

A Review of Machine Learning Applications for Credit Default Risk Prediction and Early Warning Systems

Bojun Chen

*School of Mathematics and Physics, Xi'an Jiaotong-Liverpool University, Suzhou, China
bojun.chen24@gmail.com*

Abstract. Credit risk management (CRM) is a fundamental pillar of financial systems, attracting attention from practitioners and researchers. Traditional credit assessment methods have limitations in today's complex, fast-changing financial environment. Advances in machine learning (ML) and behavioral data analytics offer new possibilities for improving CRM through better models and performance. This paper provides a systematic review of ML applications in credit default prediction and early warning systems, critically synthesizing recent literature. It discusses three major dimensions: the evolution of ensemble learning algorithms, the use and issues of behavioral data in feature engineering, and advances in model explainability (XAI). The paper shows that ensemble learning models have superior predictive power and argues that behavioral data complement traditional datasets for underbanked populations, such as "credit invisibles." It makes a strong case for XAI being essential for model transparency, combating bias, and meeting regulatory requirements. The review also addresses class imbalance, data privacy, and ethical issues, with mitigation strategies. The review offers theoretical guidance and practical implications for financial institutions to improve risk control and build reliable early warning systems, outlining directions for future research.

Keywords: Machine learning, Credit default prediction, Ensemble learning, Explainable AI, Behavioral data

1. Introduction

Credit risk—the possibility of borrower default—represents the most significant risk facing financial institutions worldwide. Effective credit risk management is essential for reducing potential losses, maintaining capital adequacy, and ensuring financial system stability. Traditional credit assessment methods, which rely heavily on expert judgment, manual inspection, and classical statistical models such as logistic regression and Altman's Z-score model, often suffer from limited adaptability and weak predictive performance in today's complex, dynamic, and data-rich financial environment. These conventional approaches depend on structured financial statement data and historical credit records, making them less effective for assessing borrowers with limited credit histories or in rapidly changing market conditions.

In comparison, machine learning (ML) algorithms have emerged as a transformative approach in credit default prediction because of their superior capabilities in processing large-scale, high-

dimensional data and identifying complex nonlinear patterns. By integrating alternative data sources, such as behavioral data, digital footprints, and social network information, these advanced models can effectively alleviate information asymmetry between lenders and borrowers, achieve higher prediction accuracy, and expand financial inclusion to previously underserved populations. Moreover, the integration of explainable artificial intelligence (XAI) techniques addresses the critical need for model transparency, regulatory compliance, and stakeholder trust in automated credit decisions.

This paper provides a systematic literature review of ML applications in credit default prediction and early warning systems. Specifically, it examines the evolution of ensemble learning algorithms, the use of behavioral and alternative data in feature engineering, and advances in model explainability through XAI methods. The review also addresses practical challenges, including class imbalance, data privacy, and algorithmic fairness. The theoretical contribution of this study lies in synthesizing recent advances across these three dimensions and identifying promising directions for future research. For practitioners, this study offers actionable guidance for constructing efficient, reliable, and regulatory-compliant early warning systems that can support proactive risk management strategies. The main contributions of this review are threefold: (1) identifying hybrid resampling-ensemble strategies as the most effective approach for imbalanced credit data; (2) analyzing behavioral and alternative data's role in expanding financial inclusion for underbanked populations; and (3) proposing an integrated REWS framework that combines ML prediction, XAI explainability, and proactive risk management into a practical end-to-end solution.

2. Theoretical background and traditional approaches

2.1. Credit risk and default: definitions and concepts

Credit risk refers to the uncertainty that a debtor will fail to fulfill contractual obligations, such as repaying principal and interest, thereby exposing the creditor to potential financial losses. Credit risk can be categorized into corporate credit risk (e.g., default risk on loans and bonds), personal credit risk (e.g., credit card and mortgage defaults), and sovereign credit risk. A credit default is generally defined as failure to meet payment obligations beyond a specified threshold period (typically 60-90 days), or the occurrence of other significant adverse credit events. Regulatory frameworks such as Basel II and III provide standardized definitions linking default to Probability of Default (PD) estimates, risk-weighted assets, and minimum capital requirements [1-3]. Credit scoring converts observable borrower characteristics into quantitative measures of creditworthiness, enabling institutions to make consistent, repeatable decisions across the credit lifecycle.

