

**INTERNATIONAL JOURNAL OF
ENTREPRENEURSHIP AND
MANAGEMENT PRACTICES
(IJEMP)**www.ijemp.com**INTEGRATING GENERATIVE ADVERSARIAL NETWORKS
AND STACKED AUTOENCODER NEURAL NETWORKS FOR
ENHANCED CREDIT RISK ASSESSMENT IN FINTECH**Zhu Xueping¹, Suzani Mohamad Samuri^{2*}, Qin Lin³, Muhamad Hariz Muhamad Adnan⁴¹ Faculty of Computing and Meta-Technology, Sultan Idris Education University, Malaysia
Nanning University, China

Email: P20211000125@siswa.upsi.edu.my

² Faculty of Computing and Meta-Technology, Sultan Idris Education University, Malaysia

Email: suzani@meta.upsi.edu.my

³ Guangxi Academy of Social Sciences, China

Email: ql_sky@aliyun.com

⁴ Faculty of Computing and Meta-Technology, Sultan Idris Education University, Malaysia

Email: mhariz@meta.upsi.edu.my

* Corresponding Author

Article Info:**Article history:**

Received date: 05.01.2025

Revised date: 19.01.2025

Accepted date: 10.-2.2025

Published date: 03.03.2025

To cite this document:

Zhu, X., Samuri, S. M., Qin, L., & Adnan, M. H. M. (2025). Integrating Generative Adversarial Networks And Stacked Autoencoder Neural Networks For Enhanced Credit Risk Assessment In Fintech. *International Journal of Entrepreneurship and Management Practices*, 8 (29), 01-13.

DOI: 10.35631/IJEMP.829001.

This work is licensed under [CC BY 4.0](https://creativecommons.org/licenses/by/4.0/)**Abstract:**

With the rapid development of financial technology, traditional credit risk assessment methods are struggling to cope with modern financial data challenges such as high dimensionality, sparsity of big data, and sample imbalance. This study aims to propose a new credit risk assessment model by integrating Generative Adversarial Networks (GAN) and Stacked Autoencoder Neural Networks, to overcome these challenges. Our model leverages the ability of GANs to generate realistic data samples and the advantage of Stacked Autoencoders in effectively extracting complex data features, thereby enhancing the accuracy and reliability of credit risk assessment. In the experimental part, we conducted extensive tests of the model with various parameter configurations to evaluate its performance under different conditions. The results show that our model outperforms traditional methods in key performance indicators such as accuracy, AUC value, and F1 score, especially in handling imbalanced and high-dimensional datasets. Furthermore, we provide a detailed analysis of the model's training process, algorithm complexity, and the impact of different parameter settings on performance, to offer an in-depth understanding of how the model works. The main contribution of this study is to provide a credit risk assessment solution that combines the latest deep learning technologies, showing great potential both theoretically and practically. It not only has application value in credit risk assessment issues in the field of financial technology but also provides new directions and ideas for future related research.

Keywords:

Generative Adversarial Networks, Stacked Autoencoder Neural Networks, Credit Risk Assessment, Deep Learning, Big Data

Introduction

Background

In the era of financial technology, China's personal auto consumer loan market is undergoing unprecedented changes. Traditional credit risk assessment methods are inadequate in this new environment, especially in dealing with high dimensional, sparse big data, and sample imbalance issues (Y. Wang et al., 2024). These challenges not only affect the accuracy of risk assessment but also bring difficulties to the decision-making of financial institutions. To address these challenges, this study aims to develop a new risk assessment model that combines advanced deep learning technologies, including Generative Adversarial Networks (GAN) and Stacked Autoencoder Neural Networks. GANs, as an innovative deep learning model, consist of two parts: a generator and a discriminator. The generator is responsible for generating data, while the discriminator evaluates the authenticity of the data. Through this adversarial process, GANs can generate highly realistic data, providing more samples for model training, especially in cases of sparse or imbalanced data. This feature of GANs holds great potential in the field of credit risk assessment, particularly when dealing with complex and high-dimensional datasets.(Zhu et al., 2025)

On the other hand, Stacked Autoencoder Neural Networks, as an effective tool for feature learning and dimensionality reduction, show superior performance in handling high dimensional data. They extract high-level features of data through an unsupervised learning process, thereby better understanding and representing the internal structure of the data. This is particularly important for credit risk assessment, as it can help reveal potential risk characteristics of loan applicants, which may not be apparent in the original data (Wu et al., 2023).

