

X-RANS: A Smote-Enhanced Stacking Ensemble Model for Credit-Card-Fraud Detection

R. Sujeetha^{1,*}, V. Sahaya Sakila²

^{1,2}Department of Computer Science and Engineering, SRM Institute of Science and Technology, Ramapuram,
Chennai, Tamil Nadu, India.
sujeethasrm@gmail.com¹, shylalipna1992@gmail.com²

Abstract: Credit-card fraud continues to be an ongoing and ever-changing threat to the financial sector, further fueled by the increase in online transactions and ever-changing fraud schemes. Traditional rule-based or individual machine learning approaches may not be effective at detecting fraud due to continually evolving patterns and imbalanced datasets, where legitimate transactions greatly outnumber fraudulent ones. To overcome the challenges mentioned above, this research proposes the X-RANS, a hybrid stacking ensemble framework to improve fraud detection accuracy and robustness. Using SMOTE (Synthetic Minority Over-sampling Technique) to address class imbalance and merging the base learners with a meta-classifier to maximise final predictions, the proposed framework combines a stacked architecture of Random Forest, XGBoost, and an Artificial Neural Network (ANN) as base learners. A publicly available credit card transaction dataset is used to train and test the system, revealing that only 0.172% of observations exhibit fraudulent behaviour. X-RANS outperforms baseline samples generated by conventional models on evaluation metrics such as Precision, Recall, F1, and AUC-ROC across all sample sizes. The X-RANS offers a scalable, flexible mechanism for real-world credit card fraud detection that better generalises across the uneven distribution of credit card transactions by leveraging both data augmentation and multi-model learning.

Keywords: Credit Card; Fraud Detection; Imbalanced Data Classification; Stacking Ensemble Learning; Artificial Neural Network; Synthetic Minority; Oversampling Technique; Deep Learning; Machine Learning.

Received on: 24/12/2024, **Revised on:** 19/03/2025, **Accepted on:** 09/05/2025, **Published on:** 09/12/2025

Journal Homepage: <https://www.fmdbpublish.com/user/journals/details/FTSCS>

DOI: <https://doi.org/10.69888/FTSCS.2025.000525>

Cite as: R. Sujeetha and V. S. Sakila, "X-RANS: A Smote-Enhanced Stacking Ensemble Model for Credit-Card-Fraud Detection," *FMDB Transactions on Sustainable Computing Systems*, vol. 3, no. 4, pp. 233-241, 2025.

Copyright © 2025 R. Sujeetha and V. S. Sakila, licensed to Fernando Martins De Bulhão (FMDB) Publishing Company. This is an open access article distributed under [CC BY-NC-SA 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/), which permits copying, remixing, redistributing, transformation, and building upon the material in any medium so long as the original work is properly cited.

1. Introduction

The challenge of credit card fraud remains a crucial and growing threat to financial ecosystems worldwide, costing billions and eroding trust in digital payment systems [1]. The fraud attack surface has increased significantly as financial services transitioned to digital offerings and e-commerce grew rapidly. It's estimated that, in 2023 alone, total global losses from card fraud will exceed \$35 billion, driven in part by organised criminal enterprises using cybercrime and the rapid pace of technological advancement [2]. Fraud detection systems traditionally use rules-based methods to trigger alerts when threshold variables are met and/or anomalies in transactional data are identified. While rule-based systems generally provide

*Corresponding author.

explainability and can analyse alerts in real time, their static nature limits their ability to capture the dynamic nature of fraud schemes. Fraudsters adapt quickly, and many traditional detection approaches fail to detect new patterns [3]; [4]. As a result, traditional fraud detection systems often have high false-negative rates, leaving financial institutions exposed to undetected fraud [18]. The maturation of machine learning (ML) models has made them a viable alternative to traditional approaches, primarily because ML can model and identify nonlinear and other complex relationships within transactional data [5]. In binary classification tasks, such as fraud detection, several algorithms, including Random Forests, Support Vector Machines (SVMs), and Neural Networks, have shown promising results. Nonetheless, every model encounters constraints when applied to actual credit card datasets, which exhibit significant imbalance, as fraudulent transactions make up less than 0.2% of the total entries [6]. The significant class imbalance significantly alters the model's learning process, leading it to favour the larger class (non-fraud). As a result, it becomes harder for the model to detect real fraud correctly [7]. By generating synthetic samples of the minority class, widely used methods such as the Synthetic Minority Over-sampling Technique (SMOTE) effectively address these issues. By creating a more balanced dataset, this method seeks to enable a fairer classification procedure [8].