2.2. Classical statistical models for credit scoring

For decades, classical statistical models served as the foundation of credit scoring. The Altman Z-score model [4] applied multivariate discriminant analysis to bankruptcy prediction, combining financial ratios (working capital/total assets, retained earnings/total assets) into a single composite score. While simple and interpretable, its linear structure limits applicability to complex modern datasets. Logistic regression became the industry standard due to its strong explanatory power—coefficients directly indicate each feature's marginal effect on default probability, computational efficiency, and seamless integration with existing IT systems. Other classical approaches include linear and quadratic discriminant analysis, Probit models, and survival analysis methods. Despite

their interpretability, these models share fundamental limitations: inability to capture nonlinear relationships, sensitivity to outliers, and poor handling of high-dimensional correlated features.

2.3. Limitations of traditional methods

Traditional credit assessment faces three critical limitations [1-3]. First, reliance on structured financial data makes these methods ineffective for "thin-file" borrowers (e.g., small businesses, rural populations), exacerbating financial exclusion. Second, parametric assumptions prevent the capture of nonlinear relationships and higher-order interactions. Third, fixed calibration cycles cannot reflect dynamic market conditions. These limitations have driven the adoption of ML and alternative data sources.

3. Machine learning models for credit default prediction

With the digital transformation of credit operations and the rapid expansion of data availability, ML models have become important supplements—and increasingly alternatives—to traditional methods in credit default prediction. Their inherent advantages in capturing nonlinear relationships, handling high-dimensional features, and adapting to complex classification problems make them particularly well-suited to modern credit risk management. This section reviews basic machine learning classifiers, advanced ensemble learning methods, and appropriate performance evaluation metrics for imbalanced credit datasets.

3.1. Basic machine learning classifiers

Support Vector Machines (SVM) [5] construct optimal classification hyperplanes in feature space, maximizing margins between default and non-default samples. Through kernel functions (RBF, polynomial) [6], SVMs can model complex nonlinear boundaries, showing particular strength with small sample sizes and high-dimensional data. However, SVMs require careful hyperparameter tuning and kernel selection, and their training scales poorly to very large datasets. Decision trees [7] offer intuitive, easily interpretable decision paths through recursive feature-based splitting (information gain, Gini index) [8], naturally handling nonlinear relationships and mixed variable types. Their primary weakness—tendency to overfit—motivated the development of ensemble methods that combine multiple trees to improve stability and generalization performance.

3.2. Advanced ensemble learning methods

Ensemble learning techniques [9] have become the dominant paradigm in credit default prediction due to their superior predictive performance and robustness. Random Forest [10] employs Bagging (bootstrap aggregating) to construct multiple decorrelated decision trees [11], reducing overfitting while providing stable predictions and valuable feature importance measures. This makes it particularly effective for high-dimensional, noisy credit data. Gradient Boosting Decision Trees (GBDT) and their optimized implementations—XGBoost, LightGBM and CatBoost—take a different approach by sequentially fitting new trees to the residuals of previous iterations, progressively correcting prediction errors. XGBoost adds regularization (L1/L2) to control model complexity and prevent overfitting [10], while LightGBM [12] uses leaf-wise tree growth and histogram-based algorithms for superior efficiency with large datasets and sparse features. Empirical studies consistently demonstrate that ensemble models—particularly XGBoost and LightGBM—outperform both classical statistical models and single ML classifiers across standard credit scoring

benchmarks (AUC, accuracy, recall and F1-score) [13]. However, this predictive superiority comes at a cost: ensemble models significantly increase complexity and exacerbate the black-box problem, making their decision processes difficult to interpret through simple rules. This interpretability challenge directly motivates the integration of XAI techniques discussed in subsequent sections.

3.3. Performance evaluation metrics for imbalanced datasets

A defining characteristic of credit default prediction is severe class imbalance: default events typically represent only 1-5% of total observations [14], making overall accuracy misleading. For imbalanced data, AUC-ROC provides aggregate discrimination, while the Precision-Recall (PR) curve is more informative for rare positive classes. Complementary metrics include F1-score, KS statistic, G-mean, and Brier score [13]. In practice, fraud detection prioritizes recall, while portfolio optimization emphasizes AUC, KS, and calibration quality.

