



Explainable Ensemble Learning for Loan Approval Prediction Using XGBoost, LightGBM, and Random Forest with SHAP Analysis

Mikaria Gultom^{1,*}, Chininta Rizka Angelia²,
Vijayaletchumy Krishnan³

¹Faculty of Business, Universitas Multimedia Nusantara, Indonesia

²Faculty of Communication, Universitas Multimedia Nusantara, Indonesia

³Faculty of Business and Communication INTI International University, Nilai, Malaysia

ABSTRACT

Loan approval prediction plays an important role in financial decision-making, as accurate and transparent assessments are essential for effective risk management. However, imbalanced datasets and the limited interpretability of machine learning models remain major challenges in developing reliable loan approval systems. This study proposes an explainable machine learning framework that evaluates Random Forest, XGBoost, and LightGBM for loan approval prediction. A dataset containing 5,000 records and 10 demographic, financial, and employment-related variables was used. To address class imbalance, the Synthetic Minority Over-sampling Technique (SMOTE) was applied to the training data. The models were evaluated using Accuracy, Precision, Recall, F1-score, and ROC-AUC. The results demonstrate that all three models achieved strong predictive performance, with Random Forest providing the best overall results, achieving an Accuracy of 0.959, Precision of 0.932, Recall of 0.887, F1-score of 0.909, and ROC-AUC of 0.9386. To improve model transparency, SHapley Additive exPlanations (SHAP) was applied to the best-performing model. The analysis identified CreditScore, EmploymentType, and Income as the most influential features in loan approval predictions. SHAP dependence analysis further revealed non-linear relationships and interactions between CreditScore and EmploymentType, while local explanations demonstrated how individual features contributed to specific predictions. These findings indicate that combining ensemble learning, SMOTE, and SHAP can provide a highly accurate and interpretable framework for supporting transparent and reliable loan approval decisions.

Keywords Loan Approval Prediction, Ensemble Learning, Explainable AI, SHAP, SMOTE

INTRODUCTION

Loan approval prediction is a critical task in the financial sector, as it directly influences risk management, operational efficiency, and profitability of lending institutions. Financial organizations must evaluate applicants based on multiple factors such as credit history, income, and employment status to determine their creditworthiness. Traditional decision-making processes are often time-consuming and subject to human bias, which has led to the increasing adoption of machine learning techniques to automate and improve the accuracy of loan approval decisions. In recent years, various classification models have been

Submitted 28 January 2026

Accepted 10 March 2026

Published 27 August 2026

Corresponding author
Mikaria Gultom,
mikaria.gultom@umn.ac.id

Additional Information and
Declarations can be found on
[page 208](#)

DOI: [10.47738/jdmdc.v3i3.74](#)

© Copyright
2026 Gultom, et al.

Distributed under
Creative Commons CC-BY 4.0

How to cite this article: M. Gultom, C. R. Angelia, V. Krishnan, "Explainable Ensemble Learning for Loan Approval Prediction Using XGBoost, LightGBM, and Random Forest with SHAP Analysis," *J. Digit. Mark. Digit. Curr.*, vol. 3, no. 3, pp. 195-210, 2026.

developed to support this process, aiming to reduce default risk while maintaining fairness and efficiency [1].

Among these approaches, ensemble learning methods such as Random Forest, XGBoost, and LightGBM have gained significant attention due to their strong predictive performance and ability to handle complex, non-linear relationships in structured data. These models combine multiple learners to improve generalization and reduce overfitting, making them particularly suitable for financial datasets. In addition, several studies have demonstrated that ensemble methods outperform traditional single models in credit risk prediction tasks [2]. However, despite their high accuracy, these models are often considered black-box systems, which limits their transparency and raises concerns regarding trust and accountability in real-world financial applications.

Another important challenge in loan approval prediction is the presence of imbalanced datasets, where the number of approved and non-approved cases is not evenly distributed. This imbalance can lead to biased model predictions, where the majority class dominates the learning process and the minority class is often misclassified. To address this issue, various resampling techniques have been proposed, among which the Synthetic Minority Over-sampling Technique (SMOTE) is widely used. SMOTE generates synthetic samples for the minority class, helping to balance the dataset and improve the model's ability to detect minority instances. Although many studies have applied SMOTE to improve classification performance, its integration with ensemble models in the context of explainable systems remains relatively underexplored [3].

In recent years, the concept of Explainable Artificial Intelligence (XAI) has emerged as a crucial component in machine learning applications, particularly in high-stakes domains such as finance. Techniques such as SHapley Additive exPlanations (SHAP) provide a unified framework to interpret model predictions by quantifying the contribution of each feature. SHAP enables both global and local interpretability, allowing stakeholders to understand not only which features are important but also how they influence individual predictions. Several studies have applied SHAP in credit risk analysis; however, most of them focus on a single model or do not provide a comparative analysis of feature importance across multiple ensemble methods [4].

