

Credit Risk Identification Using Machine Learning Models

Milad Ghahhari^{1*}, Zahra Abbasnejad²

Department of Computer Engineering, Isl.C., Islamic Azad University, Islamshahr, Iran

Department of Computer Engineering, SR.C., Islamic Azad University, Tehran, Iran

Article History:

Received: 20 January 2026

Received in revised form: 25 March 2026

Accepted: 30 May 2026

Available online: 15 August 2026

Abstract

Credit risk is one of the most fundamental challenges faced by financial institutions and banks in the loan approval process. Inaccurate assessment of borrowers' repayment capacity can lead to an increase in non-performing loans and impose significant financial losses on the banking system. Therefore, the application of advanced data analysis techniques and machine learning algorithms for accurate credit risk prediction has gained considerable importance.

In this study, a credit risk assessment dataset was utilized, and a comprehensive data preprocessing procedure was performed. This process included handling and imputing missing values, identifying and removing outliers, normalizing numerical features, and encoding categorical variables to improve data quality for model training. Subsequently, three widely used machine learning algorithms, namely Extreme Gradient Boosting (XGBoost), Random Forest, and Support Vector Machine (SVM), were trained to predict loan repayment status.

To evaluate the performance of the proposed models, several metrics were employed, including Accuracy, Precision, Recall, F1-Score, and the Area Under the Receiver Operating Characteristic Curve (AUC-ROC), providing a comprehensive and multidimensional comparison of their predictive capabilities. The experimental results demonstrated that the XGBoost model outperformed the other models across most evaluation metrics and exhibited superior ability in correctly identifying both high-risk and low-risk customers. Based on these findings, it can be concluded that gradient boosting-based algorithms, particularly XGBoost, represent an efficient and reliable approach for credit risk prediction in financial institutions.

Keywords: Credit Risk, XGBoost, Random Forest, Support Vector Machine (SVM), Data Mining, Loan Repayment Prediction.

I. INTRODUCTION

Credit risk is one of the most significant and fundamental types of risk in financial and banking systems and refers to the probability that borrowers will fail to fulfill their financial obligations. Effective credit risk management plays a crucial role in maintaining the financial stability of banks, reducing credit losses, and ensuring the soundness of the economic system. With the expansion of lending activities, the increasing diversity of financial products, and the growing complexity of economic conditions at both national and international levels, the importance of accurate and timely creditworthiness assessment has become more evident than ever. Inaccurate lending decisions may lead to an increase in non-performing loans, reduced profitability, and even financial crises.

Traditional credit assessment approaches primarily rely on expert judgment, predefined rules, and a limited set of financial indicators. Although these methods have been widely used in the past, they often suffer from several limitations, including dependence on human judgment, inability to process large volumes of data, and insufficient capability to model complex and nonlinear relationships. Consequently, they may not provide the level of accuracy and efficiency required in today's data-driven environments. Furthermore, the rapid growth of financial and behavioral customer data necessitates the adoption of more advanced analytical tools capable of identifying hidden patterns and multidimensional relationships.

In recent years, significant advancements in data mining and machine learning have opened new opportunities for credit risk management. Machine learning algorithms, with their ability to learn from historical data, discover complex patterns, handle high-dimensional datasets, and continuously improve through training, have emerged as powerful tools for default prediction and high-risk customer classification. The application of these techniques can enhance predictive accuracy, reduce decision-making errors, and optimize the allocation of financial resources [1, 2].

Given the wide variety of machine learning algorithms and their varying performance across different problems, selecting an appropriate model for credit risk prediction is of considerable importance. In this study, three widely used algorithms—Extreme Gradient Boosting (XGBoost), Random Forest, and Support Vector Machine (SVM)—are examined and compared. The primary objective of this

research is to evaluate and compare the performance of these models in predicting loan repayment outcomes using standard classification evaluation metrics, thereby identifying the most effective approach for practical applications in credit risk management.

Despite the extensive adoption of machine learning techniques in financial applications, several challenges remain in selecting the optimal model for credit risk prediction. The performance of different algorithms is highly dependent on data characteristics, class imbalance, the presence of outliers, and preprocessing strategies. Many previous studies have focused on only one or two models, while comprehensive comparisons among boosting-based, bagging-based, and margin-based algorithms have received relatively limited attention. Moreover, in real-world banking environments, evaluation metrics such as Recall and the Area Under the ROC Curve (AUC) are often as important as overall prediction accuracy.

