

A New RF-MLP Framework Using Stacking Technique for Financial Fraud Detection

Hussam Mezher Merdas*

Department of Accounting, College of Management and Economics, University of Warith Al-Anbiyaa, Kerbala, Iraq.

Received 25 December 2024; Accepted 2 March 2025

Abstract

With the huge financial transfers taking place over the Internet, it is necessary to provide precise mechanisms to ensure that financial fraud does not occur. In this study, a new framework was proposed that combines deep learning and Ensemble techniques using the stacking mechanism to detect financial fraud. The Random Forest (RF) and Multilayer Perceptron (MLP) algorithms were employed in the first stage of the stacking represented by base models, and the Gradient Boosting algorithm was employed in the Meta-Model stage. Where a huge dataset consisting of 31 columns and 284,807 rows was processed to simulate large financial flows. The results were excellent, as the model gave 100% prediction accuracy with $F1 = 100$, as it was able to detect all cases of financial fraud in the testing stage and distinguished them from non-suspicious cases. This indicates that the mechanism used is promising and can be employed in the service of companies and countries in general to avoid financial fraud.

Keywords: Ensemble techniques, Random Forest, Multilayer Perceptron, Gradient Boosting

1. Introduction

In light of the progress made in the technical financial field at present, countries suffer from the problem of financial fraud. Where individuals or specific groups strive to attack the technical data of countries or banks to carry out suspicious financial transfers. These countries or financial institutions are exposed to huge losses [1]. Despite the global effort to eliminate these frauds, some of them go undetected. This requires a combination of technical and scientific efforts to build artificial intelligence models capable of detecting and preventing financial fraud. Artificial intelligence is a broad field, one of its branches is machine learning [2]. In this study, a machine learning technique will be employed. This technique is concerned with integrating several diverse machine-learning models into one working environment. The goal of this integration is to extract the capabilities of machine learning algorithms in detecting financial fraud. This technique is called Stacking Technique. Stacking is one of the machine learning methods that belongs to the Ensemble Learning category of techniques, which aims to improve the accuracy of models by combining the results of multiple models. In Stacking, several models (known as base models) are trained separately on the same data, and then the outputs of these models are used as inputs to train another model called Meta-Model or Blender, which learns how to combine the outputs of the base models to obtain a more accurate final result [3]. The proposed model in this study relied on the Random Forest (RF) and Multi-layer Perceptron (MLP) algorithms in the basic stage of stacking, as these two algorithms were not chosen arbitrarily, but because the first algorithm is considered one of the most important Ensemble Learning algorithms that always achieves good results, while the MLP algorithm is one of the important deep learning algorithms. As for the Meta-Model stage, the Gradient

Boosting algorithm was chosen because of its great ability to deal with the results.

Zengyi Huang et al. proposed an AI model based on the K-means clustering method to detect financial fraud. It relied on collecting huge amounts of financial transaction data. Cluster analysis was performed on good amounts of financial transaction data and relied upon to detect suspicious patterns and behaviors, and thus identify potential financial fraud. Their proposed study gave good results of up to 96% in identifying financial fraud [4].

Yara et al. proposed a deep learning model for financial fraud detection based on the Long Short-Term Memory (LSTM) technique. The main goal of their method was to improve the existing detection techniques as well as improve the detection accuracy in light of big data. A real fraud dataset was used to evaluate their proposed model through the use of credit cards and the results were compared with the proposed deep learning model called Auto-encoder model and some other machine learning techniques. The obtained results showed the high performance of LSTM as it achieved an accuracy of 99.95% [5].

Joy Jong-Zong Chen et al. proposed a model for financial fraud detection using a deep convolutional neural network (DCNN) scheme. The latter technique was employed due to the existence of big data. In this study, existing machine learning models, autoencoder models, and other deep learning models were compared with the proposed model. These techniques were employed to evaluate the performance of a real-time credit card fraud dataset. The fraud detection accuracy reached 99% over 45 seconds using the proposed model [6].

