

Comparative Analysis of Explainable AI Frameworks (LIME and SHAP) in Loan Approval Systems

Isaac Terngu Adom*

Department of Mathematics and Computer Science, Benue State University, Nigeria

E-mail: iadom@bsum.edu.ng

ORCID iD: <https://orcid.org/0000-0003-1389-1255>

*Corresponding Author

Christiana O. Julius

Department of Mathematics and Computer Science, Benue State University, Nigeria

E-mail: juliuschristiana96@gmail.com

ORCID iD: <https://orcid.org/0009-0002-5903-8559>

Stephen Akuma

Department of Mathematics and Computer Science, Benue State University, Nigeria

E-mail: sakuma@bsum.edu.ng

ORCID iD: <https://orcid.org/0000-0003-1909-7618>

Samera U. Otor

Department of Mathematics and Computer Science, Benue State University, Nigeria

E-mail: sotor@bsum.edu.ng

ORCID iD: <http://orcid.org/0000-0001-6943-7137>

Received: 22 November, 2024; Revised: 05 February, 2025; Accepted: 23 May, 2025; Published: 08 December, 2025

Abstract: Machine learning models that lack transparency can lead to biased conclusions and decisions in automated systems in various domains. To address this issue, explainable AI (XAI) frameworks such as Local Interpretable Model-Agnostic Explanations (LIME) and Shapley Additive Explanations (SHAP) have evolved by offering interpretable insights into machine learning model decisions. A thorough comparison of LIME and SHAP applied to a Random Forest model trained on a loan dataset resulted in an Accuracy of 85%, Precision of 84%, Recall of 97%, and an F1 score of 90%, is presented in this study. This study's primary contributions are as follows: (1) using Shapley values, which represent the contribution of each feature, to show that SHAP provides deeper and more reliable feature attributions than LIME; (2) demonstrating that LIME lacks the sophisticated interpretability of SHAP, despite offering faster and more generalizable explanations across various model types; (3) quantitatively comparing computational efficiency, where LIME displays a faster runtime of 0.1486 seconds using 9.14MB of memory compared to SHAP with a computational time of 0.3784 seconds using memory 1.2 MB. By highlighting the trade-offs between LIME and SHAP in terms of interpretability, computational complexity, and application to various computer systems, this study contributes to the field of XAI. The outcome helps stakeholders better understand and trust AI-driven loan choices, which advances the development of transparent and responsible AI systems in finance.

Index Terms: Explainable AI, Interpretability, LIME, loan approval, ML, SHAP

1. Introduction

Technology is gradually taking over several walks of life, with Artificial Intelligence (AI) taking the lead in various domains. We find AI applicable in health, finance, education, business, and defense, among others. Today, people can utilize AI in ways believed to be impossible years ago. Loan approval systems are a strategic aspect of the financial industry, enabling individuals and businesses to access funding for various purposes. The use of machine learning algorithms in loan approval systems has become increasingly prevalent due to the growth of big data and the

increasing complexity of financial transactions. However, one of the significant challenges with the use of these algorithms is their lack of transparency and interpretability, which can lead to biased or unfair loan approval decisions. To address this problem, explainable AI (XAI) frameworks have been developed to provide interpretable explanations for machine learning models [1]. The problem of interpretability and understanding how these systems work is called the “black box problem” in loan approval systems. The black box problem refers to the lack of transparency and interpretability in the prediction process of complex machine-learning models [2]. XAI frameworks enable users to understand how the model arrived at its decisions, making it possible to detect and address biases, errors, or inconsistencies in the decision-making process [3]. XAI frameworks for machine learning models such as Local Interpretable Model-Agnostic Explanations (LIME) and Shapley Additive Explanations (SHAP) have been used in XAI tasks with distinctive capabilities. LIME provides model-agnostic explanations by approximating a complex model with a simpler one that can be more easily explained. SHAP, on the other hand, assigns values to each feature's contribution to the model's output, providing a clear and interpretable way to understand the model's decision-making process. While these XAI frameworks have been applied in various domains, including healthcare, marketing, and finance, their effectiveness in loan approval systems remains under-explored. Thus, there is a need to evaluate the performance of these frameworks in loan approval systems and compare their effectiveness in providing transparent and interpretable explanations [4]. This study seeks to fill this gap by performing a comparative analysis of LIME and SHAP in loan approval systems to determine their strengths and limitations answering the following research questions: (1). How do LIME and SHAP perform in terms of explaining loan approval decisions made by a Random Forest model? (2). Which explainability framework (LIME or SHAP) provides better interpretability in loan approval systems? (3). How do LIME and SHAP compare in computational efficiency when applied to loan datasets? Overall, the study contributes to developing more trustworthy and transparent AI systems in the financial sector, promoting fairness and accountability in loan approval processes. The rest of the paper is organized into the following sections: Section 2 presents related work to the research. Section 3 covers the methodology of our study. The implementation experiments are carried out in Section 4, and Section 5 presents the conclusion and future research direction.