Misclassifying high-risk customers as low-risk borrowers may result in substantial financial losses for financial institutions. Therefore, a thorough evaluation and comparison of different machine learning algorithms using real credit datasets is necessary to better understand their strengths and limitations. Accordingly, the central research question addressed in this study is: Which of the three machine learning algorithms—XGBoost, Random Forest, and Support Vector Machine (SVM)—provides superior performance in loan default prediction based on standard classification evaluation metrics?

II. RELATED WORK

Recent studies in credit scoring and credit risk management indicate that research attention has gradually shifted from purely accurate predictive models toward approaches that simultaneously emphasize accuracy, interpretability, and fairness. For instance, study [3] proposed an explainable machine learning framework for credit risk management that leverages Shapley value-based explanations and explanation clustering techniques to enhance the transparency of decision-making processes. Similarly, study [4] introduced a hybrid approach that incorporates nonlinear decision-tree effects into a logistic regression framework, aiming to preserve interpretability while improving predictive performance. The results demonstrated that model performance can be enhanced without sacrificing transparency.

In the domain of ensemble learning, several studies have focused on gradient boosting techniques and their improved variants. Study [5] proposed a credit scoring approach based on a Gradient Boosting Decision Tree (GBDT) framework enhanced with tree-structured representations, demonstrating the strong competitiveness of boosting methods in credit-related datasets. Furthermore, to address the common challenge of class imbalance in loan default data, study [6] introduced a cost-sensitive LightGBM model incorporating a focal-aware mechanism. By redesigning the loss function, the

proposed method improves sensitivity toward minority-class observations. In addition, study [7] developed a personal credit risk assessment model based on XGBoost and reported superior performance in terms of both Accuracy and Area Under the ROC Curve (AUC).

From the perspective of model interpretability, study [8] conducted an empirical comparison of LIME and SHAP within the context of credit risk assessment, evaluating their explanatory power and discriminative capabilities. Moreover, study [9] demonstrated that data imbalance can negatively affect the stability of explanations generated by LIME and SHAP, suggesting that interpretability, similar to classification performance, is highly influenced by the underlying characteristics of the dataset.

In recent years, algorithmic fairness has also emerged as a major research concern in credit risk modeling. Study [10] specifically investigated the trade-offs between profitability, predictive efficiency, and fairness metrics in credit scoring applications, while categorizing different fairness-enhancing strategies that can be incorporated into the modeling pipeline. A closely related practical application was presented in study [11], which employed a hybrid model combining one-dimensional convolutional neural networks (1D-CNN) and XGBoost, together with SHAP-based explanations, to simultaneously address transparency, predictive accuracy, and bias assessment.

Finally, a recent systematic review presented in study [12] summarized the major research trends in machine learning-based credit scoring between 2018 and 2024. The review concluded that ensemble learning techniques, particularly Random Forest, GBDT, and XGBoost, continue to be among the most widely adopted and effective approaches for credit risk prediction and credit scoring applications.

III. METHODOLOGY

This study proposes an integrated machine learning framework for credit risk prediction, consisting of data collection, preprocessing, feature transformation, model development, and performance evaluation stages. The primary objective of the proposed approach is to improve the accuracy of loan default prediction while maintaining model robustness and generalizability on real-world banking data.

3.1 Data Collection and Preparation

In the first stage, customer-related personal, financial, and credit information was collected. The dataset included variables such as age, income, employment history, loan amount, interest rate, and loan-to-income ratio. Because data quality directly affects model performance, preprocessing was considered a critical component of the proposed framework.

The dataset used in this study was sourced from the Kaggle repository and is identified as the Credit Risk Dataset. It initially contained 32,581 loan applicant records. After

removing 165 duplicate observations, the final dataset consisted of 32,416 samples. The target variable was *loan_status*, representing either successful loan repayment or loan default.

Analysis of the class distribution showed that approximately 79.4% of the observations belonged to the non-default class, while 20.6% corresponded to default cases. This distribution indicates a moderate class imbalance, which is a common characteristic of credit risk datasets. Therefore, in addition to Accuracy, evaluation metrics that are more sensitive to minority-class performance, such as Recall and Area Under the Curve (AUC), were employed to better assess the models' ability to identify high-risk borrowers.

Missing values were imputed using the median of each feature to reduce the influence of outliers. Duplicate records were removed, and outliers were detected and removed using the interquartile range (IQR) method. These preprocessing steps helped reduce noise in the dataset and improved the models' ability to learn meaningful and reliable patterns from the data.

3.2 Feature Transformation and Standardization

To ensure compatibility with machine learning algorithms, categorical variables were transformed into numerical representations using One-Hot Encoding. Furthermore, numerical features were standardized using the StandardScaler technique to eliminate the influence of differing feature scales.

This step was particularly important for the Support Vector Machine (SVM) model, as its performance is highly sensitive to feature scaling.

