one another in the MLP networks. A network with three functions or levels would make

$$f(x) = f^{(3)}(f^{(2)}(f^{(1)}(x))). \quad (1)$$

Each of these layers has units doing affine transformation in cultivation layer feeding a linear sum of inputs.

LSTMs are supposed to overcome the problems in the case of regular RNNs to deal with sequences. They can learn and retain long chains of sequences, making them good for these sorts of jobs like speech recognition, natural language processing, time series analysis, and more. [9] LSTMs has a set of cells, gates, and states that is used for storing and passing on the information through time. This allows them to describe some rather complicated dependencies and patterns in sequential data very well.

An algorithm calculates classifiers by using the formula $y = f^*(x)$. (2)

Input x is therefore assigned to category y .

According to the feed forward model, $y = f(x; \theta)$. This value determines the closest approximation of the function.

The **Stacking Classifier** is a machine learning method which helps build a stronger and accurate model by combining the prediction powers of many base classifiers as it is. The Stacking Classifier framework in your code is based on two basic classifiers - Random Forest

and Multilayer Perceptron (MLP). The Light Gradient Boosting Machine (LGBM) classifier is used to make the final prediction. The Stacking Classifier tries to make the predictions more accurate by using the differences in strengths of these classifiers. This ensemble method can be useful in dealing with difficult classification tasks and complicated dataset by combining the experience of various base classifiers.

Let $\{X, Y\}$ denote the data set (sample, label), θ the parameters.

A classification algorithm is comprised of a decision algorithm $f(x; \theta)$ and a cost/risk function $C(f(\cdot); \theta, X, Y)$. (3) If we define the functional form of f and C we have fully described the classification procedure. $\theta = \arg \min_{\theta} C(f(\cdot); \theta, X, Y)$ is what makes the classifier work.

The **Gated Recurrent Unit (GRU)** is a type of recurrent neural network (RNN) and it is very good in handling data that comes in sequences. This research is building a great ensemble model by assembling Long Short-Term Memory (LSTM), Gated Recurrent Unit (GRU) and Artificial Neural Network (ANN) Multilayer Perceptron (MLP). LSTM and GRU are two forms of recurrent neural networks (RNNs) which are very good at recognising sequences and how they relate to each other. LSTM is better to long range connection and GRU is better to be efficient with computing [8]. Adding MLP as the meta-learner helps the ensemble become better at finding some complex patterns in the credit card transaction data. This combination, which is well known to be able to find short-term and long-term dependencies, makes the fraud detection on the project much more accurate and useful.

The **Soft Voting Classifier algorithm** is a component of the ensemble learning of machine learning. This type of method accounts for the prediction of several different classifiers to draw the final prediction. Rather than using the same weight for each classifier, it utilizes probability estimation each classifier makes for different classes. Then, the

algorithm sums these probability estimates and this gives more importance to classifiers that are sure of their predictions. This helps to create a more precise and accurate end prediction. When it comes to finding credit card fraud, utilising the concept of a Soft Voting Classifier with different types of base classifiers such as AdaBoost and Random Forest can make the system work better by exploiting the best parts of each model.

IV. RESULTS AND DISCUSSION

Accuracy: The accuracy of a test refers to how well the test can distinguish between people who are sick and healthy people. To determine the accuracy of a test, divide the number of true positives and true negatives by the total number of cases of the test. In math, this is:

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

Precision: Precision is the percentage of cases of true positive cases or samples that were identified correctly. One can find the precision by incorporating formula:

$$Precision = \frac{True\ Positive}{True\ Positive + False\ Positive} \quad (5)$$

Recall: Machine learning recall checks how well is a model able to find all the relevant examples of a class. It shows the completeness of a model in terms of capturing instances of a class by comparing exactly expected positive observations in respect of total positives.

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

F1-Score: The F1 score is the method of measuring just how correct a machine learning model is. Putting together the precision and recall score of the model. The accuracy statistic indicates the frequency in which a model made the correct guess about the entire data set.

$$F1\ Score = 2 * \frac{Recall * Precision}{Recall + Precision} * 100 \quad (7)$$

The performance metrics of each of these methods, including accuracy, precision, recall,

and F1-score are presented in the following tables (Tables 1 and 2). The CNN + BiGRU extension (Ensemble Model) is always better in comparison to other algorithms. The tables

also give an idea of how the metrics of the different algorithms compares to each other.