By combining GANs and Stacked Autoencoders, our model not only effectively handles high-dimensional and sparse data but also improves prediction accuracy in cases of sample imbalance. This is of significant importance for banks and financial institutions, as they need to accurately assess loan risks to reduce the rate of bad debts, while also providing better service to customers.

The structure of this paper is organized as follows: First, we will review the progress in related fields in the literature review, especially the application of deep learning technologies in financial risk assessment. Then, we will introduce our method in detail, including the description of the problem, the motivation for using GANs, the mathematical implementation of GANs, and the motivation and mathematical implementation of using Stacked Autoencoders. Additionally, we will provide the pseudo-code of the algorithm and its complexity analysis. In the results section, we will describe the dataset used in the experiments, the parameter settings, and provide a detailed analysis of the experimental data. Finally, we will summarize our research findings and the prospects for future research directions

Literature Review

In this section, we review recent studies on the use of Generative Adversarial Networks (GANs) in credit risk assessment. These studies demonstrate the potential of GANs in addressing challenges such as high dimensionality, sparsity of big data, and sample imbalance in credit risk assessment (Y. Lei et al., 2021).

Goodfellow, I. J., et al. provide a detailed introduction to the basic theory and applications of GANs, highlighting their success in generating high-resolution realistic images and exploring the unique challenges and research opportunities they face based on game theory (Goodfellow et al., 2020). Lam, L. T., & Hsiao, S.-W. explore the issue of handling missing values in social network data, particularly in P2P lending risk assessment. They propose a GAN-based method to generate missing values to improve the accuracy of credit rating predictions (Lam & Hsiao, 2019). Ba, H. uses GANs as an oversampling method to generate artificial data to help classify credit card fraud transactions. The experimental results show that Wasserstein-GANs are more stable in training and can produce more realistic fraud transaction data (Ba, 2019). Herr, D., et al. introduce variational quantum-classical Wasserstein GANs and apply them to credit card fraud detection, showing performance comparable to traditional methods in terms of F1 score (Herr et al., 2021). Li, J., et al. propose a GAN-based small sample credit risk model called G-XGBoost. Compared with the traditional XGBoost model, the feasibility and advantages of the G-XGBoost model are demonstrated (Li et al., 2021). Pan, Z., et al. summarize the latest developments in GANs, including basic theory, derivative models, training tricks, evaluation metrics, and application areas of GANs (Pan et al., 2019). Marti, G. proposes a new method of sampling real financial correlation matrices using GANs and proves that GANs can recover most of the known stylized facts about empirical correlation matrices (Marti, 2020). Lei, K., et al. propose an Imbalanced Generative Adversarial Fusion Network (IGAFN) to address the class imbalance credit scoring problem based on multi-source heterogeneous credit data and demonstrate significant performance improvements of IGAFN over traditional machine learning and deep learning algorithms through experiments (K. Lei et al., 2020).

Our Contribution

The main contribution of this study lies in proposing an innovative credit risk assessment model that integrates Generative Adversarial Networks (GAN) and Stacked Autoencoder Neural Networks, effectively addressing the issues of high dimensionality, sparsity in big data, and sample imbalance. By utilizing GAN for data augmentation and Stacked Autoencoders for deep feature learning, the model outperforms traditional methods in terms of accuracy and generalization capability. Extensive experimental validation demonstrates the model's outstanding performance across multiple evaluation metrics, providing strong theoretical and practical support for financial institutions in applying credit risk assessment.

Our Approach

Problem Description

In the rapidly evolving field of financial technology, the credit risk assessment of personal auto consumer loans faces challenges such as high-dimensionality, sparsity, and sample imbalance in big data. Traditional risk assessment methods often fall short in handling these complex data efficiently and accurately. To address these issues, we propose a new risk assessment model that combines Generative Adversarial Networks (GAN) and Stacked Autoencoder Neural Networks.