4. Key challenges and advanced solutions

Although ML models demonstrate superior performance compared to traditional statistical methods in credit default prediction, their practical deployment faces a series of critical challenges that must be systematically addressed. The most prominent issues include: severe class imbalance resulting from the inherently low proportion of default samples; the high complexity of feature engineering required to extract meaningful signals from massive volumes of behavioral and alternative data; and the ongoing demand for model interpretability and transparency under increasingly stringent regulatory compliance requirements and business decision-making needs. This section systematically examines these three challenges, introduces corresponding state-of-the-art solutions and practical approaches, and establishes the methodological foundation for constructing a comprehensive ML-based risk control and early warning system.

4.1. Strategies for handling imbalanced data

Class imbalance is a major challenge in credit default prediction: default events typically represent only 1%–10% of observations, causing models to optimize for the majority class and miss high-risk borrowers. The asymmetric cost structure—false negatives cost far more than false positives—makes addressing imbalance essential.

Data-level resampling approaches modify the training set distribution to balance class representation. Oversampling methods increase minority class samples: Random oversampling (ROS) duplicates existing samples but risks overfitting. Synthetic Minority Over-sampling Technique (SMOTE) generates new samples by interpolating between minority instances and their k-nearest neighbors, preserving local structure while avoiding exact replication; advanced variants include ADASYN (adaptive synthetic sampling that focuses on hard-to-learn regions) [14], Borderline-SMOTE (targeting samples near class boundaries) [15], and WSMOTE (weighted SMOTE that accounts for sample density). Undersampling methods reduce majority class size: Random Undersampling (RUS) is computationally efficient but may discard informative samples; intelligent methods like Tomek Links (removing overlapping borderline samples) and ENN (Edited Nearest Neighbors, eliminating misclassified majority instances) attempt to preserve decision boundary information while cleaning class overlap [16].

Algorithm-level approaches work within the learning process itself. Cost-sensitive learning assigns asymmetric misclassification costs—typically making false negatives substantially more

expensive than false positives—without modifying the underlying data distribution. This pushes the model to prioritize recall for the minority class. Ensemble methods for imbalanced data include Balanced Random Forest (which undersamples the majority class within each bootstrap sample) and EasyEnsemble/BalanceCascade. The most effective practical solutions typically combine multiple strategies: SMOTE+ENN first oversamples, then cleans noisy synthetic examples; WSMOTE-ensemble+RF integrates weighted oversampling with bagging and random forests. Comprehensive empirical studies by Xiao et al., Abedin et al., and Lenka et al. consistently demonstrate that such hybrid approaches outperform single-method solutions across AUC, G-mean, F1, and KS metrics on benchmark credit scoring datasets [14-16].

4.2. Feature engineering with behavioral and alternative data

Conventional credit assessment relies predominantly on structured, static information from financial statements and historical credit records [17, 18]. While this approach has been demonstrated to be effective for established borrowers with extensive credit histories, it has been observed to struggle with thin-file applicants and to fail to capture the dynamic behavioural patterns that often precede defaults. In recent years, the proliferation of digital financial services has resulted in the generation of vast quantities of behavioural data, including account usage records, repayment patterns, spending habits, mobile device usage, and social network activity [19]. This has resulted in substantial progress in the domains of risk assessment and predictive accuracy.

The academic literature demonstrates that behavioral and alternative data create value in two primary ways. Firstly, they address information gaps for underserved populations. Mobile phone usage patterns [17], digital transaction traces, and social connectivity metrics reveal a substantial correlation with repayment behaviour. Models built from call history and phone usage achieve higher AUC and economic returns than those using only traditional credit bureau data [17]; for unbanked populations without formal credit histories, mobile usage models perform comparably to traditional models trained on scorable populations. Research by Jiang et al. using online banking data demonstrated that models built exclusively on transaction data (ETD) and social stability data (SSD) significantly outperformed historical credit data (HCD) models in identifying credit-invisible micro-enterprise borrowers, suggesting behavioral data can substantially substitute for traditional credit records [20].

Secondly, behavioural data provide dynamic risk perspectives for existing borrowers. The analysis of ongoing transaction records, payment behaviors, and platform usage logs often signals risk deterioration earlier than static variables [18]. Feature engineering transforms these raw data sources into structured, default-sensitive inputs through several approaches: time-window aggregations (7, 30 and 90-day moving averages, standard deviations, and frequency counts); ratio features (debt-to-income, credit utilization, and cash flow coverage); trend indicators (balance volatility, and repayment consistency); and anomaly scores via statistical methods or autoencoders. Social network features—community structure, centrality measures, and network diversity—capture peer effects and social capital dimensions [21]. Representation learning techniques compress high-dimensional behavioral data into low-dimensional embeddings while preserving predictive signal.