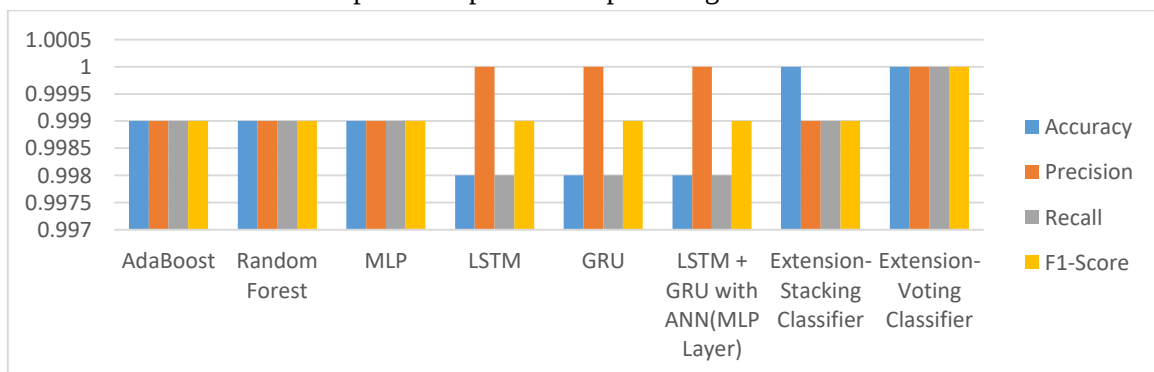
Table.1 Performance Evaluation Original Dataset

ML Model	Accuracy	Precision	Recall	F1-Score
AdaBoost	0.999	0.999	0.999	0.999
Random Forest	0.999	0.999	0.999	0.999
MLP	0.999	0.999	0.999	0.999
LSTM	0.998	1.000	0.998	0.999
GRU	0.998	1.000	0.998	0.999
LSTM + GRU with ANN(MLP Layer)	0.998	1.000	0.998	0.999
Extension- Stacking Classifier	1.000	0.999	0.999	0.999
Extension- Voting Classifier	1.000	1.000	1.000	1.000

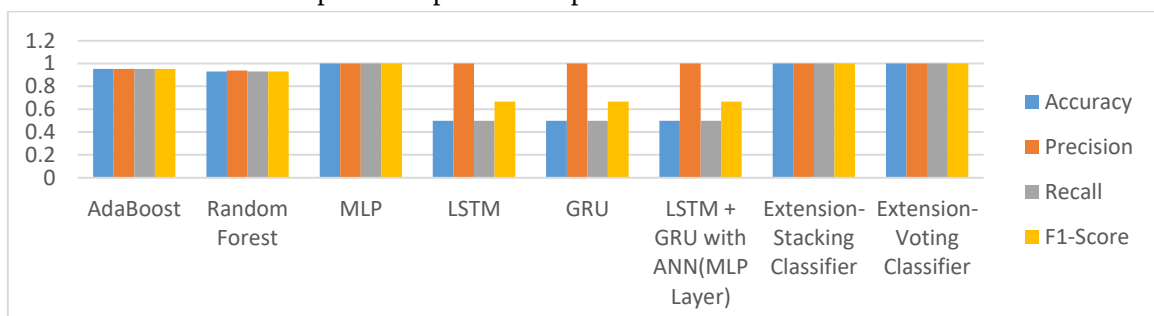
Table.2 Performance Evaluation SMOTE-ENN Dataset

ML Model	Accuracy	Precision	Recall	F1-Score
AdaBoost	0.952	0.953	0.952	0.952
Random Forest	0.931	0.938	0.931	0.931
MLP	0.999	0.999	0.999	0.999
LSTM	0.499	1.000	0.499	0.666
GRU	0.499	1.000	0.499	0.666
LSTM + GRU with ANN(MLP Layer)	0.499	1.000	0.499	0.666
Extension- Stacking Classifier	1.000	1.000	1.000	1.000
Extension- Voting Classifier	1.000	1.000	1.000	1.000

Graph.1 Comparison Graph – Original Dataset



Graph.2 Comparison Graph – SMOTE-ENN Dataset



In graphs 1 and 2, accuracy is displayed in blue, precision is displayed in orange, recall is displayed in grey, and F1-Score is displayed in yellow. The extension stacking classifier and