Based on these limitations, this study aims to fill the existing research gap by proposing an explainable ensemble learning framework for loan approval prediction that integrates Random Forest, XGBoost, and LightGBM with SMOTE and SHAP analysis. The novelty of this study lies in the combined evaluation of model performance and interpretability, as well as the comparative analysis of feature importance across different ensemble models. By addressing both predictive accuracy and

transparency, this research contributes to the development of more reliable and trustworthy decision support systems in the financial domain [5].

Literature Review

Machine learning techniques have been widely applied in loan approval and credit risk prediction to improve decision-making processes and reduce financial risk. Various classification algorithms such as logistic regression, decision trees, and support vector machines have been utilized in earlier studies. However, these traditional models often struggle to capture complex and non-linear relationships within financial data, which can lead to suboptimal predictive performance in real-world applications [6]. In addition, their limited ability to generalize across diverse datasets further motivates the need for more advanced approaches [7].

To address these limitations, ensemble learning methods have been increasingly adopted due to their superior predictive performance and robustness. Techniques such as Random Forest, XGBoost, and LightGBM combine multiple weak learners to improve generalization and reduce overfitting. Previous studies have shown that ensemble models outperform single classifiers in credit scoring and loan approval tasks [8], [9]. Furthermore, boosting-based methods such as XGBoost and LightGBM are particularly effective in handling structured data and optimizing model performance through iterative learning processes [10]. These methods have also demonstrated strong scalability when applied to large datasets [11].

Another significant challenge in loan approval prediction is the presence of imbalanced datasets, where the number of approved and non-approved cases is not equally distributed. This imbalance often leads to biased predictions that favor the majority class, reducing the model's ability to correctly classify minority instances. To mitigate this issue, resampling techniques such as the Synthetic Minority Over-sampling Technique have been widely applied [12]. SMOTE generates synthetic samples for the minority class, thereby improving class balance and enhancing model performance [13]. Several studies have reported that the use of SMOTE significantly improves recall and F1-score in imbalanced classification problems [14].

Despite the strong predictive performance of ensemble models, their lack of interpretability remains a major limitation, especially in high-stakes domains such as finance. This has led to the growing interest in Explainable Artificial Intelligence, which aims to make machine learning models more transparent and understandable. Among various techniques, SHapley Additive exPlanations has gained popularity due to its solid theoretical foundation and ability to provide consistent feature attribution [15]. SHAP enables both global and local interpretability,

allowing users to understand overall model behavior as well as individual predictions [16]. This capability is particularly important in financial applications where accountability and transparency are required [17].

Several studies have applied SHAP in credit risk analysis to interpret model predictions and identify key influencing factors. These studies consistently highlight the importance of features such as credit score, income, and employment status in determining loan outcomes [18]. However, most existing research focuses on a single model and does not provide a comparative analysis across multiple ensemble methods [19]. In addition, the integration of ensemble learning, data balancing techniques, and explainability methods within a unified framework remains relatively limited [20]. Therefore, this study aims to address this gap by proposing an integrated approach that combines multiple ensemble models with SMOTE and SHAP analysis to evaluate both predictive performance and model interpretability.

Method

Research Framework

This study follows a structured research framework consisting of several main stages, including data collection, preprocessing, data balancing, model development, evaluation, and explainability analysis. The overall research workflow is illustrated in Figure 1, which provides a systematic overview of the proposed approach. The process begins with dataset preparation, followed by data preprocessing to handle missing values and encode categorical variables. The next stage involves handling class imbalance using SMOTE, after which three ensemble learning models are trained and evaluated. Finally, SHAP is applied to interpret the model and analyze feature contributions.