The implementation process is fraught with challenges, including the protection of privacy in accordance with GDPR and analogous regulations, ensuring algorithmic fairness to prevent proxy discrimination against protected demographic groups, and ensuring cross-market transferability in light of the significant variation in feature distributions across regions and product lines. The implementation of data minimization, purpose limitation, anonymization techniques, and subgroup fairness audits are essential safeguard.

4.3. Model interpretability and explainable AI (XAI)

The rapid adoption of complex ML models—including ensemble learning, gradient-boosted trees, and deep neural networks—in credit scoring has created a fundamental tension between predictive performance and interpretability. While these models achieve substantially higher accuracy than traditional approaches, they often function as opaque black boxes, unable to clearly articulate their internal decision logic. This opacity creates multiple serious problems [22]: risk managers and business stakeholders cannot understand why particular applications are classified as high-risk; regulators cannot audit decision processes for compliance with fair lending requirements; customers receive unexplained adverse decisions that erode trust; and potential biases embedded in training data or model structure remain hidden and unaddressed. Regulatory frameworks worldwide—including the U.S. Fair Credit Reporting Act (FCRA), Equal Credit Opportunity Act (ECOA), and the EU's General Data Protection Regulation (GDPR)—increasingly mandate transparency and explainability for automated credit decisions, with GDPR specifically establishing a "right to explanation."

Explainable AI (XAI) [22] has emerged as a vital response to these challenges, meeting the diverse information requirements of multiple stakeholders. For example, risk managers need model insights for policy decisions, regulators require auditability, and customers deserve clear explanations [23]. Research shows that visual explanations through SHAP and LIME increase trust in complex models [23]. XAI also enables bias detection by identifying excessive reliance on variables correlated with protected characteristics, supporting interventions for algorithmic fairness. At the methodological level, Yeo et al. [24] categorize explainability approaches into inherently interpretable models (logistic regression, shallow decision trees, and generalized additive models) with directly readable structures, and post-hoc explanation methods designed for black-box models. Post-hoc methods—particularly SHAP (Shapley Additive Explanations) [25] and LIME (Local Interpretable Model-agnostic Explanations) [26]—have become dominant in practice. SHAP [25] provides theoretically grounded feature attribution based on cooperative game theory, offering both global importance rankings (average |SHAP value| per feature) and local explanations (individual prediction decomposition). LIME [22] approximates complex decision boundaries locally with interpretable linear models, generating concise feature contribution lists for single predictions. Studies [23] show that combined SHAP+LIME implementations on XGBoost and LSTM models in digital banking contexts enhance both model team understanding and frontline staff trust.

The FinXAI framework proposed by Yeo et al. [24] structures explainability along four dimensions—audience, data modality, explanation format, and quality evaluation—to address the multifaceted needs of stakeholders ranging from customers to regulators. Although SHAP [25] and LIME [26] remain practically valuable, their deployment in credit risk is hindered by prohibitive computational costs on large-scale deep networks, instability under varying sampling strategies, and the absence of standardized fidelity metrics. Critically, counterfactual explanations also pose security risks by potentially leaking sensitive training information. Future research must therefore move beyond post-hoc approximations toward causally grounded, explainable-by-design architectures, embedding XAI throughout the model lifecycle rather than treating it as an afterthought.