SMOTE performs well at increasing the Recall Score; however, it can also lead to overfitting if abused [9]. Beyond resampling techniques, researchers have proposed ensemble learning methods such as Bagging, Boosting, and Stacking to improve predictive accuracy. These methods combine multiple base learners into their algorithms to reduce variance and bias, leveraging model diversity to improve predictions [10]. Stacking ensembles have been shown to improve model performance over other methods by aggregating the predictions of numerous classifiers with a meta-learner in the stacking model, effectively combining the strengths of each classifier [11]. Deep Learning (DL) models—particularly Recurrent Neural Networks (RNNs), such as Long Short-Term Memory (LSTM)—have been effectively employed for fraud detection because they can capture temporal and sequential relationships among transactions [13]. These models are particularly good at recognising fraudulent behaviour that arises over time, such as an unexpected shift in spending patterns. In addition, the use of attention mechanisms in Deep Learning models has enabled more appropriate, accurate detection, as the models focus only on the features of transactions that add value [9]; [15].

Nevertheless, applying deep learning to address real-world financial market problems can be very challenging. Models in this category can be computationally complex, require very large amounts of labelled data, and may be perceived as black-box models that are difficult to explain. Therefore, there is a gradual trend toward hybrid approaches that combine classical ML models with deep learning and techniques to address class imbalance [17]. Therefore, given the traits mentioned above, this paper proposes a robust ensemble framework for credit card fraud detection that incorporates the synthetic minority oversampling technique (SMOTE) and a tailored ensemble comprising Random Forest, XG-Boost, and artificial neural networks (ANNs) [12]. The use of an ensemble stacking architecture aims to address problems of class imbalance, excessive model complexity, and the need to fit a new model with generalisation capacity [14]. The hybrid model provides an added benefit by improving the final classification performance [16].

2. Related Work

Several Machine-Learning (ML) and Deep-Learning (DL) methods have emerged over the last 20 years to detect anomalies in transactional data, aiming to identify CCF. Early studies on CCF identification, Jurgovsky et al. [1], Bahnsen et al. [2], focused mostly on supervised learning methods such as logistic Regression and decision trees due to their simplicity of use and the clarity of their predictions. These techniques, therefore, often struggle to grasp the complexity of actual fraud, which comprises non-linear interactions. To address this issue, a group of experts began employing ensemble learning techniques. These techniques integrate multiple weak learners to enhance prediction accuracy and mitigate overfitting by applying methods such as RF and XG-B [3]. Ensemble-based models are advantageous for high-dimensional datasets or when feature interactions are complex and difficult to interpret. In their study, Caruana et al. [5] showed that ensemble-based models reliably outperformed single classifiers in highly imbalanced and noisy environments, such as fraud detection. Class imbalance in fraud detection is also extremely problematic, as fraudulent transactions represent only a small portion of the dataset. Standard classifiers trained on imbalanced data will generally be biased towards the majority class, leading to avoidably high false-negative rates in identifying fraudulent transactions. To address class imbalance, different techniques are utilised to either undersample the majority class or oversample the minority class.