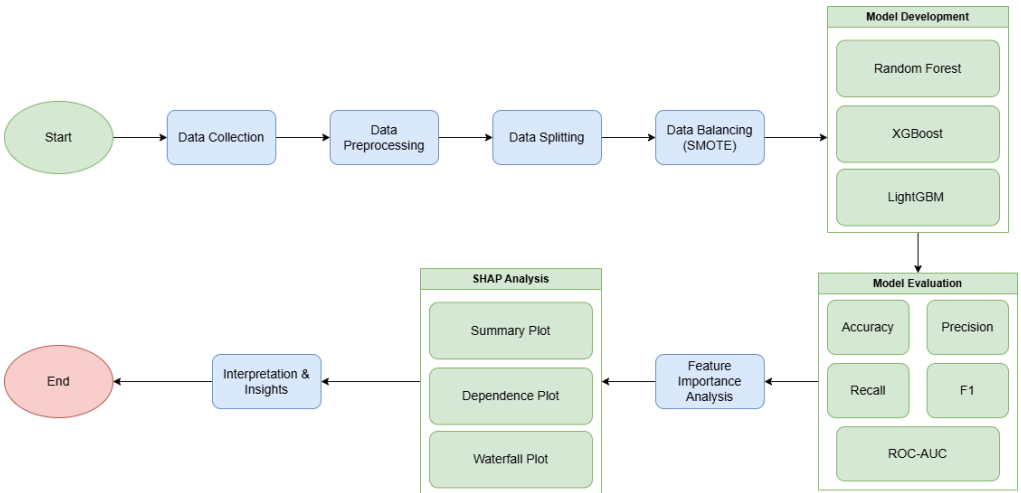


Figure 1 Research Method Flowchart

Dataset Description

This study utilizes a loan approval dataset consisting of 5,000 records and 10 variables, including demographic, financial, and employment-related features. The target variable is LoanApproved, which represents a binary classification problem where:

$$y \in \{0,1\} \quad (1)$$

0 indicates non-approved loans and 1 indicates approved loans. The input features include Age, Income, LoanAmount, CreditScore, YearsExperience, Gender, Education, City, and EmploymentType.

Data Preprocessing

Missing values in numerical features were handled using median imputation:

$$x_i = \text{median}(X) \quad (2)$$

while categorical features were imputed using the mode:

$$x_i = \text{mode}(X) \quad (3)$$

Categorical variables were transformed using label encoding, where each category is mapped into a numerical value:

$$x' = \text{LabelEncode}(x) \quad (4)$$

Data Splitting

The dataset was divided into training and testing sets using an 80:20 ratio:

$$X = X_{train} \cup X_{test} \quad (5)$$

Stratified sampling was applied to preserve class distribution:

$$P(y_{train}) \approx P(y_{test}) \quad (6)$$

Handling Imbalanced Data using SMOTE

To address class imbalance, SMOTE was applied to the training data. SMOTE generates synthetic samples using interpolation between minority class instances:

$$x_{new} = x_i + \lambda(x_{nn} - x_i) \quad (7)$$

x_i is a minority sample; x_{nn} is its nearest neighbor; $\lambda \in [0,1]$

This process increases the representation of the minority class without duplication.

Model Development

Three ensemble models were implemented:

Random Forest

$$\hat{y} = \frac{1}{N} \sum_{i=1}^N T_i(x) \quad (8)$$

XGBoost / LightGBM (Boosting concept)

$$F_m(x) = F_{m-1}(x) + \gamma_m h_m(x) \quad (9)$$

$h_m(x)$ is the weak learner and γ_m is the learning rate.

Model Evaluation

The models were evaluated using the following metrics:

Accuracy

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (10)$$

Precision

$$\text{Precision} = \frac{TP}{TP + FP} \quad (11)$$

Recall

$$\text{Recall} = \frac{TP}{TP + FN} \quad (12)$$

F1-score

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (13)$$

ROC-AUC

$$AUC = \int_0^1 TPR(FPR^{-1}(x)) dx \quad (14)$$

Explainability using SHAP

SHAP values are computed based on Shapley values from cooperative game theory:

$$\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|! (|F| - |S| - 1)!}{|F|!} [f(S \cup \{i\}) - f(S)] \quad (15)$$

ϕ_i is the contribution of feature i ; S is a subset of features; F is the full feature set.

SHAP provides both global and local interpretability by quantifying the contribution of each feature to the prediction.

Algorithm 1 Explainable Ensemble Learning for Loan Approval Prediction

Input: dataset $D = \{X, y\}$
Output: best model M^* , evaluation results, SHAP explanations
Process:
Start
Load dataset $D = \{X, y\}$
Data preprocessing
For each numerical feature $x_j \in X_{num}$:
 $x_j \leftarrow \text{median}(x_j)$
For each categorical feature $x_k \in X_{cat}$:
 $x_k \leftarrow \text{mode}(x_k)$
Encode categorical features:
 $X_{cat} \leftarrow \text{LabelEncode}(X_{cat})$
Data splitting
Split dataset into training and testing sets:
 $(X_{train}, X_{test}, y_{train}, y_{test}) \leftarrow \text{Split}(X, y, 0.8)$
Handle imbalance
Apply SMOTE on training data:
 $X'_{train}, y'_{train} \leftarrow \text{SMOTE}(X_{train}, y_{train})$
Model initialization
Initialize models:
 $M_1 \leftarrow \text{Random Forest}$
 $M_2 \leftarrow \text{XGBoost}$
 $M_3 \leftarrow \text{LightGBM}$
Training and evaluation
For each model $M_i \in \{M_1, M_2, M_3\}$:
Train model:
 $M_i \leftarrow \text{Train}(M_i, X'_{train}, y'_{train})$
Predict on test data:
 $\hat{y}_i \leftarrow M_i(X_{test})$
Compute evaluation metrics:
Accuracy, Precision, Recall, F1-score, ROC-AUC
End For
Model selection
Select best model based on AUC:
 $M^* = \arg \max(AUC_i)$
Feature importance
Compute feature importance from M^*
SHAP explainability
Select sample data $X_s \subseteq X_{test}$
Compute SHAP values:
 $\Phi \leftarrow \text{SHAP}(M^*, X'_{train}, X_s)$
Generate:
- SHAP summary plot
- SHAP dependence plot
- SHAP waterfall plot
End