The proposed model in this research paper consists of several stages that will be discussed later in this study. The model will be fed with a huge data set, after which this data will be pre-processed by examining the duplicate or missing values to process them if any. Then the second stage begins, which is entering the clean data set into the core of the

*E-mail address: hussam.mezher@uowa.edu.iq

ISSN: 1791-2377 © 2025 School of Science, DUTH. All rights reserved.

doi:10.25103/jestr.182.14

system, which consists of two main parts. These parts are based on the stacking technique, as the data was entered into the base part that contains the RF and MLP algorithms, and then their results were displayed on the Meta-Model part that includes the Gradient Boosting algorithm. In the last stage, the results obtained from the full frame are measured. The obtained accuracy was 100%.

2. Materials and Methods

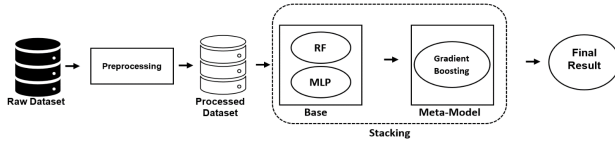


Fig. 1. Main Steps of RF-MLP Framework.

2.1. Dataset

RF-MLP Framework was fed a massive dataset to simulate the movement of real financial data. This dataset was sourced from Kaggle and contained 31 columns and 284,807 rows. It was based on financial entries at specific points in time. Figure 2 below shows how the time series are distributed in the dataset used. Figure 3 below shows the volume of financial transactions that are classified as financial fraud versus transactions that do not contain financial fraud. The figure shows that the training and test data contain (0.2%) financial fraud data and (99.8%) non-financial fraud data. The model will have to predict the first percentage, which is financial fraud operations. Figure 4 below is a 3D chart showing the relationship between the three main variables in this dataset. These variables are (time, class, and amount).

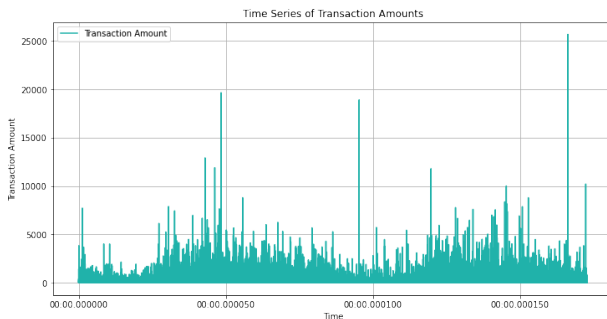


Fig. 2. Financial transactions are indicated in the proposed dataset at specified periods.



Fig. 3. Percentage distribution of financial data between fraud and non-fraud.

The proposed model in this study is based on a framework that consists of two main parts: the RF and MLP algorithms in the first part, which is the Stacking part, and the Gradient Boosting algorithm in the second part, which is the Meta-Model. But before entering the core of the model, it is fed with a large dataset and it is processed and displayed on the framework. Then, in the next step, the data is processed in this framework, followed by measuring the accuracy of the model to obtain the final results, Figure 1 below shows main steps and will be detailed as follows:

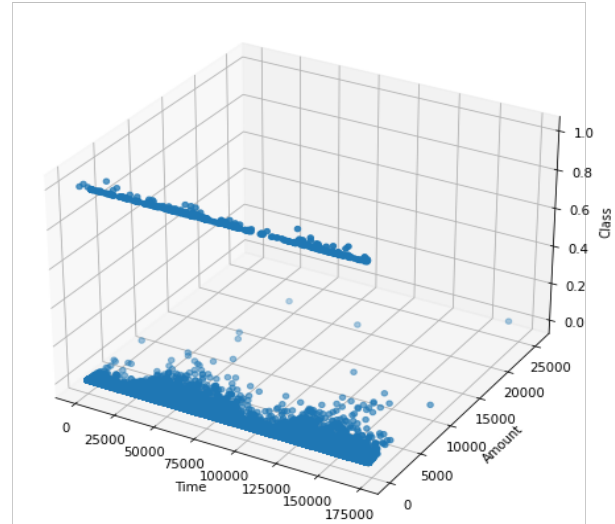


Fig. 4. The relationship between the three important elements in the data set, which are (time, class, and amount).

2.2 Stacking Technique

Stacking is a machine learning technique that relies on another technique, which is the clustering technique. The main goal of the stacking technique is to improve the accuracy of the model by combining the results of multiple models. The stacking technique relies on the diversity of the base models and a meta-model that learns from the outputs of these models to improve the overall performance [7]. The principle of this technique is to focus on training a set of base models independently on the same dataset, and then their predictions are used as inputs to a meta-learner. The meta-model learns the relationship between the outputs of the base models and the true values of the prediction and aims to reduce the errors caused by the individual models. The diversity in this technique is beneficial, as each model may have different strengths and weaknesses. When these models are combined, the impact of individual weaknesses is reduced, which leads to improved overall performance. Figure 5 below shows the stages of the stacking technique.