2. Related Works

The use of automated processes or systems in critical decision-making has evolved. Critical systems and services such as health care diagnosis, infrastructural systems, high-risk financial services, and several others have adopted systems and solutions that do not only rely on human efforts [5]. In financial critical systems like loan credit scoring and predicting loan approvals and rejections, artificial intelligence models have been used for such tasks. What remains a challenge is how transparent the system's results are to the concerned stakeholders [6]. This has given rise to Explainable AI which seeks to make the models more interpretable. Explainable AI is applicable in several fields today. In healthcare, it assists healthcare professionals in understanding and rationalizing the decisions made by AI systems when it comes to diagnosing ailments and suggesting treatment strategies [7]. Trust and safety in autonomous vehicles are strengthened by offering comprehensible explanations for the decisions made by the AI algorithms that govern these vehicles, thereby increasing acceptance and reducing uncertainty among passengers [8]. Explaining AI systems's decisions in the legal profession enables judges and lawyers to understand the factors contributing to AI-generated outcomes. This ensures fairness, accountability, and transparency in legal proceedings [9]. Interpretation approaches are useful in unraveling the complexities of understanding AI systems, especially in the generation of synthetic data and other generative processes enhancing transparency [10]. By offering interpretable explanations, AI systems enable the identification of suspicious patterns, clarification of risk classifications, and the establishment of transparency in regulatory reporting. Financial institutions can assess creditworthiness more effectively by providing explanations for loan underwriting and credit scoring decisions. This understanding of contributing factors and variables helps ensure fairness, reduce bias, and improve transparency in the loan approval process [11]. To help users understand why specific products or services are being recommended to them, recommender systems have used SHAP to provide explanations. This increases transparency and user trust [12]. There are several XAI frameworks, such as LIME, SHAP, ELI5, and so on, that have been used for interpretability tasks. Local Interpretable Model-Agnostic Explanations were first proposed in [4]. It is one of the major explainable AI frameworks due to its simplicity and compatibility with different machine learning models. [13] identified LIME as one of the XAI frameworks that make machine learning models interpretable and transparent. By providing explanations for individual predictions in text classification, sentiment analysis, and other text-based tasks, LIME is designed to handle both textual and tabular data, making it applicable for interpreting predictions in various domains like natural language processing (NLP) and structured data analysis [4]. By generating explanations in the form of features or rules that are easily understood by humans, LIME provides high interpretability, allowing users to gain an understanding of the model's decision-making process and develop confidence in the model's predictions [3]. Users can evaluate the impact of individual features on the model's output by using LIME's feature importance weights, which indicate the relative contribution of each feature to the prediction. Shapely Additive Explanations, on the other hand, was introduced in 2017 by Lundberg and Sun-In. Just like LIME, it also uses local explanations for predictions; however, its basic working principle is based on game theory. It functions by adding shapely values to individual features. SHAP, like any other explainable AI framework, has its unique features. Shapley values are a fair way of assigning importance to individual features, taking into account all