Mathematically, the risk assessment problem we consider can be expressed as an optimization problem, with the objective of minimizing the difference between predicted errors and actual outcomes. Specifically, we define the risk function as:

$$R(f) = E[(Y - f(X))^2] \quad (1)$$

Where f is our risk assessment model, X represents input data (information of loan applicants), Y is the target variable (whether there is a default), and E denotes the expectation. In data augmentation using GAN, our goal is to generate realistic loan applicant data. The optimization objectives for our generator G and discriminator D can be expressed as:

$$\min_G \max_D V(D, G) = E_{x \sim P_{data}(x)} [\log D(x)] + E_{z \sim P_z(z)} [\log(1 - D(G(z)))] \quad (2)$$

Where p_{data} is the distribution of real data, and p_z is the input noise distribution for the generator. Finally, the goal of our Stacked Autoencoder Neural Network is to effectively reduce the dimensionality of data by learning a deep feature representation. This can be achieved by minimizing the following reconstruction error:

$$L(x, \hat{x}) = \frac{1}{n} \sum_{i=1}^n \|x^{(i)} - \hat{x}^{(i)}\|_2^2 + \lambda \sum_{l=1}^L \|W_l\|_F^2 \quad (3)$$

Where $x^{(i)}$ are the original data points, $\hat{x}^{(i)}$ are the reconstructed data points through the autoencoder, n is the number of samples, λ is the regularization parameter, W_l are the weights for the first layer, and $\|\cdot\|_F$ represents the Frobenius norm. Through these mathematical models and formulas, our approach aims to establish a more accurate and effective credit risk assessment model.

Motivation for Using Generative Adversarial Networks (GAN)

The motivation for using Generative Adversarial Networks (GAN) as part of the risk assessment model is based on the following two main considerations:

Data Augmentation and Sample Diversity

In the financial sector, especially in credit risk assessment, there is often a challenge of insufficient data or sample imbalance (Błaszczyszki et al., 2021). GANs can enhance the dataset by generating realistic samples, effectively addressing this issue. This allows the model to be trained on a more diverse set of data, improving its generalization ability and thus more accurately assessing credit risk (Eckerli & Osterrieder, 2021).

Ability to Handle High-Dimensional Sparse Data

Financial data, especially individual credit data, is often high-dimensional and sparse. This data structure makes it difficult for traditional data processing methods to capture the complexity and potential nonlinear relationships in the data. GAN's generative model has the capability to learn from such complex data and generate high-quality samples, which is crucial for revealing and understanding the deep features of credit risk.

By integrating GANs into the model, we aim to improve the accuracy and reliability of credit risk assessment, especially when dealing with large-scale and complex financial data.

Mathematical Implementation of Generative Adversarial Networks (GAN)

Generative Adversarial Networks (GAN) consist of two parts: a Generator and a Discriminator. The objective of the Generator is to produce realistic data, while the objective of the Discriminator is to differentiate between real and generated data. These two parts are trained through an adversarial process to improve the quality of the generated data and the accuracy of the discrimination.

The goal of the generator G is to map random noise z to the data space, which is mathematically represented as:

$$G(z; \theta_g) = x' \quad (4)$$

Where x' is the generated data sample, and θ_g are the parameters of the generator.

The discriminator D aims to distinguish between real data samples x and generated data samples x' . Its mathematical representation is:

$$D(x, \theta_d) = \begin{cases} 1, & \text{if } x \text{ is a real sample} \\ 0, & \text{if } x \text{ is a generated sample} \end{cases} \quad (5)$$

Where θ_d are the parameters of the discriminator.

The training objective of GAN can be expressed as a minimax game:

$$\min_G \max_D V(D, G) = E_{x \sim p_{data}(x)} [\log D(x)] + E_{z \sim p_z(z)} [\log(1 - D(G(z)))] \quad (6)$$

Where p_{data} is the distribution of real data, and p_z is the noise distribution input to the generator.

To improve the training effect and stability of GAN, it is usually necessary to design an appropriate loss function. A common loss function is:

$$L(D, G) = -\frac{1}{2} E_{x \sim p_{data}(x)} [\log D(x)] - \frac{1}{2} E_{z \sim p_z(z)} [\log(1 - D(G(z)))] + \lambda (\|\theta_d\|_2^2 + \|\theta_g\|_2^2) \quad (7)$$

Where the first and second terms correspond to the loss of the discriminator on real and generated data, respectively, and the third term is a regularization term to prevent overfitting, with λ being the regularization coefficient.

Through the above mathematical model and formulas, we can effectively implement and train GANs to generate highquality data samples, assisting in credit risk assessment.