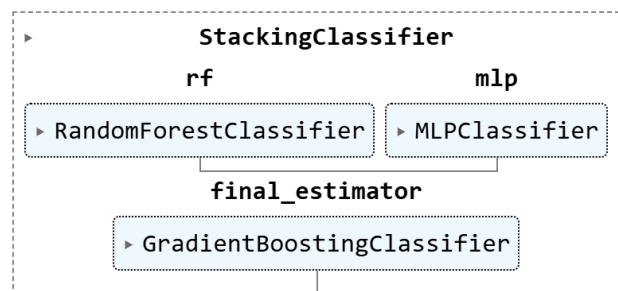


Fig. 5. The stacking technique consists of two stages, the first stage is the RF and MLP algorithms and the second stage is the Gradient Boosting algorithm.

The main steps of the stacking process:

1. Training the base models:

Several base models are trained. In this study, the RF and MLP algorithms were trained using the same training data set.

2. Generating intermediate predictions:

The base models are used to generate predictions on the training data or validation data. These predictions become the inputs for the upper model.

3. Training the upper model:

The upper model (Meta-Model) learns the relationship between the intermediate predictions and the true values of the prediction. The upper model In this study, the Gradient Boosting algorithm was employed to be the upper model.

4. Final prediction:

When predicting new data, this data was passed through the base models, and then their predictions were fed to the upper model to obtain the final result.

2.3 Random Forest

Random Forest is an ensemble technique that belongs to machine learning algorithms and is widely used in classification and regression problems [8]. In this proposed model, it was employed for classification purposes. The mechanism followed in this algorithm is to merge a large number of independent decision trees [9], where each tree contributes to improving the accuracy of the model and reducing variance and bias. The main goal here is to combine the diversity of decision trees by training them on different random parts of the data. Each tree in the forest is created using the Bagging technique, where random samples are selected from the data with replacement (Bootstrapping). In addition, a random subset of features is selected at each split of the tree, which reduces the bias resulting from strong features and increases the diversity of the trees. The basic steps in the work of this algorithm were initially to select a random dataset (with replacement) from the original dataset to train each tree. This allowed each tree to focus on different parts of the data. Each tree was then created using the selected random dataset. At each split node in the tree, a random subset of features was selected instead of using all the features, which reduced the correlation between the trees. Finally, the majority voting technique was adopted to determine the final class (Majority Voting) [10]. This represented the results of the RF algorithm.

2.4 Multi-layer Perceptron

Multilayer neural network (MLP) is one of the deep learning algorithms [11]. MLP is a model of artificial neural networks (ANN) where the nodes (Neurons) are organized into multiple layers. MLP is used to solve classification and regression problems and is capable of handling linear and non-linear data thanks to its advanced structure. In this study, this algorithm was employed for classification purposes.

This algorithm is based on a structure of fully connected layers. These layers consist of an Input Layer: Which receives the raw data or features and hidden layers which process the data using nodes (Neurons) that apply non-linear activation functions. The last layer is the output layer that produces the final output based on the type of problem (classification or regression) [12]. This algorithm consists mainly of nodes (Neurons) and each node or unit in the layers contains a weight that determines the importance of the inputs. A bias to

modify the calculated value. And an activation function that adds non-linearity. The values are passed between the layers through a process called forward propagation, where the inputs are passed through the layers to generate the outputs. The weights and biases are then updated using a backpropagation algorithm that calculates the error between the predicted and actual values [13].