Results

Class Distribution and Data Balancing

The dataset initially exhibited a notable class imbalance, where the non-approved class (Class 0) contained 3,079 instances, while the approved class (Class 1) consisted of only 921 instances. This uneven distribution indicates that the majority class significantly dominates the dataset, which can lead to biased model learning. In such conditions, machine learning models tend to favor the majority class during training, resulting in reduced sensitivity toward the minority class. As a consequence, the model may achieve high overall accuracy while failing to correctly identify loan approvals, which are critical in this context. Therefore, addressing class imbalance is essential to ensure fair and reliable model performance, particularly in financial decision-making systems.

To mitigate this issue, the Synthetic Minority Over-sampling Technique (SMOTE) was applied to the training dataset. SMOTE generates synthetic samples of the minority class by interpolating between existing minority instances, thereby increasing its representation without simply duplicating data. After applying SMOTE, both classes were balanced with 3,079 instances each, creating a more uniform class distribution. This balanced dataset allows the models to learn patterns from both classes more effectively and improves their ability to detect minority class instances. The comparison of class distributions before and after applying SMOTE is illustrated in figure 2 and figure 3, respectively.

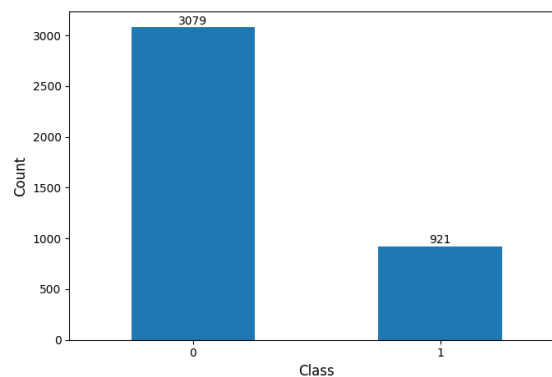


Figure 2 Class distribution before applying SMOTE

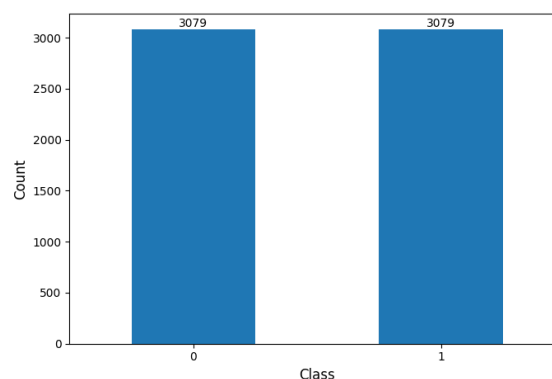


Figure 3 Class distribution after applying SMOTE

Model Performance Evaluation

The performance of the three ensemble models, namely Random Forest, XGBoost, and LightGBM, was evaluated using multiple classification metrics, including Accuracy, Precision, Recall, F1-score, and ROC-AUC, to provide a comprehensive assessment of their predictive capabilities. The results, as presented in Table 1, indicate that all models achieved strong performance, with accuracy values above 0.94 and ROC-AUC scores exceeding 0.92, demonstrating their effectiveness in distinguishing between approved and non-approved loan applications. Among the evaluated models, Random Forest consistently outperformed the others across all metrics, achieving the highest accuracy, precision, and F1-score, as well as the highest ROC-AUC value of 0.9386. This superior performance suggests that Random Forest is more effective in capturing complex patterns within the dataset and provides better generalization in classification tasks. Additionally, its relatively high recall for the minority class indicates improved capability in correctly identifying approved loan cases, which is critical in imbalanced classification problems.