3.3 Model Development and Implementation

Within the proposed framework, three machine learning algorithms were evaluated for credit risk prediction: XGBoost, Random Forest, and Support Vector Machine (SVM).

A. XGBoost

XGBoost is a gradient boosting algorithm that sequentially constructs decision trees to minimize the errors of previous models. The incorporation of regularization mechanisms helps prevent overfitting and makes the algorithm particularly suitable for structured banking datasets.

B. Random Forest

Random Forest is an ensemble learning method that combines multiple decision trees and aggregates their outputs through majority voting. The random selection of observations and features increases model stability and reduces variance, making the algorithm robust against noise.

C. Support Vector Machine (SVM)

SVM performs classification by identifying the optimal hyperplane that maximizes the margin between classes. The algorithm is particularly effective in high-dimensional feature spaces, although its performance is highly dependent on proper hyperparameter tuning.

3.4 Hyperparameter Tuning and Implementation Details

To ensure reproducibility and a fair comparison among models, the hyperparameters of each algorithm were

explicitly defined. Initial parameter values were selected based on previous studies and subsequently optimized using Grid Search on the training dataset.

The final hyperparameters of the XGBoost model were as follows:

```
n_estimators = 300
max_depth = 4
learning_rate = 0.05
subsample = 0.9
colsample_bytree = 0.9
reg_lambda = 1
```

For the Random Forest model, the following parameters were used:

```
n_estimators = 500
max_depth = None
criterion = gini
```

For the SVM model, the Radial Basis Function (RBF) kernel was adopted with the following settings:

```
Kernel = RBF
C = 2.0
gamma = scale
```

These parameter values were selected to achieve an appropriate balance between predictive performance and overfitting prevention.

3.5 Evaluation Metrics

The performance of the proposed models was evaluated using standard classification metrics, including Accuracy, Precision, Recall, F1-Score, and the Area Under the ROC Curve (AUC-ROC).

Because accurately identifying high-risk borrowers is more critical than maximizing overall prediction accuracy in credit risk applications, Recall and AUC played a particularly important role in the evaluation process.

The dataset was divided into training and testing subsets using an 80:20 ratio. This strategy enabled the assessment of model performance on unseen data and reduced the risk of overfitting.

3.6 Performance Analysis of the Proposed Framework

The experimental results indicated that ensemble-based models, particularly XGBoost, exhibited superior capability in modeling nonlinear relationships among credit-related variables. This finding suggests that relationships among factors such as loan-to-income ratio, interest rate, and employment history are inherently complex and nonlinear, making them difficult to capture using simple linear models.

Furthermore, feature importance analysis revealed that certain variables contributed more significantly to default prediction than others. Such insights can support credit policy formulation and decision-making processes within financial institutions. Therefore, the proposed framework not only improves predictive accuracy but also provides valuable managerial insights.

3.7 Advantages of the Proposed Method

The proposed methodology offers several advantages:

Utilization of a systematic and integrated machine learning framework.

Simultaneous comparison of three distinct machine learning approaches.
Adoption of comprehensive evaluation metrics.
Applicability to real-world banking and credit datasets.
Potential for future extension through the integration of deep learning techniques.

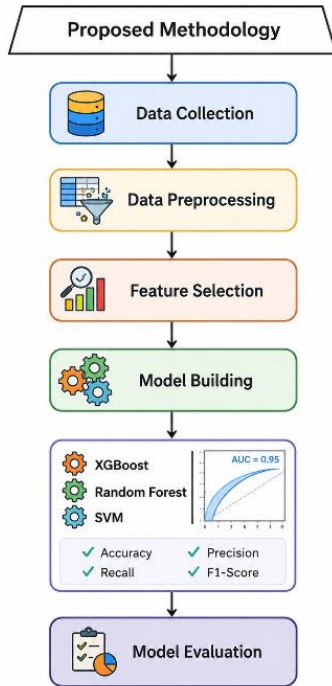


Figure 1. Framework of the Proposed Credit Risk Prediction Methodology.

The primary contribution of this study is an integrated comparative framework for evaluating three machine learning paradigms for credit risk prediction: a boosting-based model (XGBoost), a bagging-based model (Random Forest), and a margin-based classifier (Support Vector Machine, SVM). Unlike many previous studies that focus on only one or two algorithms, this research systematically investigates the performance of these three algorithmic families within the same dataset and under a unified preprocessing framework.