possible feature combinations, and SHAP uses them to determine the contribution of each feature in predictions in deep neural networks, random forests, and support vector machines, regardless of their architecture or type. In addition, it uses a fair attribution method that upholds fundamental properties like local accuracy and proper handling of missing values. SHAP ensures consistency in feature attributions, ensuring that the sum of feature importance matches the discrepancy between the expected and actual outputs for a given instance [14]. Several studies have compared the performance of LIME and SHAP in explaining the predictions made by machine learning models. These studies have shown that both techniques are effective in generating explanations, but there are some differences in their performance. [15] evaluated the performance of LIME and SHAP in explaining the predictions made by a variety of machine learning models on different datasets. The authors reported that both techniques were effective in generating explanations, but SHAP was more consistent and provided more accurate explanations than LIME. They concluded that soon, the focus will be more on analyzing models than data, and interpretability will propel machine learning. [16] used the explainable frameworks of LIME and SHAP for retinoblastoma diagnosis with the interpretation of deep learning models. They were able to show that with the help of LIME and SHAP, input image features have the highest impact on the model's prediction. Despite the limited dataset, their approach yielded promising results. [17] used explainable AI for the classification of multivariate time series in the maritime sector. They classified ship types and provided explanations corresponding to different time intervals. [18] analyzed machine learning models for credit scoring with explainable AI by using credit scoring systems to determine how credit risk analysis was analyzed. The study concluded that LIME and SHAP were very suitable for explaining black box classifiers using deep neural networks to increase transparency and automation in machine learning. The interpretability of a machine learning-based credit scoring model can be improved using surrogate models and post-hoc explanations. The credit risk estimation is based on a probability estimate of the likelihood that a borrower or debtor will default on their obligations. This probability serves as a measure of the risk associated with lending or investing, and it is desirable to have it as close as possible to the true or effective level of risk. An accurate evaluation of the probability of default enables better decision-making and risk-management strategies [19]. [20] seeks to assess the predictive capacity of several ML models in the context of P2P lending platforms' credit scoring, after applying the Shapley method to provide explainability to the prediction. As machine learning models become more complex and powerful, the need for explainability and interpretability will only increase. LIME and SHAP are two of the most widely used techniques for generating explanations for machine learning models, and they are constantly being improved and updated. We can expect to see more studies that compare the performance of LIME and SHAP on different types of models and datasets as the one here with varying techniques and approaches. [21] presented an approach to how class imbalance affects the interpretability of LIME and SHAP in credit scoring applications. In our study, the emphasis is on the loan approval system application of explainable AI usage which is not extensively addressed by past research.

3. Experimental Setup

The comparative analysis of LIME and SHAP as XAI frameworks for the loan approval system was carried out. The machine learning model used is Random Forest for the prediction of loan approval outcomes. LIME and SHAP are then applied to generate explanations for individual loan application predictions and evaluate the quality and interpretability of these explanations. The evaluation of LIME and SHAP is based on stability, consistency, and computational efficiency. Their performance is compared in terms of identifying the most influential features in loan approval decisions and assessing their consistency across different instances. Figure 1 shows the methodology of our system.

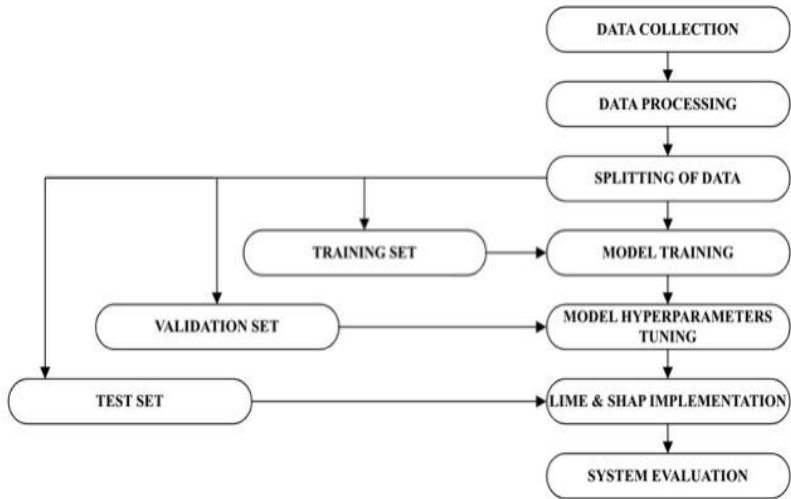


Fig. 1. Methodology of the system

3.1 Data Preprocessing

The loan dataset obtained from Kaggle has 613 columns with key features used in the study, including credit history, applicant income, dependents, loan ID (primary key of the dataset), gender, loan amount, loan term, self-employed, married, co-applicant income, and education. The preprocessing of the dataset involves handling missing values, removing duplicates, addressing outliers, and standardizing data formats.