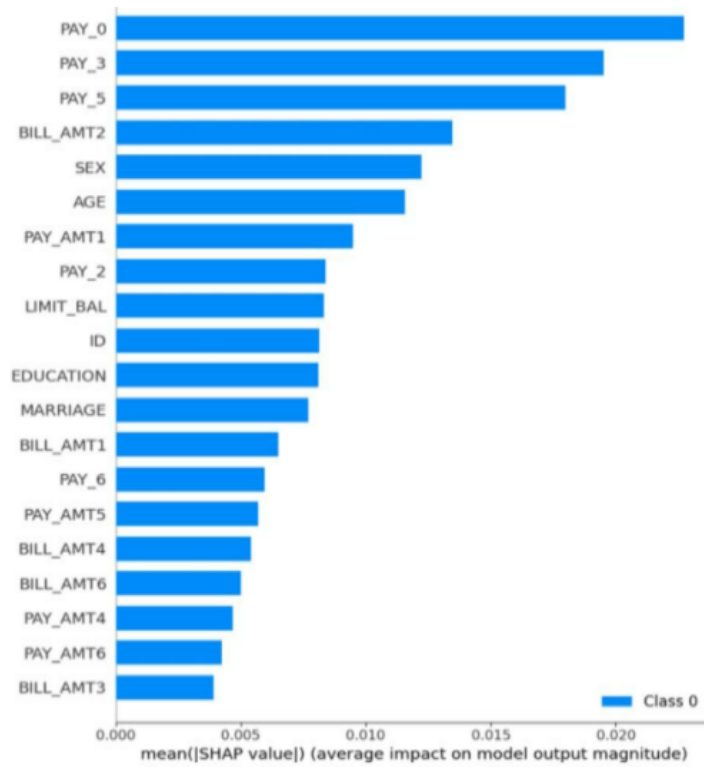


Figure 1. Feature importance for credit default prediction using SHAP model [25]

5. Construction and application of Risk Early Warning Systems

While previous sections addressed individual methodological components in isolation, this section proposes an integrated machine learning-based Risk Early Warning System (REWS) that combines ensemble learning models, behavioral data feature engineering, imbalanced data handling strategies, and XAI techniques into a comprehensive, practical, and auditable end-to-end solution for credit risk management [22, 27]. The framework translates individual technical capabilities into coordinated organizational processes that support proactive risk identification and intervention.

5.1. Framework of ML-based Risk Early Warning Systems

REWS is architected as a six-layer pipeline [22, 27] (Figure 2). Layer 1 (Data Collection) integrates traditional credit, transaction, and alternative data. Layer 2 (Feature Engineering) constructs time-window aggregations, ratio features, and trend indicators. Layer 3 (Model Training) uses ensemble models (XGBoost, and LightGBM) with imbalanced learning [9]. Layer 4 (Explainability) applies SHAP and LIME for global and local interpretation [24-26]. Layer 5 (Decision & Early Warning) classifies borrowers into risk tiers and triggers tiered alerts. Layer 6 (Feedback Loop) monitors drift, retrains models, and adjusts thresholds [18].

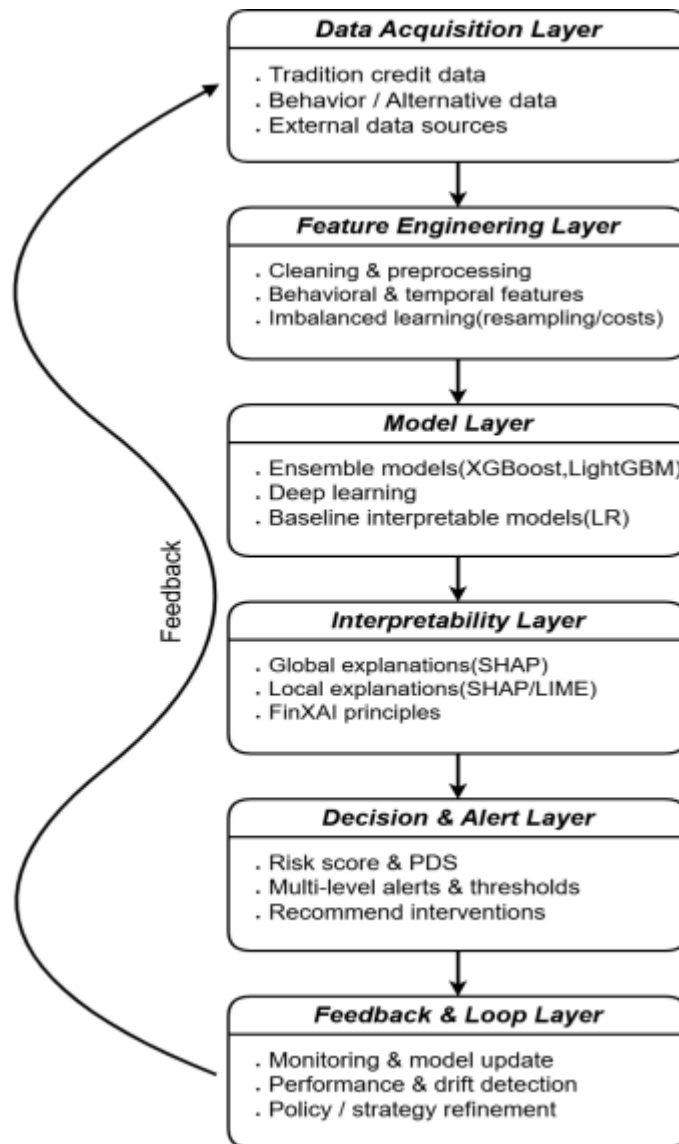


Figure 2. Conceptual framework of the ML-XAI-based credit Risk Early Warning System (REWS)