2.5 Gradient Boosting

Gradient Boosting is an advanced machine learning algorithm, employed in this study for classification purposes. This algorithm works on the concept of ensemble learning to develop a strong model by integrating weak learners such as decision trees [14]. Instead of building models independently as in techniques such as Random Forest, models in Gradient Boosting are built sequentially where each tree attempts to correct the errors of the previous model. The mechanism followed in this algorithm is to build a strong predictive model by training weak models one by one [15]. Each new model focuses on reducing the residuals made by the previous model, and it does this by using gradient descent optimization to minimize the loss function. The processing in this algorithm is done by starting with a simple model (such as a mean value for regression or a random prediction model for classification). This initial model provides the initial predictions $F_0(x)$. The errors (r_i) are then calculated based on the difference between the true values (y_i) and the current predictions $F(x_i)$ using the following equation [16]:

$$r_i = y_i - F(x_i) \quad (1)$$

A new decision tree is then trained to minimize the errors (r_i). The goal is for the new tree to predict the gradients of the loss function. The current model is then updated by adding the new model, adjusting its effect using the learning rate η :

$$F_{t+1}(x) = F_t(x) + \eta h_t(x) \quad (2)$$

The previous steps are repeated several times, leading to a gradual improvement in performance.

3. Result

The dataset used is enormous, as mentioned above, it consists of 284,807 rows and 31 columns. Despite this, the RF-MLP model was able to successfully detect and identify all cases of fraud in the training and testing phases, since, as Figure 6 below shows, the model identified all cases of fraud. It correctly identified fraud and gave 100% accuracy, identifying 56,864 cases as not fraud and 98 cases as fraud.

Accuracy: 1.0				
Confusion Matrix:				
[[56864 0]				
[0 98]]				
Classification Report:				
	precision	recall	f1-score	support
0	1.00	1.00	1.00	56864
1	1.00	1.00	1.00	98
accuracy			1.00	56962
macro avg	1.00	1.00	1.00	56962
weighted avg	1.00	1.00	1.00	56962

Fig. 6. The results obtained from the RF-MLP model, show its efficiency and ability to detect financial fraud with 100% accuracy.

Figure 6 above shows us the following results:

Accuracy: The accuracy level is 1.0 (or 100%), which indicates that the model classified all cases in the test set correctly without any errors. This reflects the efficiency of the model and that it worked perfectly and gave excellent results.

Confusion matrix: The confusion matrix shows that out of 56,864 transactions that did not contain financial fraud (category 0), the model correctly identified all 56,864 transactions as legitimate, and 98 transactions (category 1) were accurately identified as fraudulent transactions, without false negative results. This ideal classification also confirms the high accuracy of the model and its ability to classify.

Recall: Recall of 1.00 for both categories indicates that the model identified all actual cases of fraud and non-fraud, without error.

F1-score: The F1-score across both classes combines the model's precision and recall and was 100%, reinforcing that the model balances accuracy and completeness.

Support: The support values, which are the number of each class in the test set (56,864 for class 0 and 98 for class 1), show that the model was able to maintain perfect performance across both classes despite the dataset being highly unbalanced.

Overall, the model demonstrated excellent effectiveness and efficiency in identifying financial fraud, with perfect accuracy and perfect scores on all metrics, which is a great result for fraud detection scenarios where the cost of misclassification can be high for businesses, but such models reduce the occurrence of fraud.

4. Conclusions

Big data of financial transactions requires accurate techniques to detect financial fraud. In this study, the stacking technique was employed as a powerful approach that combines the outputs of multiple algorithms to achieve higher accuracy. The research demonstrated the effectiveness of this technique by obtaining 100% accuracy, indicating the system's ability to distinguish between healthy and suspicious financial transactions with great accuracy. The Random Forest algorithm is one of the best algorithms for classifying high-dimensional and non-linear data. It showed strong performance when used as a base model in the system. The MLP neural network algorithm added additional depth to the model, as it can absorb complex patterns in the data. Its contribution was significant in improving the stacking accuracy. The Gradient boosting algorithm was used as a final layer (Meta-Model) to aggregate the outputs of the base models. This layer showed great effectiveness in improving the results by leveraging the strengths of each model individually and reducing the impact of weaknesses.

The use of additional algorithms such as XGBoost or SVM may improve the stacking performance. Model performance can be improved by using data augmentation techniques to expand the training set. Even with 100% accuracy, it is important to analyze cases that were close to misclassification to identify potential weaknesses.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License.