3.2 Model Building

Random Forest was used for building the machine learning model because it reduces the risk of overfitting by aggregating predictions from multiple decision trees. It also effectively handles datasets with a large number of features (high-dimensional data) without requiring dimensionality reduction techniques. Random Forest handles imbalanced datasets well, as they are not heavily influenced by the majority class, can capture patterns in the minority class effectively, and provide insights into feature importance. Apart from the Random Forest model having a clearer way of measuring the features' importance and contribution by reducing uncertainty in making decisions, it is used for explainability tasks because, unlike other models, it provides global interpretability and not just individual predictions. Moreso, Random Forest models typically have higher precision due to their ability to handle non-linearity and interactions between features. Logistic Regression, being linear, on the other hand, may struggle in cases where the decision boundary is complex, which can lower its precision. Support Vector Machine can perform similarly to Random Forest, especially in well-separated classes, but may fall behind when feature interactions become complex. For accuracy, Random Forest performs better because they use multiple decision trees. Random Forest for loan approval systems gives a higher recall, accuracy, and precision when compared with Logistic Regression, SVM achieves precision similar to Random Forest when tuned properly but may require more computational efforts. For the F1 score, Random Forest achieves a high F1 score because of its ability to balance precision and recall compared with other models. The loan dataset is split into training, validation, and testing sets. The training set is used to train machine learning models; the validation set helps tune model hyperparameters; and the testing set is used to evaluate and compare LIME and SHAP explanations.

3.3 LIME Implementation Setup

A loan application instance is selected from the testing set to implement LIME. LIME is used to generate local explanations for the chosen loan application. LIME approximates the behavior of the Random Forest model around the instance of interest by generating a local interpretable model. Next, the features of the loan application are perturbed, and predictions are obtained from the Random Forest model for these perturbed instances. This is followed by fitting a local interpretable model to the perturbed instances and their corresponding model predictions, after which the feature importance values are calculated from the coefficients of the interpretable model. 100 perturbed instances were generated to model the local decision boundary. A linear surrogate model was used with a kernel weight of 0.75. Finally, an explanation highlighting the features and their contributions to the loan approval decision is generated for the selected loan application.

3.4 SHAP Implementation Setup

To implement SHAP, the loan application instance used is selected from the same testing set. SHAP is then used to generate explanations for the chosen loan application. TreeSHAP variant was used with 100 background samples randomly selected from the training set with pairwise interaction effects between features. To achieve this, SHAP assigns each feature in the loan application a Shapley value, which represents the average contribution of that feature across all possible feature combinations. Next, the Shapley values for the loan application using the Random Forest model were computed as the machine learning model. An explanation that shows the feature contributions and their impact on the loan approval decision is then generated for the selected loan application. Unlike LIME, SHAP takes into consideration all the loan instances and does not use a single instance to generate an explanation for the prediction.

4. Results & Discussion

The primary objective of this comparative analysis is to examine how LIME and SHAP contribute to enhancing the transparency and interpretability of loan approval systems, specifically when combined with the Random Forest machine learning model. The analysis involved implementing both frameworks, generating explanations for loan approval decisions, and evaluating the quality and utility of the explanations. The comparison was conducted using a representative loan dataset obtained from Kaggle with an extract dataset shown in Figure 2.

From the experiment, the system recorded an accuracy of 85%, precision-84%, recall-97%, and an F1 score of 90% for the Random Forest model. A high precision of 84% for this system indicates a low false positive rate, which implies that the risk of granting loans to undeserving borrowers is greatly reduced and there are fewer situations where a customer who ought to be rejected for a loan gets approved. Recall, also known as true positive rate, measures the degree of correctly predicted approved loans out of all actual approved loans. In loan approval systems, recall portrays the model's ability to identify all creditworthy borrowers and avoid false negatives, ensuring that only deserving applicants are approved and not overlooked. A recall of 97% represents a low false negative meaning that the chances

	Loan_ID	Gender	Married	Dependents	Education	Self_Employed	ApplicantIncome	CoapplicantIncome	LoanAmount	Loan_Amount_Term	Credit_History	Property_Area
1	LP001003	Male	Yes	1	Graduate	No	4583	1508.0	128.0	360.0	1.0	Ru
2	LP001005	Male	Yes	0	Graduate	Yes	3000	0.0	66.0	360.0	1.0	Urt
3	LP001006	Male	Yes	0	Not Graduate	No	2583	2358.0	120.0	360.0	1.0	Urt
4	LP001008	Male	No	0	Graduate	No	6000	0.0	141.0	360.0	1.0	Urt
5	LP001011	Male	Yes	2	Graduate	Yes	5417	4196.0	267.0	360.0	1.0	Urt
...	...	...	...	...	...	...	...	...	...	...	...	...
609	LP002978	Female	No	0	Graduate	No	2900	0.0	71.0	360.0	1.0	Ru
610	LP002979	Male	Yes	3+	Graduate	No	4106	0.0	40.0	180.0	1.0	Ru
611	LP002983	Male	Yes	1	Graduate	No	8072	240.0	253.0	360.0	1.0	Urt
612	LP002984	Male	Yes	2	Graduate	No	7583	0.0	187.0	360.0	1.0	Urt
613	LP002990	Female	No	0	Graduate	Yes	4583	0.0	133.0	360.0	0.0	Semiurt