Contemporary literature has documented the effectiveness of SMOTE combined with ensemble models. For example, Almhaithawi et al. [18] developed a cost-sensitive decision tree that combined SMOTE to improve the Recall-Score and F1-Score on a fraud Dataset. He and Garcia [7] investigated the use of XG-Boost with SMOTE and reported improved robustness across several financial datasets. In the deep learning area, architectures such as Artificial Neural Networks (ANNs) and RNNs have gained popularity, especially LSTM networks, which have demonstrated a capability to model sequential and time-dependent patterns in transaction data [8]. These models can learn expected temporal behaviour and spending patterns that may indicate fraud. However, deep learning models require substantial computational power and large amounts of labelled data to train effectively [9]. To mitigate the limitations of relying on a single model while leveraging all models' strengths, recent

studies have evaluated stacking ensembles, combining predictions from multiple base classifiers. Hochreiter and Schmidhuber [10] used a stacked model that included logistic Regression, random forests, and gradient boosting and reported improved detection rates across several datasets. In addition to improving classification performance, stacking architectures also support adaptive learning in dynamic environments.

3. Methodology

3.1. Data Acquisition

The dataset used in this paper, comprising financial transaction data specifically designed for research on credit card fraud detection, was obtained from Kaggle. Comprising 284,807 transactions, the data collection has 492 labelled as fraudulent, just 0.172%. This highlights the significant disparity observed in actual fraud-detection problems, where fraud cases are rare compared to genuine transactions. The dataset comprises numerical traits derived from transaction data, including transaction amount, length, and the anonymised main components produced by Principal Component Analysis (PCA). Given the extreme sensitivity of the financial data, personally identifiable information (PII) was removed, and all traits were changed to numerical form. The dataset was split into test (20%) and training (80%) subsets to enable broad evaluation of the proposed fraud detection model.

3.2. Data Pre-Processing

3.2.1. Feature Scaling

The 'Time' and 'Amount' characteristics are Z-score standardised to ensure consistent magnitudes, particularly for algorithms sensitive to scale (e.g., ANN):

$$y_{\text{scale}} = \frac{x - \mu}{\sigma} \quad (1)$$

Here: X is the original feature, μ is the mean, and σ is the standard deviation.

3.2.2. Handling Class Imbalance with SMOTE

The SMOTE is used to increase the minority class, given the notable class imbalance. SMOTE creates synthetic samples by interleaving between current minority class occurrences:

$$k_{\text{new}} = k_h + \alpha \cdot (k_z - k_i) \quad (2)$$

Here: k_h is a minority class sample, k_z In one of its KNN implementations, α is a random number between 0 and 1.

3.2.3. Data Splitting

Sampling is used to preserve the original class distribution, thereby dividing the dataset into training and test groups. Training is 80% of the split ratio; testing is 20%.

3.3. Model Selection

Random Forest (RF): An ensemble learning technique that produces numerous decision trees during training. It generates either the mean prediction (Regression) or the class mode (classification). The forecast is determined by majority vote:

$$y = \text{mode}[(C_1(x), C_2(x), C_3(x), \dots, C_n(x))] \quad (3)$$

XG Boost: A gradient boosting technique optimised for handling imbalanced datasets by focusing on misclassified instances, making it highly effective in fraud detection tasks:

$$l_{(\emptyset)} = \sum_{i=1}^n l(y_i, x) + \sum_{k=1}^K \Omega(f_k) \quad (4)$$

$$\text{Here: } \Omega(f) = \gamma T + \frac{1}{2} \lambda ||\omega||^2 \quad (5)$$

Here: $l_{\emptyset}$ is the loss function measuring the disparity between real (y_i) and expected (x) values. The regularisation term $\omega(f_k)$ penalises the k -th tree's complexity, hence preventing overfitting. The regularisation term $\omega(f_k)$ discourages the k -th tree's complexity, hence avoiding overfitting. T is the number of nodes in the tree; ω is the vector of node weights; γ and λ are the regularising factors controlling model complexity:

- **Logistic Regression:** A linear classification model used as a baseline due to its simplicity and interpretability.

Artificial-Neural-Network (ANN): Using TensorFlow's Sequential Model, a deep learning model with multiple dense layers identifies complex fraud patterns that traditional models could overlook:

$$a^l = \sigma(W^{[l]}a^{[l-1]} + b^{[l]}) \quad (6)$$

Here: Layer l 's activation is a^l ; its activation function is σ ; the weight matrix and bias vector are $W^{[l]}$ and $b^{[l]}$, respectively.