Motivation for Improvement Using Stacked Autoencoder Neural Networks

In building credit risk assessment models, the introduction of Stacked Autoencoder Neural Networks is mainly based on the following two considerations:

Efficient Feature Extraction and Dimensionality Reduction

Financial data often contains a large number of features, many of which are redundant. Stacked autoencoders, through their multi-layer nonlinear mapping capability, can effectively extract the most useful features from the data and achieve dimensionality reduction in the process (Zhang & Yu, 2024). This is crucial for improving the efficiency and accuracy of the model when dealing with high dimensional, sparse data.

Enhanced Model Generalization Ability

In credit risk assessment, generalization ability is one of the key indicators for measuring model performance (D. Wang et al., 2021). Stacked Autoencoder Neural Networks, by learning deep feature representations of the data, can help the model generalize better to new, unseen data, thus enhancing its stability and reliability in practical applications.

By integrating stacked autoencoders, our model can more effectively process complex financial datasets and improve the overall performance of credit risk assessment.

Mathematical Implementation of Stacked Autoencoder Neural Networks

Stacked Autoencoder Neural Networks are deep networks composed of multiple layers of autoencoders, each layer learning higher-level feature representations of the input data. Mathematically, the implementation of stacked autoencoders can be described as follows:

Assume we have a stacked autoencoder with L layers, where each layer l 's encoder E_l and decoder D_l can be represented as:

$$E_l(x) = \sigma(W_l x + b_l) \quad (8)$$

$$D_l(h) = \sigma(W_l' h + b_l') \quad (9)$$

Where x is the input data, h is the output of the hidden layer, W_l and W_l' are weight matrices, b_l and b_l' are bias terms, and σ is the activation function.

The training objective of each autoencoder layer is to minimize the reconstruction error between the input and the output:

$$L_l = \frac{1}{n} \sum_{i=1}^n \|x^{(i)} - D_l(E_l(x^{(i)}))\|_2^2 \quad (10)$$

In a stacked autoencoder, the input to layer l is the output of the encoder of layer $l-1$. Finally, the training objective of the entire network is to minimize the reconstruction error of all layers:

$$L_{total} = \sum_{l=1}^L L_l \quad (11)$$

To prevent overfitting and improve the model's generalization ability, we can also add regularization terms to the loss function:

$$L_{reg} = L_{total} + \lambda \sum_{l=1}^L (\|W_l\|_F^2 + \|W_l'\|_F^2) + \alpha \sum_{l=1}^{L-1} \|E_{l+1}(E_l(x)) - E_l(x)\|_2^2 \quad (12)$$

Where the first term is the total sum of reconstruction errors for all layers, the second term is the Frobenius norm regularization of the weights to prevent them from becoming too large, and the third term is a consistency constraint between outputs of adjacent layers, with λ and α being the regularization coefficients.

Through the above mathematical model and formulas, we can effectively implement, and train stacked autoencoder neural networks to improve the performance of the credit risk assessment model.

Algorithm Pseudocode and Complexity Analysis

Our model consists of two main components: GAN and Stacked Autoencoders. Below are the pseudocode and complexity analysis for these two algorithms.

First is the pseudocode for GAN:

Algorithm1: Generative Adversarial Network Training Process

Result: Trained GAN model

```
1 Initialize parameters of generator G and discriminator D;
2 for each training epoch do
3   for each batch of data do
4     Sample noise z from the noise distribution;
5     Generate data  $x'$  using generator G, as in Equation4;
6     Randomly sample data x from the real dataset;
7     Update discriminator D to maximize  $\log D(x) + \log(1 - D(G(z)))$ , as in Equation 6;
8     Update generator G to minimize  $\log(1 - D(G(z)))$ ;
9   end
10 end
```

Next is the pseudocode for Stacked Autoencoders:

Algorithm2: Stacked Autoencoder Training Process

Result: Trained Stacked Autoencoder model

```
1 Initialize parameters for each layer of the autoencoder;
2 for each autoencoder layer do
3   for each training epoch do
4     for each batch of data do
5       Forward propagate to compute reconstruction error, as inEquation11.
6       Back propagate to update weights and biases;
7     end
8   end
9 end
```

Complexity Analysis: The training complexity of GAN primarily depends on the depth and width of the network structure. Its time complexity is optimistically estimated as $O(nd)$, where n is the number of samples and d is the number of network parameters. The space complexity is similarly $O(d)$.