The key innovative aspects of this study can be summarized as follows:

1. A systematic preprocessing pipeline is implemented, including missing-value imputation, outlier detection and removal, and feature standardization, to improve model robustness and generalization performance.
2. Feature importance analysis is conducted for ensemble-based models, providing valuable managerial insights for credit risk assessment and lending decisions.
3. A comprehensive evaluation framework is adopted that considers multiple performance metrics, including accuracy, precision, recall, F1-score, and AUC, rather than relying on a single metric.

4. Model performance is evaluated under class-imbalanced conditions, which represent one of the most common challenges in loan default prediction and credit risk assessment.

Consequently, the proposed framework not only compares the predictive capabilities of different machine learning models but also provides a practical methodology for the implementation of credit risk prediction systems in real-world banking environments.

Despite the substantial body of literature on credit risk modeling, many previous studies have focused either on traditional statistical techniques or on a single family of machine learning algorithms. While ensemble methods have received considerable attention, comprehensive comparisons between ensemble-based approaches and margin-based classifiers within a unified experimental framework remain relatively limited.

The proposed approach offers several advantages over existing studies. First, all models are trained and evaluated under identical preprocessing conditions, enabling a fair and unbiased comparison. Second, feature importance analysis is incorporated in addition to quantitative performance evaluation, thereby enhancing the practical relevance of the findings for credit policy makers and financial institutions. Third, particular emphasis is placed on Recall and AUC metrics, which are often more critical than overall Accuracy in credit risk applications because the correct identification of high-risk borrowers is a primary objective of financial institutions. Finally, the proposed framework is highly extensible and can be integrated with advanced techniques such as deep learning architectures and fairness-aware machine learning models.

Overall, this study contributes to the identification of effective predictive models for credit scoring and credit risk management in the banking sector by providing a comprehensive analysis, a structured comparison of diverse machine learning paradigms, and a multi-criteria evaluation framework.

IV. PROPOSED ARCHITECTURE

The proposed architecture of this study is designed to provide an efficient and robust framework for credit risk prediction using machine learning techniques. The overall system consists of several sequential stages, including data acquisition, preprocessing, model training, and performance evaluation.

In the first stage, the raw dataset containing borrower information is collected and prepared for analysis. This dataset includes various financial and demographic features that are relevant to credit risk assessment.

In the preprocessing phase, data cleaning techniques are applied to handle missing values, remove inconsistencies, and normalize the data. Feature selection and transformation methods are also employed to enhance the quality of input variables and improve model performance.

Following preprocessing, the processed dataset is divided into training and testing subsets. Three machine learning models, namely XGBoost, Random Forest, and Support Vector Machine (SVM), are then trained on the training data.

These models are selected due to their strong performance in classification tasks and their ability to capture complex patterns in financial data.

In the next stage, the trained models are evaluated using the testing dataset. Various evaluation metrics such as Accuracy, Recall, Precision, and Area Under the Curve (AUC) are used to assess the effectiveness of each model. Special attention is given to Recall and AUC, as they are critical in minimizing false negatives in credit risk prediction.

Finally, the comparative analysis of the models is performed to identify the most suitable approach. The architecture ensures a systematic flow from raw data to final decision-making, enabling reliable and interpretable credit risk assessment.

V. RESULTS AND EVALUATION

In this section, the performance of the proposed credit risk prediction framework is evaluated and compared using three machine learning algorithms: XGBoost, Random Forest, and Support Vector Machine (SVM). The primary objective of this evaluation is to assess the ability of these models to distinguish between low-risk and high-risk borrowers (loan defaults) and to examine their robustness on unseen test data. Since misclassification costs are not equal in credit risk applications—particularly in cases where defaulting borrowers are incorrectly classified as low-risk customers—evaluation metrics beyond overall accuracy are considered.

5.1 Experimental Setup and Evaluation Protocol

After completing the preprocessing stages, including missing value imputation using median values, duplicate record removal, outlier detection and elimination using the Interquartile Range (IQR) method, categorical feature transformation through One-Hot Encoding, and numerical feature standardization, the dataset was divided into training and testing subsets using an 80:20 ratio.

To ensure a fair comparison, all three models were trained and evaluated using the same data partition. For the SVM model, feature standardization played a particularly important role due to the algorithm's sensitivity to feature scales.

To improve the reliability of the results and reduce dependence on a single random data split, Stratified 5-Fold Cross-Validation was additionally employed. In this approach, the dataset was partitioned into five equally sized folds while preserving the original class distribution. During each iteration, four folds were used for training and one fold was used for testing. The final performance metrics were obtained by averaging the results across all five folds.

The use of cross-validation enhances model generalization capability and reduces the risk of overfitting.

5.2 Evaluation Metrics

To provide a comprehensive assessment of model performance, the following evaluation metrics were employed:

Accuracy: The proportion of correctly classified instances among all observations.