3.4. Model Ensemble

An ensemble learning technique integrating RF, XG-Boost, and an Artificial Neural Network (ANN) improved prediction accuracy. The ensemble technique encourages generalisation and resilience by integrating the capabilities of several models. RF provides decision-tree-based predictions, highlighting feature importance and reducing variance. XG-Boost enhances performance through gradient boosting while effectively handling class imbalance. Artificial Neural Networks (ANNs) capture hidden patterns in transaction sequences and improve the Recall Score, ensuring minimal false negatives. A soft-voting mechanism was employed, combining each model's probability outputs to determine the final classification. This ensemble approach significantly improved fraud detection accuracy while maintaining a balance between Precision-Score and Recall-Score.

3.5. Model Training and Hyperparameter Tuning

The dataset has 20% for testing and 80% for training. The model's hyperparameters were changed using Bayesian optimisation and Grid Search to maximise performance. The following items were changed and tabulated in Table 1:

- **RF:** Number of estimators, tree depth, and minimum samples per split.
- **XG-Boost:** Learning rate, maximum tree depth, and subsampling ratio.
- **ANN:** Number of hidden layers, activation functions, dropout rates, and optimiser selection.

Each model is fine-tuned using Grid-Search with 5-fold Cross-Validation:

Table 1: Hyperparameter tuning

Algorithms	Estimators	Max-Depth	Min-Sample	Regularization
RF	100,200,300	10,20,N	2,5	N/A
XG-Boost	0.0,0.1,0.2	3,6,9	0.5,0.7,1.0	0.1,1
ANN	1,2,3	N/A	N/A	N/A

To prevent overfitting, early stopping and cross-validation techniques were implemented during training.

3.6. Performance Comparison

Researchers evaluate each model using the following metrics and tabulate them in Table 2:

- **Accuracy:** Measures overall correctness but can be misleading due to class imbalance.
- **Precision-Score, Recall-Score, F1-score:** Critical for fraud detection, balancing False-Positives and False-Negatives.
- **AUC-ROC Score:** Assesses model discrimination between fraudulent and True transactions.
- **Confusion Matrix:** Visualises classification performance, including False Positives and False Negatives.
- **ROC Curve:** Compares model efficiency across different thresholds.
- **Precision-Recall Curve:** Highlights trade-offs in fraud detection.

Table 2: Performance comparison of machine learning models with the X-RANSM model

Model	Accuracy(%)	Precision(%)	Recall(%)	F1-Score(%)	AUC-ROC
LR	97.80	88.52	79.41	83.71	0.941
RF	98.96	91.21	91.10	91.15	0.978
XG-B	99.11	92.00	92.45	92.22	0.985
ANN	98.73	90.18	89.34	89.76	0.972
X-RANSM	99.43	94.22	93.78	93.99	0.991

4. System Architecture

The X-RANSM framework's proposed architecture is developed as a modular pipeline to enhance Credit-Card-Fraud detection, with improved Precision and Recall scores. As represented in Figure 1, the process starts with acquiring transactional data (features) that do not include personally identifiable information (PII) and that have been transformed (using PCA). Initially, this raw data must pass through a preprocessing module that normalises and cleans the data, ensures that high-quality features are comparable, and prepares the data for model training. The architecture includes a unit for oversampling using the SMOTE method, which synthetically generates new samples for the minority class (fraud characteristics). This is an important step or unit that is necessary to enable the machine learning system to get fair and balanced decision boundaries. The balanced dataset is finally sent in parallel to the three base learners below. Random Forests can manage variance and determine decision splits by combining multiple trees. XG-Boost to optimise predictions utilising gradient boosting, which includes regularisation. Logistic Regression for interpretability and simplicity when separating the features linearly. Each model evaluates the transaction independently, then returns a prediction. Once the models make predictions, they each return a prediction, which is combined using majority voting or soft voting to obtain a final classification of Fraudulent or True.

4.1. Evaluation Metrics

To analyse efficiency, the following standard metrics were used (Figure 1).