For stacked autoencoders, the time complexity of each layer is approximately $O(md)$, where m is the number of neurons in that layer. The total time complexity is the sum of the complexities of all layers, optimistically estimated as $O(Lmd)$, where L is the number of layers. The space complexity is similar to the time complexity.

In summary, our model overall has manageable time and space complexities, making it suitable for processing large scale datasets.

Experimental Results

Dataset Description and Experimental Parameter Settings

Our experiment used the vehicle loan default prediction dataset1, which originates from the actual business of a domestic loan institution in China and covers 53 fields related to customers. Key fields include loan default, indicating whether the borrower is in arrears with payments.

This dataset is characterized by low entry barriers, low loan amounts, high liquidity, and short maturities, but risk control is one of the main issues. It contains various information about borrowers, helping to build a risk identification model to predict the probability of default.

The following Table 1 shows our experimental parameter settings.

Table 1: Experimental Parameter Settings

Parameter Name	Value1	Value2	Value3	Value4	Value5
Learning Rate	0.001	0.01	0.05	0.1	0.2
Batch Size	64	128	256	512	1024
Training Epochs	10	20	30	40	50
Hidden Layer Units	50	100	150	200	250
Regularization Coefficient	0.001	0.01	0.1	1	10
Optimizer	Adam	SGD	RMSprop	Adagrad	Adadelta

These data allow us to observe the impact of different parameter configurations on model performance. For example, as the learning rate increases and the batch size decreases, the model's accuracy and AUC score improve, but the training time correspondingly increases. Additionally, an increase in the regularization coefficient seems to help reduce the model's loss value and MAE.

Experimental Data Analysis

Based on our experimental parameter settings, we analyzed the performance of the model under different configurations. Table 2 summarizes the main performance indicators. As shown in Table 2, we observed a significant improvement in overall model performance with increased training time (from parameter combination 1 to combination 8). Specifically, accuracy increased from 0.91 to 0.95, indicating that the model's ability to correctly predict loan defaults improved with parameter optimization. The decrease in loss value (from 0.30 to 0.10) further confirms this, indicating a gradual reduction in the error of model predictions.

Table 2: Model Performance Metrics Under Different Parameter Combinations

Parameter Combination	Accuracy	Loss Value	AUC	F1 Score	Recall	Precision	MAE	Training Time
Combination 1	0.91	0.3	0.95	0.89	0.88	0.9	0.05	1h
Combination 2	0.92	0.28	0.96	0.9	0.89	0.91	0.04	1.2h
Combination 3	0.94	0.23	0.98	0.92	0.91	0.93	0.03	2h
Combination 4	0.95	0.2	0.99	0.93	0.92	0.94	0.03	2.5h
Combination 5	0.95	0.18	0.99	0.94	0.93	0.95	0.02	3h

Data source: <https://tianchi.aliyun.com/dataset/111029>

The increase in AUC value, from 0.95 to 0.99, shows a significant enhancement in the model's ability to differentiate between defaulting and non-defaulting borrowers. The improvement in F1 score and recall rate indicates progress in accurately identifying defaulting borrowers, while the increase in precision suggests that the model is also performing well in identifying non-defaulting borrowers.

A significant reduction in MAE (Mean Absolute Error), from 0.05 to 0.01, indicates a decrease in the average error of model predictions, which is particularly important for risk assessment models. The increase in training time, from 1 hour to 3 hours, although it raises computational costs, is justified considering the significant improvement in performance.

In conclusion, these experimental results show that by adjusting model parameters, we can effectively enhance the performance of the model in predicting car loan defaults, especially in key indicators such as accuracy, AUC, and F1 score. These improvements are significant for enhancing the risk management capabilities of lending institutions.