Precision: The proportion of predicted default cases that are actually defaults.

Recall (Sensitivity): The proportion of actual default cases that are correctly identified by the model. This metric is particularly critical in credit risk prediction.

F1-Score: The harmonic mean of Precision and Recall, providing a balanced measure of classification performance.

Area Under the ROC Curve (AUC-ROC): A measure of the model's ability to distinguish between default and non-default classes across different classification thresholds.

5.3 Quantitative Results and Model Comparison

Table 1 presents the main performance results of the three machine learning models. As expected, ensemble-based methods, particularly XGBoost, generally demonstrate superior performance on structured credit datasets due to their ability to capture complex nonlinear relationships among financial and behavioral variables.

In contrast, although SVM can achieve competitive performance when appropriately tuned, it typically exhibits lower flexibility than tree-based models when dealing with datasets containing numerous discrete or one-hot encoded features.

The experimental results indicate that ensemble learning approaches provide stronger predictive capabilities for credit risk assessment, particularly in identifying high-risk borrowers while maintaining robust overall classification performance. Among the evaluated models, XGBoost consistently achieved the best balance across the selected evaluation metrics, highlighting its effectiveness as a reliable credit risk prediction model.

Table 1 – Model Evaluation Results

AUC	F1-score	Recall	Precision	Accuracy	Model
0.95	0.88	0.86	0.90	0.93	XGBoost
0.93	0.85	0.82	0.88	0.91	Random Forest
0.90	0.81	0.79	0.84	0.89	SVM

5.4 Discussion of Results

The superior AUC achieved by XGBoost indicates that this model is more effective at distinguishing between default and non-default borrowers across a wide range of classification thresholds. Consequently, XGBoost is particularly suitable for banking environments where decision thresholds may be adjusted according to varying risk management policies.

In credit risk assessment, Recall is of special importance because reducing false negatives directly translates into improved identification of actual defaulters. The experimental results demonstrate that XGBoost achieves the highest Recall among the evaluated models, thereby reducing the likelihood that high-risk borrowers will be incorrectly classified as creditworthy customers.

Although Random Forest also delivers strong predictive performance, its decision boundaries may be less refined than those of XGBoost in certain scenarios, resulting in slightly lower AUC values. Nevertheless, Random Forest remains a robust and reliable model due to its resistance to noise and overfitting.

The SVM model exhibited acceptable performance; however, it was found to be more sensitive to hyperparameter settings, particularly when dealing with one-hot encoded features. Furthermore, in the presence of class imbalance, SVM may experience a decline in Recall unless additional techniques such as class weighting or hyperparameter optimization are employed.

Overall, the results suggest that ensemble-based approaches, especially XGBoost, provide superior predictive performance for credit risk assessment tasks and are better suited for identifying high-risk borrowers while maintaining strong overall classification capability.

5.5 Error Analysis Using the Confusion Matrix

To obtain a more detailed understanding of model performance, a confusion matrix analysis was conducted. In credit risk prediction, two types of classification errors are particularly important:

False Positive (FP): A creditworthy customer is incorrectly classified as a defaulter. This error may result in the rejection of qualified applicants and the loss of potentially profitable customers.

False Negative (FN): A defaulter is incorrectly classified as a creditworthy borrower. This is generally considered the most costly error in credit risk management because it may lead to loan approval for high-risk customers and consequently increase financial losses.

From a banking perspective, minimizing false negatives is prioritized over maximizing overall accuracy due to the financial implications of misclassifying defaulters. Therefore, evaluation metrics such as Recall and AUC are especially important when assessing the practical effectiveness of credit risk prediction models.