Fig. 2. An extract from the dataset

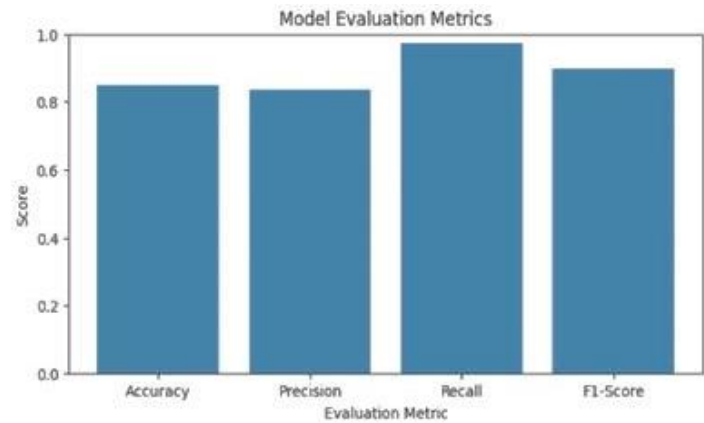


Fig. 3. Random Forest Evaluation Metrics

of trustworthy customers being rejected for loans are greatly reduced. F1 score is a combined metric that considers both precision and recall. It provides a balance between these two measures and is particularly useful when the dataset is imbalanced or when both false positives and false negatives need to be minimized. An F1 score of 90% shows the tradeoffs between recall and precision and indicates very low bias in the system explanations in which case the system has low false negatives and low false positives. Comparatively, in terms of interpretability and complexity, these evaluation metrics present LIME and SHAP as very effective explainable AI frameworks in loan approval systems in terms of the explanations and predictions generated.

4.1 LIME Explanations

By using LIME, insights can be gained into how the model arrived at a particular loan approval decision for a specific instance. This transparency helps in understanding the reasons behind the decision and detecting any potential biases or discriminatory patterns. LIME explanations can be used to validate and justify individual loan approval decisions by identifying the most influential features for individual loan instances. LIME is designed to be model-agnostic, meaning it can be applied to any machine learning model, including Random Forest, Logistic Regression, or Support Vector Machines (SVM). This flexibility makes LIME suitable for loan approval systems that employ different types of models or when the model used is unknown. It can provide explanations regardless of the underlying model's complexity. Figure 4 displays the LIME explanation for a loan instance highlighting its feature contribution and values. The probability of the loan being approved or not is shown to the left of the figure. In the middle, the most important parameters for the loan being approved (shown in orange) and not approved (shown in blue) are shown. The rightmost part shows the actual values of the parameters. In the context of LIME explanations, the prediction probability of 83% and 17% indicates the likelihood of the loan application being classified as either "NO" or "YES" by the underlying machine learning model. Also as shown in Figure 4, the parameters which increase the possibility of a loan being classified as 'yes' are 'property area', 'Gender', 'education', 'loan amount', and 'employment status' which are ordered by the weights. The features of 'loan amount term' (the period the loan is into effect), 'loan ID', 'marital status', 'number of dependents', 'applicant's income', and 'co-applicant income' decrease the possibility of the loan being approved. The prediction probabilities represent the probabilities assigned by the local model for each class label, "NO" and "YES" in this case. The given prediction probabilities suggest that based on the features and their contributions identified by LIME, there is a higher likelihood (83%) that the loan application will be classified as "NO" by the underlying model for the given individual given that the explanations are specific to individual instances and provide insights into the factors influencing the prediction outcome.