References

- [1] D. Mangala and L. Soni, "A systematic literature review on frauds in banking sector," *J. Financ. Crime*, vol. 30, no. 1, pp. 285-301, Dec. 2023. doi: 10.1108/JFC-12-2021-0263
- [2] H. M. Merdas and A. H. Mousa, "Food sales prediction model using machine learning techniques," *Int. J. Electr. Comp. Engin.* (2088-8708), vol. 13, no. 6, Dec. 2023. doi: 10.11591/ijece.v13i6.pp6578-6585
- [3] U. Muhammad, J. Laaksonen, D. Romaissa Beddiar, and M. Oussalah, "Domain generalization via ensemble stacking for face presentation attack detection," *Int. J. Comp. Vis.* pp. 1-24, Jun. 2024. doi: 10.1007/s11263-024-02152-1
- [4] Z. Huang, H. Zheng, C. Li, and C. Che, "Application of machine learning-based k-means clustering for financial fraud detection," *Academ. J. Sci. Techn.*, vol. 10, no. 1, pp. 33-39, Mar. 2024. doi: 10.54097/74414c90
- [5] Y. Alghofaili, A. Albattah, and M. A. Rassam, "A financial fraud detection model based on LSTM deep learning technique," *J. Appli. Secur. Res.*, vol. 15, no. 4, pp. 498-516, Sep. 2020. doi: 10.1080/19361610.2020.1815491
- [6] J. I.-Z. Chen and K.-L. Lai, "Deep convolution neural network model for credit-card fraud detection and alert," *J. Artif. Intellig.*, vol. 3, no. 2, pp. 101-112, Jun. 2021. doi: 10.36548/jaen.2021.2.003
- [7] A. A. Alzubaidi, S. M. Halawani, and M. Jarrah, "Towards a Stacking Ensemble Model for Predicting Diabetes Mellitus using Combination of Machine Learning Techniques," *Int. J. Adv. Comp. Sci. Applic.*, vol. 14, no. 12, Dec. 2023. doi: 10.14569/ijacsa.2023.0141236
- [8] Y. M. Mohialden, N. M. Hussien, and S. A. Salman, "Automated Water Quality Assessment Using Big Data Analytics," *Mesopotamian J. Big Dat.*, vol. 2024, pp. 211-222, Nov. 2024. doi: 10.58496/MJBD/2024/015
- [9] O. Sagi and L. Rokach, "Approximating XGBoost with an interpretable decision tree," *Inform. Sci.*, vol. 572, pp. 522-542, Sep. 2021. doi: 10.1016/j.ins.2021.05.055
- [10] V. Derbentsev, V. Babenko, K. Khrustalev, H. Obruch, and S. Khrustalova, "Comparative performance of machine learning ensemble algorithms for forecasting cryptocurrency prices," *Int. J. Engin.*, vol. 34, no. 1, pp. 140-148, Jan. 2021. doi: 10.5829/ije.2021.34.01a.16
- [11] D. C. Edara, V. Sistla, and V. K. Kishore Kolli, "Health App Recommendation System using Ensemble Multimodel Deep Learning," *J. Engin. Sci. Techn. Rev.*, vol. 13, no. 5, Aug. 2020. doi: 10.1007/978-981-16-2123-9_43
- [12] R. Yu, W. Yu, and X. Wang, "Kan or mlp: A fairer comparison," *arXiv preprint arXiv:2407.16674*, Aug. 2024. doi: 10.48550/arXiv.2407.16674
- [13] B. Ali, K. Alakkari, M. Abotaleb, M. M. Mijwil, and K. Dhoska, "MLP and RBF Algorithms in Finance: Predicting and Classifying Stock Prices amidst Economic Policy Uncertainty," *Mesopotamian J. Big Dat.*, vol. 2024, pp. 48-67, May 2024. doi: 10.58496/MJBD/2024/005
- [14] X. Sun and J. Fu, "Many-objective optimization of BEV design parameters based on gradient boosting decision tree models and the NSGA-III algorithm considering the ambient temperature," *Energy*, vol. 288, Feb., 2024, Art. no. 129840, doi: 10.1016/j.energy.2023.129840
- [15] L. Ni et al., "Streamflow forecasting using extreme gradient boosting model coupled with Gaussian mixture model," *J. Hydrology*, vol. 586, Jul. 2020, Art. no. 124901 doi: 10.1016/j.jhydrol.2020.124901
- [16] S. Park, S. Jung, J. Lee, and J. Hur, "A short-term forecasting of wind power outputs based on gradient boosting regression tree algorithms," *Energies*, vol. 16, no. 3, Jan. 2023. Art. no. 1132, doi: /10.3390/en16031132