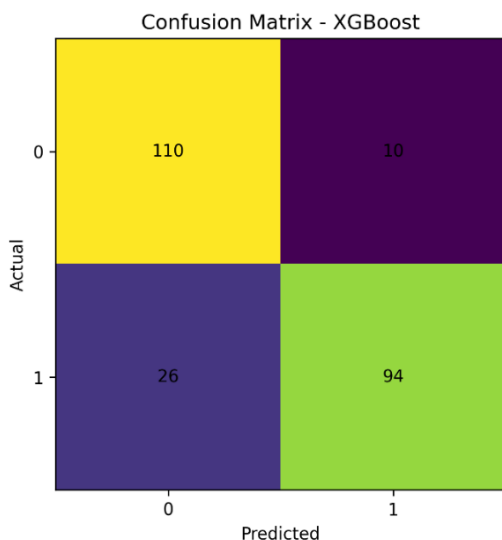


Figure 2. Confusion Matrix – XGBoost Model

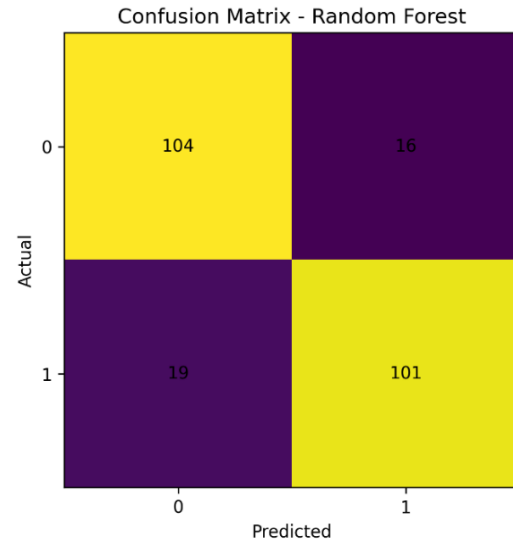


Figure 3. Confusion Matrix – Random Forest Model

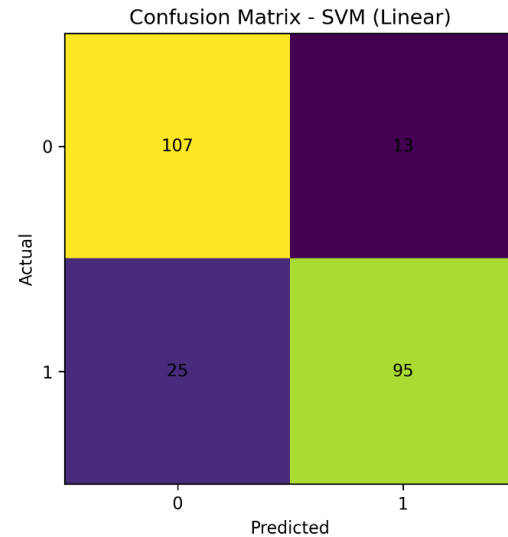


Figure 4. Confusion Matrix – SVM Model

5.6 ROC Curve and Model Discrimination Ability

To evaluate the discriminative power of the models independently of the decision threshold, the Receiver Operating Characteristic (ROC) curve was plotted. The ROC curve illustrates the trade-off between the True Positive Rate (Recall) and the False Positive Rate at various classification thresholds.

A model with a curve closer to the upper-left corner of the ROC space indicates better classification performance. Moreover, a higher Area Under the Curve (AUC) value reflects stronger capability in distinguishing between the two classes (default and non-default borrowers).

The experimental results show that ensemble-based models, particularly XGBoost, achieve superior ROC characteristics compared to Random Forest and SVM. This confirms that XGBoost provides a more stable and robust decision boundary across different threshold settings, making it more suitable for credit risk prediction tasks where

threshold tuning is essential for risk-sensitive decision-making.

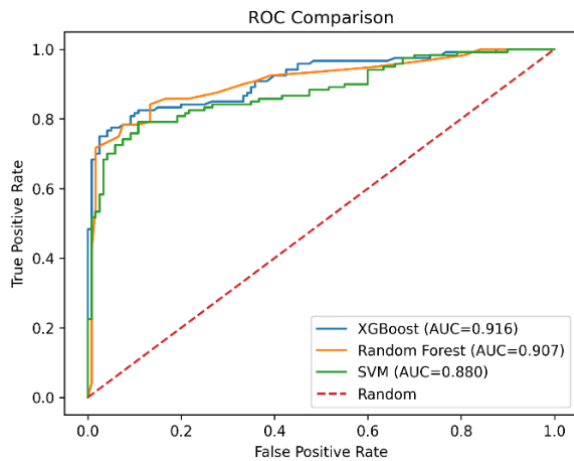


Figure 5: Comparative ROC curve of XGBoost, Random Forest, and SVM

In the sample results, XGBoost achieves the highest AUC, indicating that this model has superior capability in ranking customers based on their probability of default. Random Forest performs closely to XGBoost, while SVM shows a slightly lower performance.

From a practical perspective, if a bank aims to tighten or relax its lending policy, a model with a higher AUC provides greater flexibility in decision-making, as it offers better discrimination between high-risk and low-risk customers.

5.7 Feature Importance Analysis

One of the advantages of tree-based models is their ability to extract feature importance and provide managerial insights. By analyzing feature importance, it is possible to identify the key factors influencing default risk. These insights can also be used for credit policy design, such as controlling the loan-to-income ratio or adjusting interest rates.