5.2. Integration and practical implementation

REWS delivers tangible value by automating data pipelines to enable comprehensive portfolio coverage and shifting risk management from reactive to proactive through dynamic detection. It also leverages alternative data to reach thin-file populations, while XAI ensures regulatory compliance and builds trust [18-23]. However, realizing these benefits requires overcoming significant implementation hurdles. Organizations must establish robust data governance, upgrade IT infrastructure for real-time processing, and manage the full model lifecycle. Crucially, they need to design human-AI workflows that reserve expert judgment for complex edge cases. Ultimately, REWS should be viewed as a strategic capability investment rather than a one-time tech deployment, balancing upfront costs against long-term risk reduction.

6. Conclusion

This review systematically examined ML applications in credit default prediction and early warning systems, focusing on ensemble learning, behavioral data feature engineering, and XAI.

This review yields three principal conclusions. Firstly, ensemble learning models (e.g., Random Forest, XGBoost and LightGBM) exhibit superior predictive capability compared with traditional statistical models, particularly when combined with resampling strategies (SMOTE, and ADASYN) and cost-sensitive learning for imbalanced data [9, 14, 15]. Hybrid approaches consistently achieve the best performance across AUC, G-mean and F1 metrics [13]. Secondly, behavioral and alternative data complement traditional credit histories by providing dynamic risk signals and enabling assessment for previously unscorable populations [17,20]. Mobile usage, transaction patterns, and social network data significantly expand financial inclusion while improving prediction accuracy [21]. Thirdly, XAI methods (SHAP, LIME) and frameworks (FinXAI) are essential for regulatory compliance, bias detection, and stakeholder trust in complex ML models [23-26]. The concept of explainability has evolved from a technical feature to a business requirement [22].

The proposed REWS framework integrates these advances into a six-layer end-to-end system for dynamic, explainable, and proactive credit risk management [27]. Unlike traditional static approaches, REWS enables continuous monitoring, early risk detection, and automated intervention triggering.

The following three areas require further research. Firstly, temporal risk modeling: most studies utilize static datasets, yet real-world credit portfolios exhibit continuous concept drift; online learning and adaptive retraining strategies remain underexplored. Secondly, the concept of causal explainability is addressed. SHAP and LIME have been demonstrated to identify feature associations; however, they do not establish causal links between variables and default outcomes. The integration of causal inference into credit scoring would enhance regulatory trust. Thirdly, federated learning: as privacy regulations tighten, training models across institutions without centralizing sensitive data could balance predictive power with compliance.

In conclusion, ML+XAI-driven REWS systems offer a promising pathway for modernizing credit risk management. However, realizing their full potential requires multidisciplinary collaboration across multiple disciplines, encompassing technology, data ethics, and governance. It is imperative that financial institutions consider the adoption of REWS as a strategic, long-term investment in organizational structure, IT infrastructure, and model risk management capabilities. It is incumbent upon regulators to strike a balance between the encouragement of innovation and the mitigation of risk through the establishment of clear explainability, fairness, and privacy requirements.

References

- [1] Basel Committee on Banking Supervision. (2004). *International convergence of capital measurement and capital standards: A revised framework (Basel II)*. Basel: Bank for International Settlements. <https://www.bis.org/publ/bcbs107.htm>
- [2] Basel Committee on Banking Supervision. (2011). *Basel III: A global regulatory framework for more resilient banks and banking systems*. Basel: Bank for International Settlements. <https://www.bis.org/publ/bcbs189.htm>
- [3] Basel Committee on Banking Supervision. (2017). *Basel III: Finalising post-crisis reforms*. Basel: Bank for International Settlements. Available at: <https://www.bis.org/bcbs/publ/d424.htm>
- [4] Altman, E. I. (1968). Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *Journal of Finance*, 23(4), 589-609.
- [5] Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273-297.
- [6] Boser, B. E., Guyon, I. M., & Vapnik, V. N. (1992). A training algorithm for optimal margin classifiers. In *Proceedings of the fifth annual workshop on Computational learning theory*, July, pp. 144-152.