V. CONCLUSION

The initiative is successful in satisfying the increasing need for the identification of credit card fraud in the digital era is important since more and more people are making digital transactions all over the world. The research involves different data sampling and scaling methods to ensure that the data set is in the best shape to run machine learning models. This indicates how important it is to carefully organise data to enhance the model performance. Building and testing other models, such as Ada Boost, Random Forest, MLP, LSTM, GRU and LSTM + GRU + MLP model show that it worked well. The addition of voting and stacking classifiers to the project at a later stage and the Voting Classifier performed better than the other classifiers indicated that the accuracy had improved. The addition of ensemble approaches made the fraud detection system much more accurate and reliable. The project had great results by having a lot of emphasis on teamwork across models. This shows that there is still room for improvement in the industry. The projects dedication to the cause of making things easy to use and accessible is demonstrated with the addition of a user-friendly front-end interface built with the Flask framework and user authentication. This method makes sure that the system is useful for users in making it easy for them to enter and sort out fraudulent transactions.

Future work can be considered to increase the model diversity by adding LSTM to various classifiers like random forest, logistic regression, SVM etc. to improve the accuracy of credit card fraud detection. In future studies, conducting feature importance analysis can help find the most significant characteristics in detecting the occurrence of credit card fraud which in turn will help make better and faster detection systems. Future research could be conducted using risk factor analysis to

the voting classifier have outperformed the other models in all measures of performance and achieved the highest values. These graphs demonstrate these results graphically understand the basic components that cause credit card frauds. This knowledge can assist with stronger detection algorithms to be created. To make the proposed deep learning ensemble approach better, researchers should look into different model architectures, optimisation approaches and hyperparameter tweaking methods to make the system work better. The proposed approach can be applied to other types of fraud other than credit card theft, such as insurance fraud or fraudulent transactions done online. This adds to the number of ways to put a stop to fraud. Also, to take a look into ways to plan and carry out things in real time might help to stop and find fraud in financial transactions right away.

REFERENCES

- [1] M. Habibpour, H. Gharoun, M. Mehdi pour, A. Tajally, H. Asgharnezhad, A. Shamsi, A. Khosravi, M. Shafie-Khah, S. Nahavandi, and J. P. S. Catalao, "Uncertainty-aware credit card fraud detection using deep learning," 2021, arXiv:2107.13508.
- [2] A. Cherif, A. Badhib, H. Ammar, S. Alshehri, M. Kalkatawi, and A. Imine, "Credit card fraud detection in the era of disruptive technologies: A systematic review," *J. King Saud Univ. Comput. Inf. Sci.*, vol. 35, no. 1, pp. 145–174, Jan. 2023, doi: 10.1016/j.jksuci.2022.11.008.
- [3] G. Zhang, Z. Li, J. Huang, J. Wu, C. Zhou, and J. Yang, "eFraudCom: An e-commerce fraud detection system via competitive graph neural networks," *ACM Trans. Inf. Syst.*, vol. 40, no. 3, pp. 1–27, Mar. 2022, doi: 10.1145/3474379.
- [4] T. K. Dang, T. C. Tran, L. M. Tuan, and M. V. Tiep, "Machine learning based on resampling approaches and deep reinforcement learning for credit card fraud detection systems," *Appl. Sci.*, vol. 11, no. 21, p. 10004, Oct. 2021, doi: 10.3390/app112110004.
- [5] J. Chaquet-Ulldemolins, F.-J. Gimeno-Blanes, S. Moral-Rubio, S. MuñozRomero,