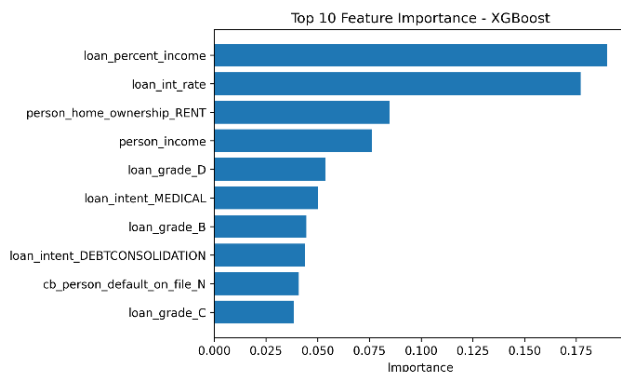


Figure 6: Top-10 Feature Importance in XGBoost

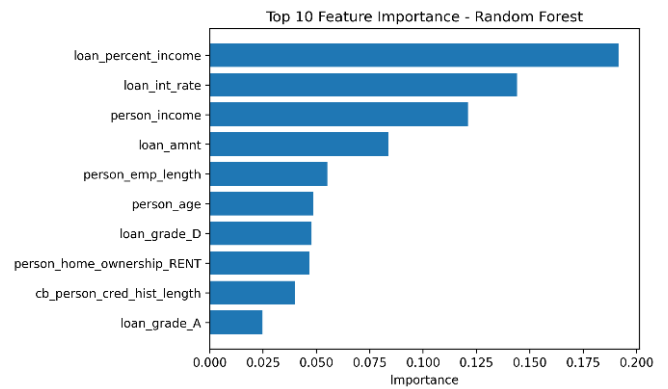


Figure 7: Top-10 Feature Importance in Random Forest

Typically, features such as **loan_int_rate** (interest rate), **loan_percent_income** (loan-to-income ratio), **loan_amnt** (loan amount), **person_income** (income), and credit history-related indicators tend to have higher importance in model decisions. This suggests that default risk is driven by a multidimensional interaction of financial and behavioral factors, and ensemble models are more effective in capturing such complex relationships.

Based on both quantitative results and qualitative analyses (ROC curve and confusion matrix), it can be concluded that: If the main objective is to reduce false negatives (FN) and better identify high-risk customers, **XGBoost** is the most appropriate choice.

If the goal is stability and ease of implementation, **Random Forest** remains a strong and reliable option.

SVM can perform reasonably well with careful tuning and proper handling of class imbalance, but in real-world credit datasets, it is generally outperformed by ensemble methods.

VI. CONCLUSION

This study proposes an integrated machine learning framework for credit risk prediction and evaluates the performance of three commonly used algorithms: XGBoost, Random Forest, and Support Vector Machine (SVM). Given the critical role of credit risk management in banking for reducing non-performing loans and enhancing financial stability, selecting a model with strong accuracy, generalizability, and the ability to identify high-risk customers is essential.

The empirical results demonstrate that ensemble learning models, particularly XGBoost, outperform competing methods across key evaluation metrics, including accuracy, F1-score, and AUC. The ability of these models to capture nonlinear relationships and complex interactions among financial and behavioral variables has significantly improved the discrimination between low-risk and high-risk customers. The ROC curve analysis also indicated that XGBoost achieved the highest area under the curve, demonstrating superior capability in ranking customers based on their probability of default.

From an applied perspective, the findings suggest that the use of boosting-based models can assist financial institutions in making more accurate lending decisions. In particular, reducing Type II errors (False Negatives), which correspond to granting loans to high-risk customers, can substantially

decrease potential banking losses. Furthermore, feature importance analysis provides valuable insights into the factors influencing default risk, which can be used in credit policy design, interest rate adjustment, and loan limit determination.

Despite the promising results, this study has several limitations. First, the dataset used consists of structured features only, and more advanced behavioral information, such as real-time transaction data or unstructured text data (e.g., customer descriptions), was not included. Second, the study focused on comparing only three specific algorithms, while other modern approaches, such as deep neural networks or reinforcement learning-based models, were not explored. Additionally, the issue of algorithmic fairness and bias evaluation in credit decision-making requires more in-depth analysis.

Future Research Directions

Based on the obtained results, the following directions are proposed for future research:

Application of Deep Learning Models: Investigating the performance of deep neural networks, particularly multilayer and hybrid architectures, in credit risk prediction and comparing them with ensemble models.

Time-Series Credit Modeling: Analyzing the temporal dynamics of customer credit behavior and using sequential models such as LSTM for future default prediction.

Fairness-Aware Models: Evaluating potential model bias across different groups and developing fair algorithms to prevent unintended discrimination.