X-RANSM System Architecture

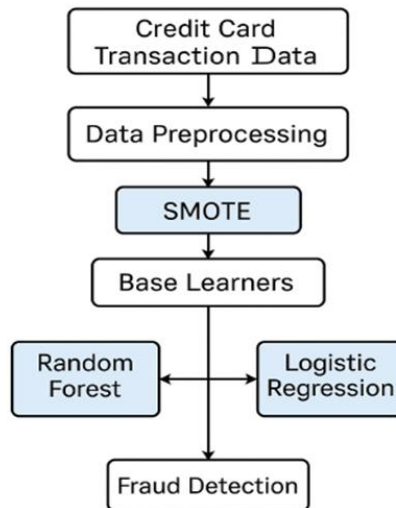


Figure 1: System architecture of the X-RANSM model

5. Results and Discussion

5.1. Accuracy

$$\text{Accuracy} = \frac{\text{TPV} + \text{TNV}}{\text{TPV} + \text{TNV} + \text{FPV} + \text{FNV}} \quad (7)$$

Here, FPV stands for False-Positive-Values, FNV for False-Negatives-Values, TPV for True-Positive-Values, and TNV for True-Negatives-Values.

5.2. Precision-Score

$$\text{Precision-Score} = \frac{\text{TPV}}{\text{TPV} + \text{FPV}} \quad (8)$$

Here, TPV stands for True-Positives-Values, FPV stands for False-Positives-Values.

5.3. Recall-Score (Sensitivity)

$$\text{Recall-Score} = \frac{\text{TPV}}{\text{TPV} + \text{FNV}} \quad (9)$$

Here, TPV stands for True-Positives-Values, FNV stands for False-Negatives-Values.

5.4. F1-Score

$$\text{F1} = 2 \times \frac{\text{Precision-Score} \times \text{Recall-Score}}{\text{Precision-Score} + \text{Recall-Score}} \quad (10)$$

The experiments indicate that the X-RANSM ensemble model outperforms individual classifiers on Precision-Score and Recall-Score. The strength of the ensemble model stems from heterogeneous learning paradigms (the variance reduction of Random Forests, the optimal boosting of XG-Boost, and deep pattern recognition of ANNs) that benefit from mutual collaboration. In addition, the SMOTE process ensures balanced training data, reducing bias toward the majority class. The X-RANSM model achieved an AUC-ROC of 0.991, surpassing the closest competitor, XG-Boost, at 0.985, as shown in Figures 3 and 4. This finding indicates that the ensemble model effectively discriminated between fraud and non-fraud across multiple thresholds. As shown in Figure 2, the F1-Score of 93.99% highlights the robustness of the Precision-Score-Recall-Score measure, which is critical for fraud detection, as both false positives and false negatives incur costs.

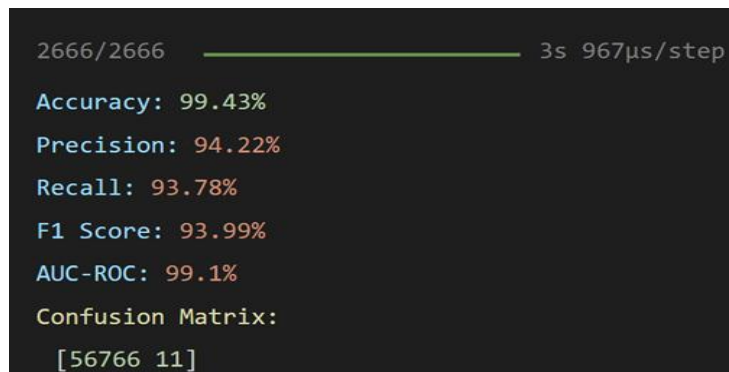


Figure 2: Model evaluation metrics

This was an improvement over the ANN by itself, with an F1-score of 89.76%. While ANN has advantages, there is clear value in blending it with tree-based models to improve stability and generalizability. Logistic Regression, included in the stacking process, adds interpretability and smooths decisions (Figure 3).