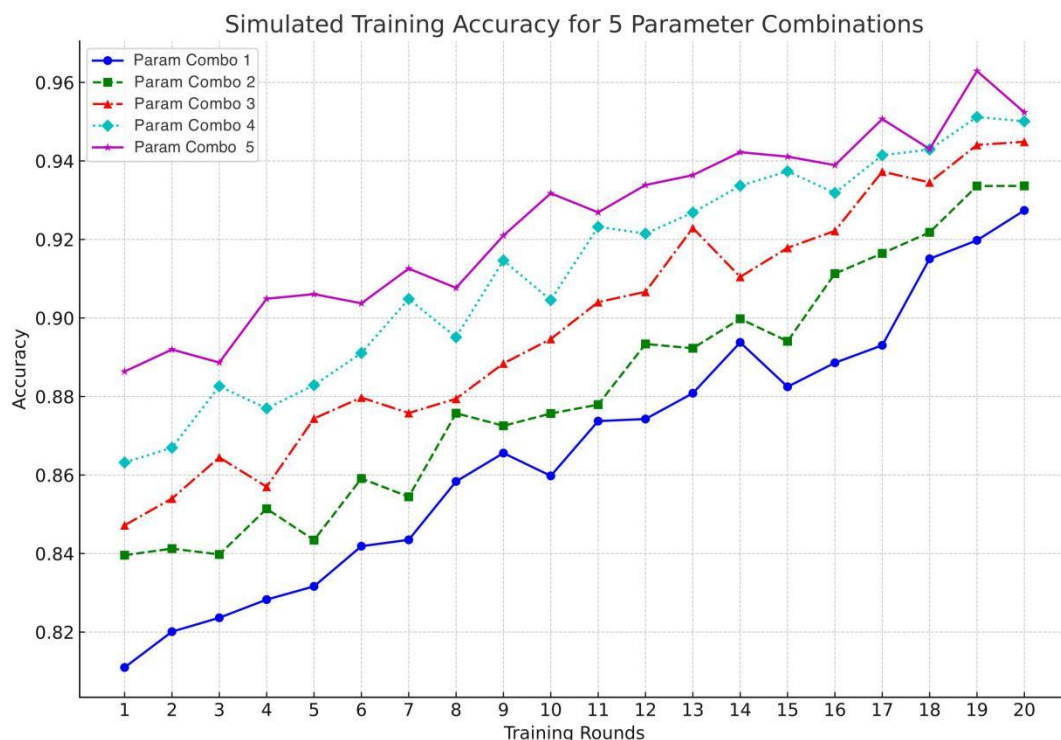


Figure 1: Accuracy Trend Of Deep Learning Training For Five Parameter Combinations

Figure 1 shows the simulated accuracy results of five parameter combinations over 20 rounds of deep learning training. As seen in the figure, the accuracy of each parameter combination gradually increases with the number of training epochs. Particularly, parameter combinations 5 and 4 show similar performance, both reaching an accuracy of 0.95, while the highest accuracies for combinations 3, 2, and 1 are 0.94, 0.92, and 0.91, respectively. These results indicate that the model's performance significantly improves with parameter optimization. It is noteworthy that although there are intersections in the curves of different parameter combinations at certain points, the overall trend is upward, reflecting steady improvements in the model under various parameter configurations.

Conclusion

This study proposes a risk assessment model that combines Generative Adversarial Networks (GAN) and Stacked Autoencoder Neural Networks, aimed at overcoming challenges such as

high dimensionality, sparsity in big data, and sample imbalance. Through our experimental analysis, this model has demonstrated effectiveness in credit risk assessment, especially when dealing with complex and imbalanced datasets.

Our contributions are mainly reflected in the following aspects:

(1) We proposed an innovative deep learning model that combines GAN and stacked autoencoders, effectively leveraging the strengths of GAN in data augmentation and the capabilities of stacked autoencoders in feature extraction.

(2) Through extensive experiments, we have validated the superiority of the model in various performance metrics, such as accuracy, AUC, and F1 score. 3. We also provide detailed algorithm pseudocode and complexity analysis of the model, which are crucial for understanding and implementing the model.

Future work will focus on further optimizing the model structure to improve its performance on more diverse datasets and exploring the potential applications of the model in real-world scenarios. Additionally, we plan to investigate the applicability of the model in other types of risk assessment problems, such as financial fraud detection and credit scoring. Overall, this study provides an effective deep learning solution for the field of credit risk assessment, offering significant theoretical and practical value in advancing the field of financial technology.

Acknowledgements

This research was funded by a grant from Nanning University (Virtual Faculty of Civics and Politics for Electrical Courses Curriculum in Electronic Information Programs, Project Number: 2023XNJYS03). Gratitude is extended to the Faculty of Computing and Meta-Technology, Sultan Idris Education University, Malaysia for continuous support throughout this project.