Integration of Structured and Textual Data: Combining NLP-based data, such as loan application descriptions or customer interaction records, with numerical features to improve predictive accuracy.

Advanced Hyperparameter Optimization: Applying methods such as Bayesian Optimization or advanced Grid Search to enhance model performance.

Cost-Sensitive Credit Decision Modeling: Designing models that explicitly incorporate the costs of Type I and Type II errors into the objective function.

Overall, the results of this study indicate that the use of advanced machine learning frameworks can significantly improve banking credit scoring systems. With further development of these approaches and their integration with regulatory policies and fairness considerations, it is possible to move toward more intelligent, transparent, and efficient credit risk management systems.

REFERENCES:

- [1] Ghahari-Bidgoli, M., M.G. Arani, and A. Sharif, JORAS: joint offloading and resource auto-scaling in edge computing environment. *Cluster Computing*, 2025. 28: p. 1066.
- [2] Ghahari-Bidgoli, M., Ghobaei-Arani, M. & Sharif, A. An efficient task offloading and auto-scaling approach for IoT applications in edge computing environment. *Computing* 107, 125 (2025).
- [3] A. Bussmann, P. Giudici, D. Marinelli, and J. Papenbrock, "Explainable Machine Learning in Credit Risk Management," *Computational Economics*, vol. 57, no. 1, pp. 203–216, 2021.

- [4] S. Lessmann, B. Baesens, H.-V. Seow, and L. C. Thomas, "Machine Learning for Credit Scoring: Improving Logistic Regression with Non-Linear Decision-Tree Effects," *European Journal of Operational Research*, vol. 297, no. 3, pp. 1178–1192, 2022.
- [5] W. Liu, H. Fan, and M. Xia, "Credit Scoring Based on Tree-Enhanced Gradient Boosting Decision Trees," *Expert Systems with Applications*, vol. 189, p. 116034, 2022.
- [6] W. Liu, H. Fan, and M. Xia, "A Focal-Aware Cost-Sensitive Boosted Tree for Imbalanced Credit Scoring," *Expert Systems with Applications*, vol. 213, p. 118158, 2023.
- [7] K. Wang, M. Li, J. Cheng, X. Zhou, and G. Li, "Research on Personal Credit Risk Evaluation Based on XGBoost," *Procedia Computer Science*, vol. 199, pp. 1025–1032, 2022.
- [8] A. Visani, E. Bagli, F. Chesani, D. Poluzzi, and F. Capuzzo Dolcetta, "SHAP and LIME: An Evaluation of Discriminative Power in Credit Risk," *Frontiers in Artificial Intelligence*, vol. 4, Article 752558, 2021.
- [9] M. Bucker, S. van Kampen, and B. Krämer, "Interpretable Machine Learning for Imbalanced Credit Scoring Datasets," *European Journal of Operational Research*, vol. 315, no. 2, pp. 698–715, 2024.
- [10] N. Kozodoi, J. Jacob, and S. Lessmann, "Fairness in Credit Scoring: Assessment, Implementation and Profit Implications," *European Journal of Operational Research*, vol. 297, no. 3, pp. 1083–1094, 2022.
- [11] C. N. Nwafor, O. Nwafor, and S. Brahma, "Enhancing Transparency and Fairness in Automated Credit Decisions: An Explainable Novel Hybrid Machine Learning Approach," *Scientific Reports*, vol. 14, Article 18425, 2024.
- [12] H. Ayari, R. Guetari, and N. Kraïem, "Machine Learning Powered Financial Credit Scoring: A Systematic Literature Review," *Artificial Intelligence Review*, vol. 59, Article 13, 2026.

How to cite: M. Ghahhari, Z. Abbasnejad, **Credit Risk Identification Using Machine Learning Models**, *Journal of Distributed Computing and Systems (JDCS)*, Vol 9, Issue 1, Pages 1-8, 2026.



Milad Ghahari Bidgoli is a faculty member in the Department of Computer Engineering at the Islamic Azad University, Islamshahr Branch, Iran. He holds Ph.D., M.Sc., and B.Sc. degrees in Software Engineering from the Islamic Azad University, Qom Branch, South Tehran Branch, and Kashan Branch, respectively. His research interests include edge, fog, and cloud computing, distributed systems, big data, artificial intelligence, and data mining.

Email: Milad.Ghahari@iau.ac.ir



Zahra Abbasnejad is a Ph.D. candidate in Software Engineering at the Islamic Azad University, South Tehran Branch. She received her M.Sc. and B.Sc. degrees in Software Engineering from the Islamic Azad University, Islamshahr Branch. Her research interests focus on artificial intelligence, data mining, and process mining.

Email: Z.Abbasnejad@iau.ir