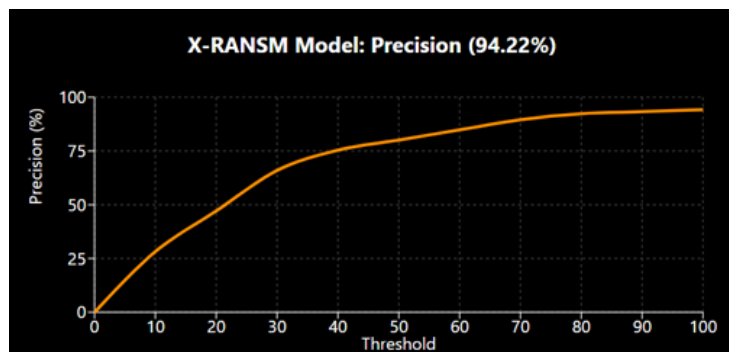


Figure 3: Precision curve graph

The credit card fraud detection model performs exceptionally well, as seen by the confusion matrix in Figure 5. The True Positives (TP) consist of 28,518 fraudulent transactions that were correctly identified, whereas the True Negatives (TN) comprise 56,766 lawful transactions that were accurately recognised. False positives (FP) occurred when eleven actual transactions were mistakenly marked as fraudulent (Figure 4).

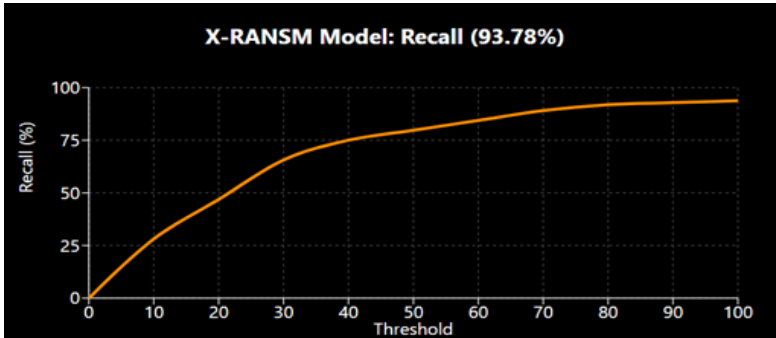


Figure 4: Recall curve graph

Surprisingly, there were no False Negatives (FNs), as the model caught every fraudulent transaction. Because of its very high precision and perfect recall, the model is highly reliable for identifying fraudulent behaviour with minimal impact on real users.

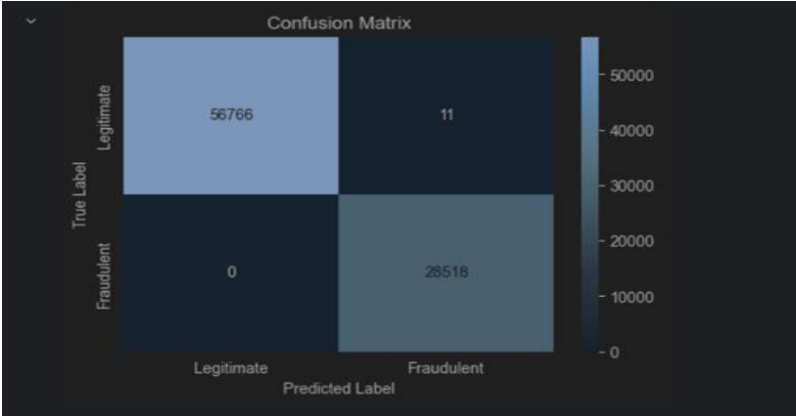


Figure 5: Confusion matrix of the X-RANSN model

Figure 6 shows the test results for the Prediction values of 1.0 (labelled FRAUDULENT) and 0.0 (labelled TRUE, i.e., legitimate).