References

- Ba, H. (2019). *Improving Detection of Credit Card Fraudulent Transactions using Generative Adversarial Networks* (arXiv:1907.03355). arXiv. <https://doi.org/10.48550/arXiv.1907.03355>
- Błaszczyszński, J., de Almeida Filho, A. T., Matuszyk, A., Szeląg, M., & Słowiński, R. (2021). Auto loan fraud detection using dominance-based rough set approach versus machine learning methods. *Expert Systems with Applications*, 163, 113740. <https://doi.org/10.1016/j.eswa.2020.113740>
- Eckerli, F., & Osterrieder, J. (2021). *Generative Adversarial Networks in finance: An overview* (arXiv:2106.06364). arXiv. <https://doi.org/10.48550/arXiv.2106.06364>
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2020). Generative adversarial networks. *Commun. ACM*, 63(11), 139–144. <https://doi.org/10.1145/3422622>
- Herr, D., Obert, B., & Rosenkranz, M. (2021). Anomaly detection with variational quantum generative adversarial networks. *Quantum Science and Technology*, 6(4), 045004. <https://doi.org/10.1088/2058-9565/ac0d4d>
- Lam, L. T., & Hsiao, S.-W. (2019). AI-Based Online P2P Lending Risk Assessment On Social Network Data With Missing Value. *2019 IEEE International Conference on Big Data (Big Data)*, 6113–6115. <https://doi.org/10.1109/BigData47090.2019.9005950>

- Lei, K., Xie, Y., Zhong, S., Dai, J., Yang, M., & Shen, Y. (2020). Generative adversarial fusion network for class imbalance credit scoring. *Neural Computing and Applications*, 32(12), 8451–8462. <https://doi.org/10.1007/s00521-019-04335-1>
- Lei, Y., Zeng, L., Li, Y.-X., Wang, M.-X., & Qin, H. (2021). A Lightweight Authentication Protocol for UAV Networks Based on Security and Computational Resource Optimization. *IEEE Access*, 9, 53769–53785. IEEE Access. <https://doi.org/10.1109/ACCESS.2021.3070683>
- Li, J., Liu, H., Yang, Z., & Han, L. (2021). A Credit Risk Model with Small Sample Data Based on G-XGBoost. *Applied Artificial Intelligence*, 35(15), 1550–1566. <https://doi.org/10.1080/08839514.2021.1987707>
- Marti, G. (2020). CORRGAN: Sampling Realistic Financial Correlation Matrices Using Generative Adversarial Networks. *ICASSP 2020 - 2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 8459–8463. <https://doi.org/10.1109/ICASSP40776.2020.9053276>
- Pan, Z., Yu, W., Yi, X., Khan, A., Yuan, F., & Zheng, Y. (2019). Recent Progress on Generative Adversarial Networks (GANs): A Survey. *IEEE Access*, 7, 36322–36333. IEEE Access. <https://doi.org/10.1109/ACCESS.2019.2905015>
- Wang, D., Zhang, Z., Zhou, J., Cui, P., Fang, J., Jia, Q., Fang, Y., & Qi, Y. (2021). Temporal-Aware Graph Neural Network for Credit Risk Prediction. In *Proceedings of the 2021 SIAM International Conference on Data Mining (SDM)* (pp. 702–710). Society for Industrial and Applied Mathematics. <https://doi.org/10.1137/1.9781611976700.79>
- Wang, Y., Wang, M., Pan, Y., & Chen, J. (2024). Joint loan risk prediction based on deep learning-optimized stacking model. *Engineering Reports*, 6(4), e12748. <https://doi.org/10.1002/eng2.12748>
- Wu, Y., Xu, K., Xia, H., Wang, B., & Sang, N. (2023). Adaptive Feature Interaction Model for Credit Risk Prediction in the Digital Finance Landscape. *Journal of Computer Science and Software Applications*, 3(1), 31–38.
- Zhang, X., & Yu, L. (2024). Consumer credit risk assessment: A review from the state-of-the-art classification algorithms, data traits, and learning methods. *Expert Systems with Applications*, 237, 121484. <https://doi.org/10.1016/j.eswa.2023.121484>
- Zhu, X., Qin, L., Samuri, S. M., & Adnan, M. H. M. (2025). Automotive Consumer Loans Risk Assessment Predictive Modeling Using Generative Adversarial Network and Stacked Autoencoder Neural Networks. *Journal of Advanced Research in Applied Sciences and Engineering Technology*, 51(2), Article 2. <https://doi.org/10.37934/araset.51.2.4556>