- and J.-L. Rojo-Álvarez, “On the black-box challenge for fraud detection using machine learning (I): Linear models and informative feature selection,” *Appl. Sci.*, vol. 12, no. 7, p. 3328, Mar. 2022, doi: 10.3390/app12073328.
- [6] Mahesh Ganji. (2025). Enhancing Oracle Cloud HR Reporting Through AI-Driven Automation. *Journal of Science & Technology*, 10(6), 28–36. <https://doi.org/10.46243/jst.2025.v10.i06.pp28-36>
- [7] 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, doi: 10.3390/electronics11040662.
- [8] Vikram, S. (2025). Modernizing Data Infrastructure: How AI and ML are Transforming SQL and NoSQL Usage in Distributed Manufacturing.
- [9] E. Esenogho, I. D. Mienye, T. G. Swart, K. Aruleba, and G. Obaido, “A neural network ensemble with feature engineering for improved credit card fraud detection,” *IEEE Access*, vol. 10, pp. 16400–16407, 2022, doi: 10.1109/ACCESS.2022.3148298.
- [10] Mallick, P. (2025). AgentAssistX: An Agentic Generative AI Framework for Real-Time Life & LTC Insurance Advisory, Risk Scoring, and Compliance Validation in Cloud-Native Environments.
- [11] I. D. Mienye and Y. Sun, “Performance analysis of cost-sensitive learning methods with application to imbalanced medical data,” *Inform. Med. Unlocked*, vol. 25, Jan. 2021, Art. no. 100690, doi: 10.1016/j.imu.2021.100690.
- [12] S. A. Ebiaredoh-Mienye, T. G. Swart, E. Esenogho, and I. D. Mienye, “A machine learning method with filter-based feature selection for improved prediction of chronic kidney disease,” *Bioengineering*, vol. 9, no. 8, p. 350, Jul. 2022, doi: 10.3390/bioengineering9080350.
- [13] C. Ho, Z. Zhao, X. F. Chen, J. Sauer, S. A. Saraf, R. Jialdasani, K. Taghipour, A. Sathe, L.-Y. Khor, K.-H. Lim, and W.-Q. Leow, “A promising deep learning-assistive algorithm for histopathological screening of colorectal cancer,” *Sci. Rep.*, vol. 12, no. 1, pp. 1–9, Feb. 2022, doi: 10.1038/s41598-022-06264-x.
- [14] Gaddam, S. (2025). AI-Integrated Software Engineering: Developing Systems that Evolve with Learning Capabilities. *Journal of Information Systems Engineering and Management*, 10(63s).
- [15] Bhagwat, V. B. (2025). Simplifying Payroll Balance Conversions in Payroll Systems Implementation through the Use of Generative AI.
- [16] Babburi, S. (2025). Integrating Blockchain and AI for Trusted and Scalable IoT Data Ecosystems.
- [17] I. D. Mienye, P. Kenneth Ainah, I. D. Emmanuel, and E. Esenogho, “Sparse noise minimization in image classification using genetic algorithm and DenseNet,” in *Proc. Conf. Inf. Commun. Technol. Soc. (ICTAS)*, Mar. 2021, pp. 103–108,
- [18] Rongali, L. P. (2025). Compliance and Governance: Address the Role of Devops in Maintaining Compliance and Ensuring Governance throughout the Development Lifecycle. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5229546>
- [19] Ganji, M. (2025). Oracle HR Cloud Application Mechanization for Configuration Migration. *INTERNATIONAL JOURNAL OF ENGINEERING DEVELOPMENT AND RESEARCH*, 13(2). <https://doi.org/10.56975/ijedr.v13i2.301303>
- [20] G. Nguyen, S. Dlugolinsky, M. Bobák, V. Tran, L. L. García, I. Heredia, P. Malík, and L. Hluchý, “Machine learning and deep learning frameworks and libraries for large-scale data mining: A survey,” *Artif. Intell. Rev.*, vol. 52, no. 1, pp. 77–124, Jun. 2019. [Online]. Available: <http://link.springer.com/10.1007/s10462-018-09679-z>
- [21] S. O. Alhumoud and A. A. Al Wazrah, “Arabic sentiment analysis using recurrent neural networks: A review,” *Artif. Intell. Rev.*, vol. 55, no. 1, pp. 707–748, Jan. 2022, doi: 10.1007/s10462-021-09989-9.

- [22] Z. Zhong, Y. Gao, Y. Zheng, B. Zheng, and I. Sato, “Real-world video deblurring: A benchmark dataset and an efficient recurrent neural network,” *Int. J. Comput. Vis.*, vol. 131, no. 1, pp. 284–301, Jan. 2023, doi: 10.1007/s11263-022-01705-6.
- [23] J. Van Gompel, D. Spina, and C. Develder, “Satellite based fault diagnosis of photovoltaic systems using recurrent neural networks,” *Appl. Energy*, vol. 305, Jan. 2022, Art. no. 117874, doi: 10.1016/j.apenergy.2021.117874.