- [7] Breiman, L., Friedman, J., Olshen, R. A., & Stone, C. J. (1984). *Classification and regression trees*. Chapman & Hall Publishing, New York.
- [8] Mienye, I. D., & Jere, N. (2024). A survey of decision trees: Concepts, algorithms, and applications. *IEEE Access*, 12, 86716-86727.
- [9] Mienye, I. D., & Sun, Y. (2022). A survey of ensemble learning: Concepts, algorithms, applications, and prospects. *IEEE Access*, 10, 99129-99149.
- [10] Breiman, L. (2001). Random forests. *Machine learning*, 45(1), 5-32.
- [11] Zou, R. Y., & Schonlau, M. (2018). The random forest algorithm for statistical learning with applications in Stata. *Stata J*, 20(3).
- [12] Ponsam, J. G., Gracia, S. J. B., Geetha, G., Karpaselvi, S., & Nimala, K. (2021, December). Credit risk analysis using LightGBM and a comparative study of popular algorithms. *The 4th International Conference on Computing and Communications Technologies (ICCCCT)*. IEEE, pp. 634-641.
- [13] Lenka, S. R., Bisoy, S. K., Priyadarshini, R., & Sain, M. (2022). Empirical analysis of ensemble learning for imbalanced credit scoring datasets: A systematic review. *Wireless Communications and Mobile Computing*, (1), 6584352.
- [14] Xiao, J., Wang, Y., Chen, J., Xie, L., & Huang, J. (2021). Impact of resampling methods and classification models on the imbalanced credit scoring problems. *Information Sciences*, 569, 508-526.
- [15] Abedin, M. Z., Guotai, C., Hajek, P., & Zhang, T. (2023). Combining weighted SMOTE with ensemble learning for the class-imbalanced prediction of small business credit risk. *Complex & Intelligent Systems*, 9(4), 3559-3579.
- [16] Tomek, I. (1976). An experiment with the edited nearest-neighbor rule. *IEEE Transactions on Systems, Man, and Cybernetics*, SMC-6(6), 448-452.
- [17] Björkegren, D., & Grissen, D. (2020). Behavior revealed in mobile phone usage predicts credit repayment. *The World Bank Economic Review*, 34(3), 618-634.
- [18] Malik, P., Kaul, V., Jagthap, G., Jamley, H., & Chopra, H. (2025). Dynamic credit scoring with real-time transactional & behavioral data using deep reinforcement learning. *International Journal of Science, Engineering and Technology*, 13(5), 253-267.
- [19] Razavi, R., & Elbahnasawy, N. G. (2025). Unlocking credit access: Using non-CDR mobile data to enhance credit scoring for financial inclusion. *Finance Research Letters*, 73, 106682.
- [20] Jiang, M., Shi, J., Zheng, Y., & Zhou, W. (2026). The role of alternative data in micro-enterprises' credit risk assessment in China—Empirical evidence based on machine learning. *Journal of Behavioral and Experimental Finance*, 101154.
- [21] Zhang, L., Wang, L., & Zhu, Y. (2025). Impact of social networks on digital credit assessment for rural residents: A study using machine learning methods. *Information Systems Frontiers*, pp. 1-11.
- [22] Al Shiam, S. A., Hasan, M. M., et al. (2024). Credit risk prediction using explainable AI. *Journal of Business and Management Studies*, 6(2), 61-66.
- [23] Pamisetty, A., Paleti, S., Adusupalli, B., Singireddy, J., Inala, R., & Nagabhyru, K. C. (2025). Explainable AI systems for credit scoring and loan risk assessment in digital banking platforms. *2025 IEEE 13th International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS)*. IEEE, September, pp. 1478-1483.
- [24] Yeo, W. J., Van Der Heever, W., Mao, R., Cambria, E., Satapathy, R., & Mengaldo, G. (2025). A comprehensive review on financial explainable AI. *Artificial Intelligence Review*, 58(6), 189.
- [25] Talaat, F. M., Aljadani, A., Badawy, M., & Elhosseini, M. (2024). Toward interpretable credit scoring: Integrating explainable artificial intelligence with deep learning for credit card default prediction. *Neural Computing and Applications*, 36(9), 4847-4865.
- [26] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?" Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, August, pp. 1135-1144.
- [27] Gentyala, R. (2022). A hybrid machine learning approach for credit scoring: Integrating traditional financial history with mobile phone behavioral metrics. *International Journal of Artificial Intelligence and Machine Learning Research and Development (QITP-IJAIMLRD)*, 3(1), 13-40.