Table 1 Performance comparison of ensemble models					
Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Random Forest	0.959	0.932	0.887	0.909	0.9386
LightGBM	0.956	0.919	0.887	0.903	0.9273
XGBoost	0.948	0.897	0.874	0.885	0.9257

ROC Curve Analysis

The Receiver Operating Characteristic (ROC) curves for all models are presented in Figure 4, providing a visual comparison of their classification performance across different threshold values. The curves for all three models are positioned close to the top-left corner of the plot, indicating high true positive rates and low false positive rates, which reflect strong discriminative ability. Among the evaluated models, Random Forest achieved the highest area under the curve with an AUC value of 0.9386, outperforming both XGBoost and LightGBM. This result confirms that Random Forest is more effective in distinguishing between approved and non-approved loan applications across various decision thresholds. Furthermore, the relatively small gap between the ROC curves suggests that all models are robust and capable of handling the classification task effectively, although Random Forest maintains a consistent advantage in overall predictive performance.

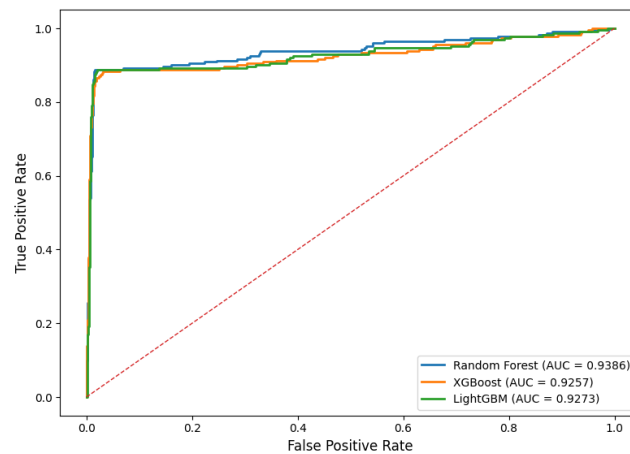


Figure 4 ROC curve comparison of ensemble models

Feature Importance Analysis

The feature importance derived from the Random Forest model is illustrated in Figure 5, highlighting the relative contribution of each feature to the prediction of loan approval. The results show that CreditScore is the most influential feature, with a substantially higher importance value compared to the other variables, indicating that it plays a dominant role in the model's decision-making process. This is followed by EmploymentType and Income, which also exhibit significant contributions, suggesting that employment stability and financial capacity are key determinants in assessing loan eligibility. In contrast, other features such as LoanAmount, Age, YearsExperience, City, Gender, and Education have relatively lower importance values, indicating a smaller impact on the model's predictions. These findings suggest that the model prioritizes indicators of financial reliability and economic stability when making approval decisions, which aligns with common practices in credit risk assessment.

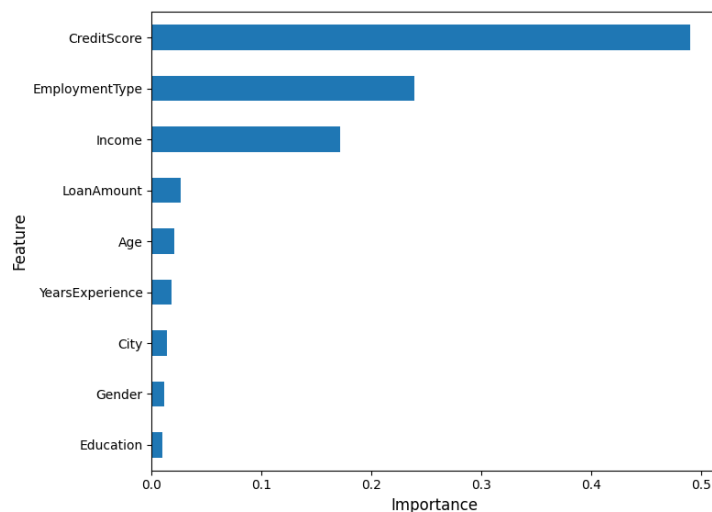


Figure 5 Feature importance based on Random Forest

SHAP-Based Explainability Analysis

To enhance model interpretability, SHAP (SHapley Additive exPlanations) was applied to the best-performing model, namely Random Forest, in order to provide a comprehensive understanding of how individual features contribute to the prediction outcomes. Unlike traditional feature importance methods, SHAP is based on cooperative game theory and assigns each feature a contribution value that reflects its impact on the model's output for both global and local interpretations. This approach allows for a more transparent analysis of the model by quantifying not only which features are important but also how they influence predictions in positive or negative directions. By applying SHAP, the study is able to uncover the underlying decision-making process of the model, thereby improving trust and reliability in the prediction results, which is particularly important in high-stakes domains such as financial decision-making.