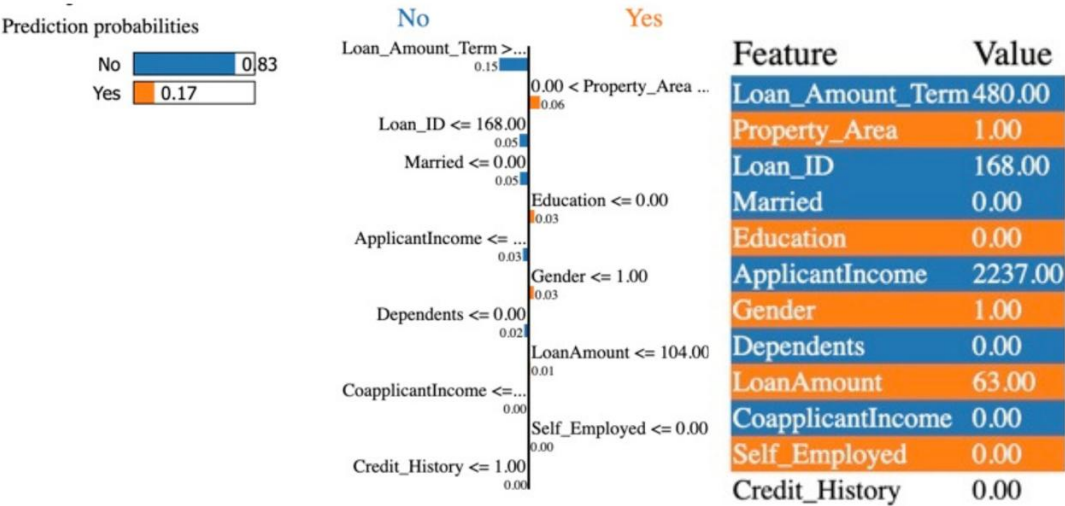


Fig. 4. A loan approval prediction explained by LIME

4.2 SHAP Explanations

SHAP provides a comprehensive understanding of feature importance, allowing stakeholders to evaluate the influence of different factors on loan approval decisions. This transparency ensures that the decision-making process is not solely based on arbitrary or hidden criteria, but rather on quantifiable and interpretable feature contributions. SHAP values enable stakeholders to assess the relative importance of different features in loan approval decisions by highlighting the relative importance of features at a global level. SHAP requires access to a specific model's internals. It utilizes game theory concepts to attribute the prediction outcome to each feature. This makes SHAP more suitable when the loan approval system uses a specific model, as in this case the Random Forest classifier. SHAP can capture the interactions between features specific to the chosen model, providing detailed insights into the decision-making process. SHAP explanations are provided in four dimensions namely; summary plot, force plots, dependence plot, and interaction plot.

4.2.1 Summary Plot Explanation

As shown in Figure 5, the y-axis indicates the variable name, in order of importance from top to bottom. The value next to it is the mean on the x-axis showing the SHAP value. For all the features and samples in the selected range, the plot aggregates SHAP values sorted in a way where the most important feature is at the top of the list. The important features of the dataset are shown with the help of a bar diagram, where the features are categorized according to their precedence with credit_history which tops the list having the highest precedence, and self_employed having the lowest precedence at the bottom of the list.