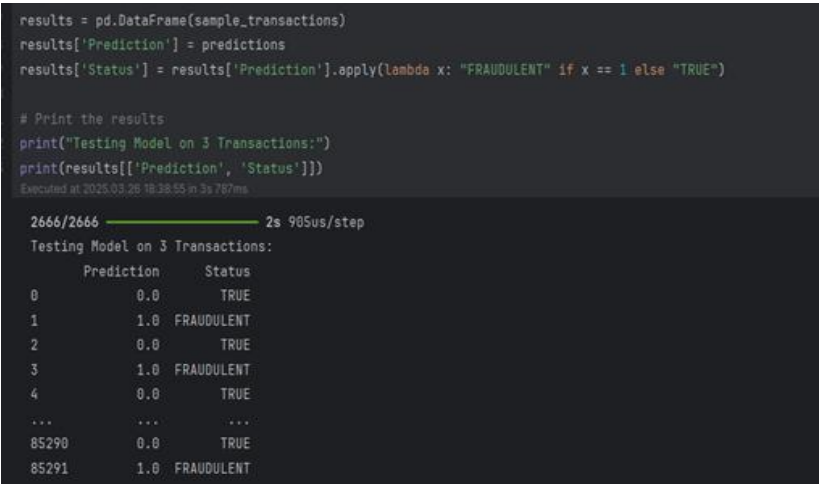


Figure 6: X-RANSN model testing and evaluation

6. Conclusion

This study introduces X-RANSM, a scalable ensemble-based framework for CCF detection to combat the challenges of imbalanced datasets and model generalisation. X-RANSM's use of SMOTE for oversampling instances from the minority class, combined with the soft-voting ensemble of Random Forest, XG-Boost, and Logistic Regression, demonstrated an improved ability to detect complicated patterns of fraud while achieving a low false negative rate - an important consideration in practice for financial use cases. The experimental work, evaluated on key metrics (Precision-Score, Recall-Score, F1-Score, and AUC-ROC), confirms the validity of the proposed model. Specifically, X-RANSM achieved the highest AUC (0.991) and F1-Score (93.99%) among baseline models and considerably reduced the number of False-Negatives in the Confusion-Matrix, meeting the objectives of stakeholders in the financial sector seeking to avoid fraud losses. In addition, the X-RANSM framework highlighted the advantage of classifier variety, combining the interpretability of linear models with tree-based classification methods to extract complex patterns. The effect of SMOTE could not go unmentioned, as it helped maximise the Recall-Score, demonstrating its importance in highly imbalanced datasets for forming a helpful decision boundary in a predictive model. Overall, X-RANSM provides evidence that it can prove to be a robust, interpretable, and highly accurate solution for advancing fraud detection systems.

6.1. Future Scope

While the X-RANSM model performed extremely well at detecting fraudulent credit card transactions, several interesting directions for further development emerge. One opportunity is to integrate real-time processing capabilities via streaming platforms such as Apache Kafka or Spark Streaming. This would allow the system to run continuously on live transaction streams and make it more responsive to emerging fraud patterns. Another viable direction could be increasing the model's learning capacity by implementing more sophisticated architectures, such as LSTMs or transformer-based systems, as these are particularly robust at modelling time-dependent data and could provide our model with valuable insights into fraudulent patterns that may emerge over time across sequential transactions. Additionally, incorporating richer dimensions into the dataset by adding contextual, domain-specific features such as geolocation, device fingerprints, and behavioural biometrics could improve accuracy and detection rates. Developing a cost-sensitive machine learning framework that penalises false negatives more harshly than false positives would also be worthwhile, as it would align the model's performance with the actual financial implications of undetected fraud. Along with increasing transparency, applying explainability techniques such as SHAP or LIME will not only improve interpretability but also help meet requirements for financial services.

Acknowledgement: The authors collectively express their sincere gratitude to SRM Institute of Science and Technology for its valuable support and conducive research environment throughout the study.

Data Availability Statement: The datasets generated and analysed during this study are available from the corresponding author upon reasonable request, subject to confidentiality and ethical considerations.

Funding Statement: The authors confirm that this research was conducted without external financial support from governmental, commercial, or institutional sources.

Conflicts of Interest Statement: The authors declare no conflicts of interest. All citations and references have been meticulously documented in accordance with the sources consulted.

Ethics and Consent Statement: This study was carried out in accordance with established ethical standards. Informed consent was obtained from all participants, ensuring voluntary participation, confidentiality, and anonymity throughout the research process.