The SHAP summary plot presented in Figure 6 illustrates the overall impact of each feature on the model output by combining both feature importance and the direction of influence across all samples. The plot ranks features based on their average absolute SHAP values, indicating their global significance in the prediction process. The results confirm that CreditScore, EmploymentType, and Income are the most influential features, which is consistent with the findings from the Random Forest feature importance analysis. In addition, the distribution of SHAP values within the plot reveals how variations in feature values affect the prediction outcome, where higher values of CreditScore and Income generally contribute positively toward loan approval, while lower values tend to decrease the likelihood of approval. This consistency between SHAP analysis and traditional feature importance strengthens the reliability of the model interpretation and provides deeper insight into how key financial and employment-related factors drive the decision-making process.

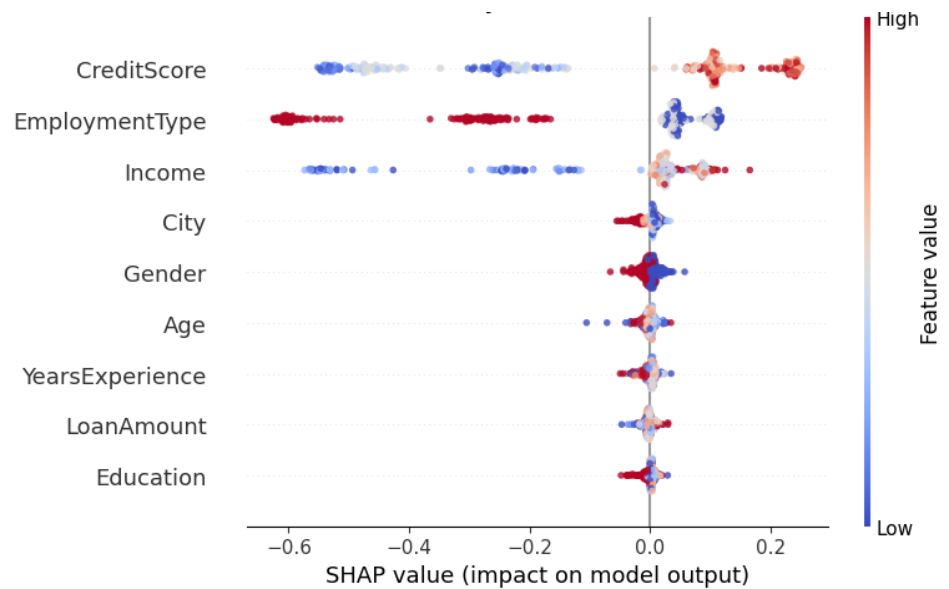


Figure 6 SHAP summary plot

The SHAP dependence plot for CreditScore, as presented in Figure 7, provides a detailed visualization of how variations in CreditScore influence the model output at the individual sample level. The plot demonstrates a non-linear relationship, where increases in CreditScore generally lead to higher SHAP values, indicating a stronger contribution toward loan approval, although the rate of increase is not uniform across all ranges. This suggests that the effect of CreditScore on the prediction is more complex than a simple linear relationship. In addition, the color gradient representing EmploymentType reveals interaction effects between these two features, indicating that the impact of CreditScore on the prediction varies depending on the applicant's employment status. For example, individuals with similar credit scores may receive different prediction outcomes based on their employment type, highlighting the importance of feature interactions in the model. This analysis confirms that the model captures not only individual feature contributions but also the combined effects of multiple features in determining loan approval decisions.

Discussion

The results of this study demonstrate that ensemble learning models are highly effective for loan approval prediction, particularly when combined with appropriate data preprocessing techniques. All evaluated models achieved strong performance, with ROC-AUC values exceeding 0.92, indicating a high ability to distinguish between approved and non-approved loan applications. Among the models, Random Forest consistently delivered the best performance across all evaluation metrics, suggesting its superior capability in capturing complex and non-linear relationships within structured financial data. This performance

advantage can be attributed to its ensemble mechanism, which aggregates multiple decision trees to reduce variance and improve generalization [21], [22]. In addition, the application of SMOTE significantly contributed to performance improvement by addressing class imbalance in the dataset. By generating synthetic samples for the minority class, the models were able to better learn its underlying patterns, leading to improved recall and more balanced classification outcomes [23], [24]. This highlights the importance of incorporating data balancing techniques in classification tasks involving skewed distributions, especially in financial decision-making contexts [23], [25].