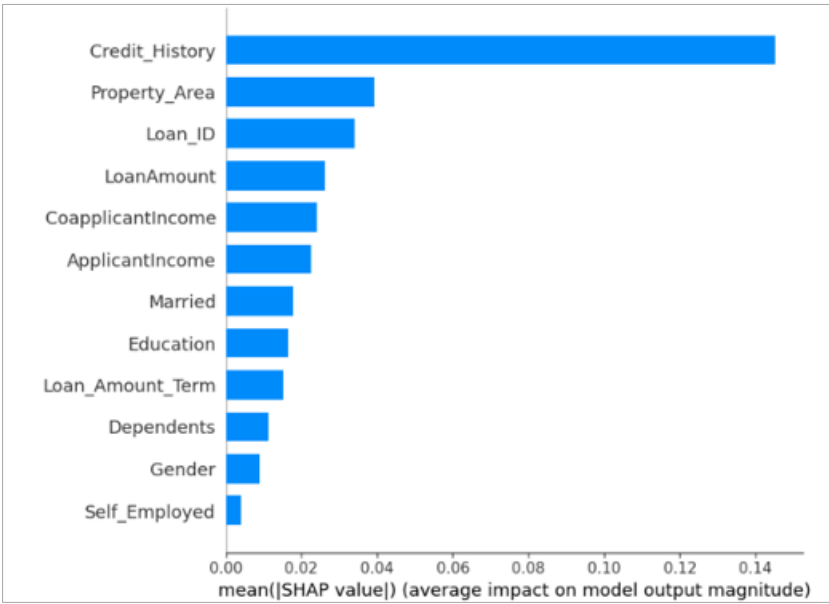


Fig. 5. Bar diagram showing important features of the loan approval dataset

4.2.2 Force Plot Explanation

The force plot as shown in Figure 6 indicates the model’s score of 0.74 for the observation. Higher scores lead the model to predict 1 (YES) and lower scores lead the model to predict 0 (NO). The features that were important to predicting this observation are shown in red and blue, with red representing features that pushed the model score higher such as; education, coapplicant income, loan amount, nth term, and the credit_history. While, blue represents features that pushed the score lower such as; gender, marital status, applicant income, property area, and self_employed. The red color typically represents positive contributions, indicating that the feature value increases the prediction compared to the baseline or expected value. Features highlighted in red have a positive impact on the output and contribute to pushing the prediction higher. On the other hand, the blue color represents negative contributions, indicating that the feature value decreases the prediction compared to the baseline or expected value. Features highlighted in blue have a negative impact on the output and contribute to pulling the prediction lower. Features that had more of an impact on the score are located closer to the dividing boundary between red and blue, and the size of that impact is represented by the size of the bar. This particular loan instance was ultimately classified as a yes because they were pushed higher by all the factors shown in red and had a high score of 0.74. The scatter plot shows the different features and the extent of their contributions to the explanation of a given instance.

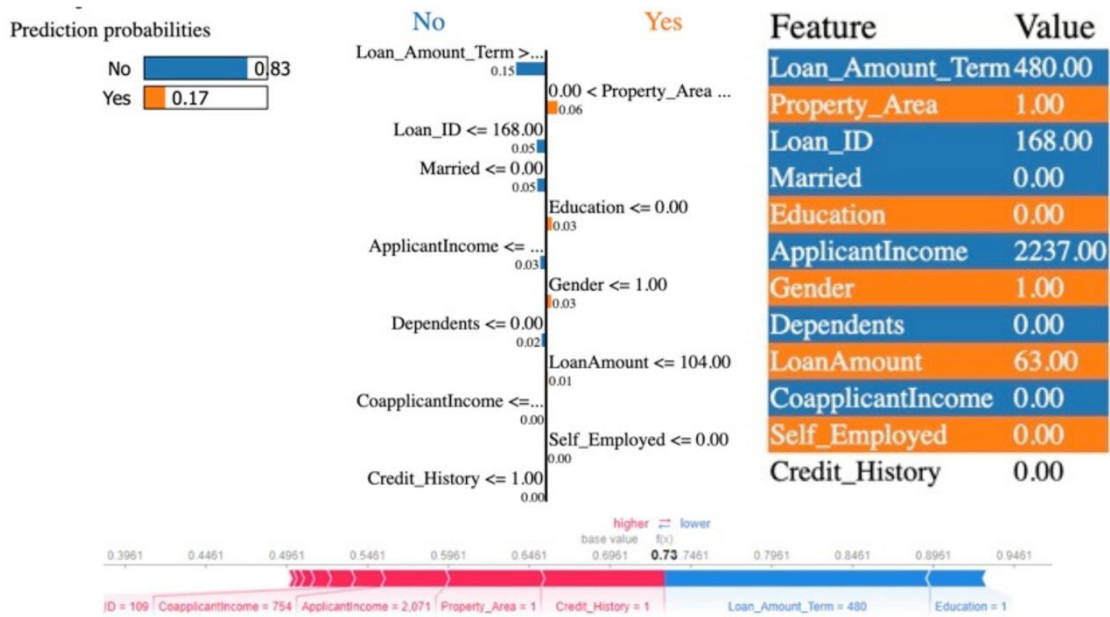


Fig. 6. SHAP force plot for a loan instance

4.2.3 Dependence Plot Explanation

The dependence plot in Figure 7 visualizes the relationship between a specific feature and its corresponding SHAP values. In this dependence plot, the feature selected is credit_history. The plot shows how changes in the credit_history feature influence the predicted outcome of the loan approval system. It also helps one to understand the direction and strength of the relationship between the credit_history feature and the loan approval decision. The plot can be interpreted by observing the distribution of points and the trend line. Each dot is a single prediction (row) from the dataset. The x-axis is the value of the feature (credit_history). While, the y-axis is the SHAP value for that feature, which represents how much knowing that feature's value changes the output of the model for that sample's prediction. The color effect corresponds to a second feature that may have an interaction effect with the feature that is plotted (by default this second feature is chosen automatically) in this case, it is the applicant's income. If an interaction effect is present between this other feature and the feature under consideration, it will show up as a distinct vertical pattern of coloring. For instance, the credit history here has no direct interaction with the applicant's income. However, a sudden loss or reduction in earnings could affect the credit_history indirectly hence, the dependence plot shown.