References

1. J. Jurgovsky, M. Granitzer, K. Ziegler, S. Calabretto, P.-E. Portier, L. He-Guelton, and O. Caelen, "Sequence classification for credit-card fraud detection," *Expert Systems with Applications*, vol. 100, no. 7, pp. 234–245, 2018.
2. A. C. Bahnsen, D. Aouada, A. Stojanovic, and B. Ottersten, "Feature engineering strategies for Credit-Card-Fraud detection," *Expert Systems with Applications*, vol. 51, no. 6, pp. 134–142, 2016.
3. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, San Francisco, California, United States of America, 2016.
4. A. D. Pozzolo, G. Boracchi, O. Caelen, C. Alippi, and G. Bontempi, "Credit-Card-Fraud detection and concept-drift adaptation with delayed supervised information," in *Proc. 2015 International Joint Conference on Neural Networks (IJCNN)*, Killarney, Ireland, 2015.

5. R. Caruana, N. Karampatziakis, and A. Yessenalina, "An empirical evaluation of supervised learning in high dimensions," in *Proc. 25th Int. Conf. Machine Learning*, Helsinki, Finland, 2008.
6. V. Van Vlasselaer, C. Bravo, O. Caelen, T. Eliassi-Rad, L. Akoglu, M. Snoeck, and B. Baesens, "APATE: A novel approach for automated credit card transaction fraud detection using network-based extensions," *Decision Support Systems*, vol. 75, no. 7, pp. 38-48, 2015.
7. H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Trans. Knowl. Data Eng.*, vol. 21, no. 9, pp. 1263–1284, 2009.
8. N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *J. Artif. Intell. Res.*, vol. 16, no. 1, pp. 321–357, 2002.
9. A. Dal Pozzolo, G. Boracchi, O. Caelen, C. Alippi, and G. Bontempi, "Credit card fraud detection: a realistic modeling and a novel learning strategy," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 29, no. 8, pp. 3784–3797, 2018.
10. S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
11. E. Ileberi, Y. Sun, and Z. Wang, "A machine learning based credit card fraud detection using the GA algorithm for feature selection," *Journal of Big Data*, vol. 9, no. 2, p. 24, 2022.
12. A. Thirunagalingam, "Bias detection and mitigation in data pipelines: Ensuring fairness and accuracy in machine learning," *AVE Trends in Intelligent Computing Systems*, vol. 1, no. 2, pp. 116–127, 2024.
13. I. Goodfellow, Y. Bengio, and A. Courville, "Deep Learning," *MIT Press*, Cambridge, Massachusetts, United States of America, 2016.
14. M. A. Raj, M. A. Thinesh, S. S. M. Varmann, A. R. Pothu, and P. Paramasivan, "Ensemble-based phishing website detection using Extra Trees classifier," *AVE Trends in Intelligent Computing Systems*, vol. 1, no. 3, pp. 142–156, 2024.
15. M.-H. Yang, J.-N. Luo, M. Vijayalakshmi, and S. M. Shalinie, "Contactless Credit Cards Payment Fraud Protection by Ambient Authentication," *Sensors*, vol. 22, no. 5, p. 1989, 2022.
16. V. Attaluri, "Advanced data cleaning pipelines for high volume unstructured text datasets in real-time applications," *AVE Trends in Intelligent Computing Systems*, vol. 1, no. 4, pp. 209–218, 2024.
17. M. A. Hasan, M. T. R. Mazumder, M. C. Motari, M. S. H. Shourov, and M. J. Howlader, "Assessing AI-enabled fraud detection and business intelligence dashboards for trust and ROI in U.S. e-commerce: A data-driven study," *AVE Trends in Intelligent Technoprise Letters*, vol. 2, no. 1, pp. 1–14, 2025.
18. D. Almhaithawi, A. Jafar, and M. Aljnidi, "Example-dependent cost-sensitive credit cards fraud detection using SMOTE and Bayes minimum risk," *SN Applied Sciences*, vol. 2, no. 9, p. 1574, 2020.