The interpretability analysis further strengthens the reliability of the proposed approach by providing both global and local insights into the model's decision-making process. Feature importance and SHAP analysis consistently identified CreditScore, EmploymentType, and Income as the most influential factors, aligning with domain knowledge in credit risk assessment [26], [27]. The SHAP summary plot revealed not only the relative importance of these features but also their directional impact on predictions, while the dependence plot highlighted the presence of non-linear relationships and interaction effects between features, particularly between CreditScore and EmploymentType. Moreover, the local explanation using the SHAP waterfall plot demonstrated how individual features contribute to specific predictions, offering transparency at the instance level [26]. This level of interpretability is essential in financial applications, where decisions must be explainable and justifiable to stakeholders [27]. Overall, the integration of high predictive performance with explainable AI techniques provides a robust framework for developing transparent and trustworthy loan approval systems [21], [26], [27].

Conclusion

This study demonstrates that the integration of ensemble learning and explainable artificial intelligence provides an effective and reliable approach for loan approval prediction. Among the evaluated models, Random Forest consistently achieved the best performance across all evaluation metrics, including Accuracy, F1-score, and ROC-AUC, indicating its strong capability in capturing complex patterns within the dataset and delivering robust classification results. The application of SMOTE played a critical role in addressing class imbalance by generating synthetic samples for the minority class, which improved the model's ability to correctly identify approved loan applications and ensured a more balanced predictive performance. In addition to high predictive accuracy, the use of SHAP enabled a comprehensive interpretation of the model by identifying CreditScore, EmploymentType, and Income as the most influential features, while also providing insights into both the magnitude and direction of their contributions. Furthermore, SHAP analysis allowed for both global and local explanations, offering transparency into how

individual predictions are formed and revealing interaction effects between features. This combination of strong predictive performance and interpretability highlights the potential of the proposed approach to support transparent, data-driven decision-making in financial systems, where both accuracy and explainability are essential for building trust and ensuring accountability.

Declarations

Author Contributions

Conceptualization: M.G., V.K.; Methodology: M.G., C.R.A.; Software: V.K.; Validation: C.R.A., M.G.; Formal Analysis: M.G.; Investigation: V.K., C.R.A.; Resources: M.G., V.K.; Data Curation: C.R.A.; Writing – Original Draft Preparation: M.G.; Writing – Review and Editing: C.R.A., V.K.; Visualization: V.K.; All authors have read and agreed to the published version of the manuscript.

Data Availability Statement

The data presented in this study are available on request from the corresponding author.

Funding

The authors received no financial support for the research, authorship, and/or publication of this article.

Institutional Review Board Statement

Not applicable.

Informed Consent Statement

Not applicable.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

References

- [1] T. Hastie, R. Tibshirani, and J. Friedman, "The elements of statistical learning: Data mining, inference, and prediction," *Springer*, vol. 2009, no. Jan., pp. 1–745, Jan. 2009.
- [2] L. Breiman, "Random forests," *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, Oct. 2001, doi: 10.1023/A:1010933404324.
- [3] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *J. Artif. Intell. Res.*, vol. 16, no. 2002, pp. 321–357, Jun. 2002, doi: 10.1613/jair.953.
- [4] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," *NeurIPS*, vol. 2017, no. Dec., pp. 4765–4774, Dec. 2017, doi: 10.48550/arXiv.1705.07874.

- [5] N. Bussmann, P. Giudici, D. Marinelli, and J. Papenbrock, "Explainable machine learning in credit risk management," *Comput. Econ.*, vol. 57, no. 1, pp. 203–216, Jan. 2021, doi: 10.1007/s10614-020-10042-0.
- [6] A. Tasnim, M. Saiduzzaman, M. A. Rahman, J. Akhter, and A. S. M. M. Rahaman, "Performance evaluation of multiple classifiers for predicting fake news," *J. Comput. Commun.*, vol. 10, no. 9, pp. 1–15, Sep. 2022.
- [7] I. Goodfellow, Y. Bengio, and A. Courville, "Deep learning," *MIT Press*, vol. 2016, no. Jan., pp. 1–800, Jan. 2016.
- [8] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," *KDD*, vol. 2016, no. Aug., pp. 785–794, Aug. 2016, doi: 10.1145/2939672.2939785.
- [9] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T.-Y. Liu, "LightGBM: A highly efficient gradient boosting decision tree," *NeurIPS*, vol. 2017, no. Dec., pp. 3149–3157, Dec. 2017.
- [10] J. H. Friedman, "Greedy function approximation: A gradient boosting machine," *Ann. Statist.*, vol. 29, no. 5, pp. 1189–1232, Nov. 2001, doi: 10.1214/aos/1013203451.
- [11] L. Yu, S. Wang, and K. K. Lai, "Credit risk assessment with a multistage neural network ensemble learning approach," *Expert Syst. Appl.*, vol. 34, no. 2, pp. 1434–1444, Feb. 2008, doi: 10.1016/j.eswa.2007.01.009.
- [12] G. E. Batista, R. C. Prati, and M.-C. Monard, "A study of the behavior of several methods for balancing machine learning training data," *ACM SIGKDD Explor. Newsl.*, vol. 6, no. 1, pp. 20–29, Jun. 2004, doi: 10.1145/1007730.1007735.
- [13] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Trans. Knowl. Data Eng.*, vol. 21, no. 9, pp. 1263–1284, Sep. 2009, doi: 10.1109/TKDE.2008.239.
- [14] M. Buda, A. Maki, and M. A. Mazurowski, "A systematic study of the class imbalance problem in convolutional neural networks," *Neural Netw.*, vol. 106, no. 2018, pp. 249–259, Oct. 2018, doi: 10.1016/j.neunet.2018.07.011.
- [15] S. M. Lundberg, G. Erion, H. Chen, A. DeGrave, J. M. Prutkin, B. Nair, R. Katz, J. Himmelfarb, N. Bansal, and S.-I. Lee, "From local explanations to global understanding with explainable AI for trees," *Nat. Mach. Intell.*, vol. 2, no. 1, pp. 56–67, Jan. 2020, doi: 10.1038/s42256-019-0138-9.
- [16] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you?: Explaining the predictions of any classifier," *KDD*, vol. 2016, no. Aug., pp. 1135–1144, Aug. 2016, doi: 10.1145/2939672.2939778.
- [17] B. Goodman and S. Flaxman, "European Union regulations on algorithmic decision-making and a 'right to explanation'," *AI Mag.*, vol. 38, no. 3, pp. 50–57, Sep. 2017, doi: 10.1609/aimag.v38i3.2741.
- [18] S. Lessmann, B. Baesens, H.-V. Seow, and L. C. Thomas, "Benchmarking state-of-the-art classification algorithms for credit scoring: An update of research," *Eur. J. Oper. Res.*, vol. 247, no. 1, pp. 124–136, Oct. 2015, doi: 10.1016/j.ejor.2015.05.030.
- [19] D. Martens, B. Baesens, T. Van Gestel, and J. Vanthienen, "Comprehensible credit scoring models using rule extraction from support vector machines," *Eur. J. Oper. Res.*, vol. 183, no. 3, pp. 1466–1476, Dec. 2007, doi: 10.1016/j.ejor.2006.04.051.
- [20] S. Maldonado, J. López, and C. Vairetti, "An alternative SMOTE oversampling strategy for high-dimensional datasets," *Appl. Soft Comput.*, vol. 76, no. 2019, pp. 380–389, Mar. 2019, doi: 10.1016/j.asoc.2018.12.024.

- [21] M. Ghahremani, H. Otieno, M. Kazreen, and S. Shiaeles, "Machine learning-based loan approval automation: Enhancing efficiency, accuracy and fairness in credit decision-making," *Expert Syst.*, vol. 43, no. 2026, pp. 1–15, 2026, doi: 10.1111/exsy.70307.
- [22] X. Cai, W. Dai, and J. Lu, "Loan default prediction based on machine learning approaches," *ICAIGIS*, vol. 2025, no. Aug., pp. 1–10, Aug. 2025, doi: 10.1145/3728725.3728813.
- [23] N. N. A. Nanda, Y. Farida, and W. D. Utami, "Implementation of SMOTE to improve the performance of Random Forest classification in credit risk assessment in banking," *INTENSIF J. Ilm. Penelit. Penerapan Teknol. Sist. Inf.*, vol. 9, no. 2, pp. 1–15, Aug. 2025, doi: 10.29407/intensif.v9i2.23930.
- [24] S. Xie and T. Shingadia, "Explainable machine learning framework for predicting auto loan defaults," *Risks*, vol. 13, no. 9, pp. 1–15, Sep. 2025, doi: 10.3390/risks13090172.
- [25] S. Shreya and H. Pathak, "Explainable artificial intelligence credit risk assessment using machine learning," *arXiv*, vol. 2025, no. Jun., pp. 1–15, Jun. 2025, doi: 10.48550/arXiv.2506.19383.
- [26] S. Yang, Z. Huang, W. Xiao, and X. Shen, "Interpretable credit default prediction with ensemble learning and SHAP," *ICAHN*, vol. 2025, no. Aug., pp. 102–106, Aug. 2025, doi: 10.1109/ICAHN67688.2025.00027.
- [27] M. Rambauli, T. Ravele, and C. Sigauke, "Application of explainable AI and uncertainty quantification in credit risk assessment," *Risks*, vol. 14, no. 7, pp. 1–15, Jul. 2026, doi: 10.3390/risks14070152.