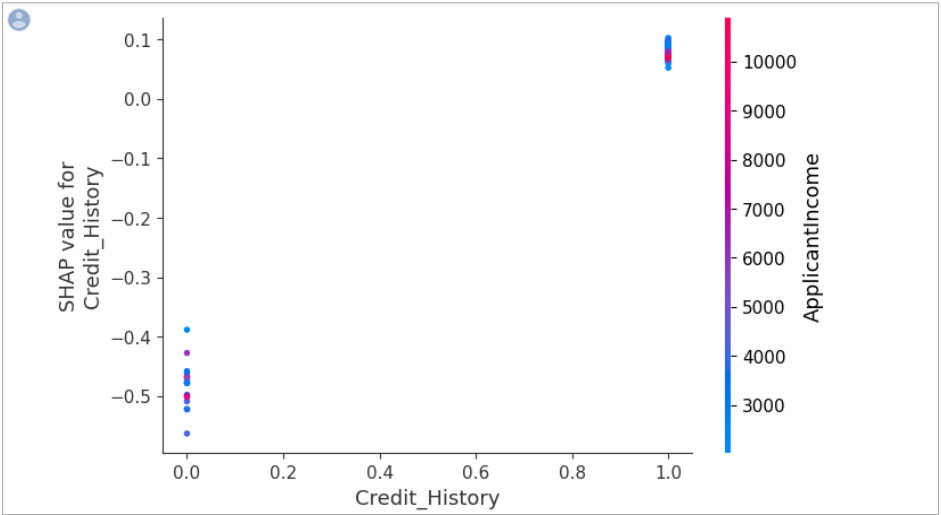


Fig. 7. Dependence plot for a loan instance.

4.2.4 Interaction Plot Explanation

Finally, the interaction plot shown in Figure 8 visualizes the interactions between different features and their combined effects on the model predictions. In this loan instance, the interaction plot is generated using the SHAP interaction values. This provides insights into how two or more features interact with each other to influence the loan approval decision. The plot displays summary information about the interaction effects across the dataset. Each dot on the plot represents a specific combination of feature values, and its position represents the impact of that combination on the loan approval decision. The color of the dots represents the value of the feature that is being interacted with. It helps to identify which combinations of features have the most significant impact on the loan approval decision.

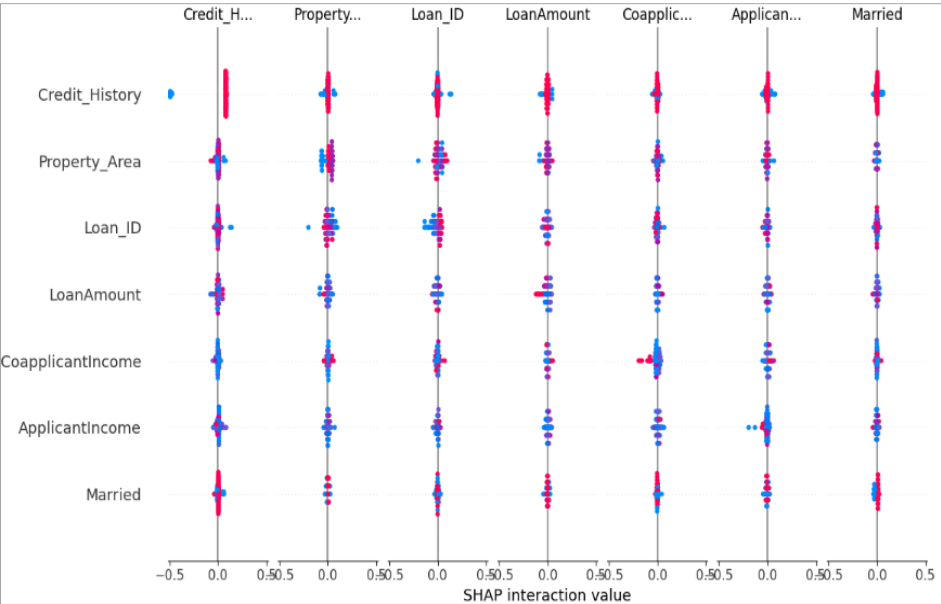


Fig. 8. SHAP Interaction plot for a given loan instance

4.3 Evaluation & Application

The explanations were shown to loan officers comprising 10 participants (six males and four females) and their evaluation result shows SHAP generated 10 important features than LIME which generated 8 important features that could aid them in making loan application decisions. It took LIME 0.1486 seconds to generate explanations using 9.14 MB of memory. SHAP used 0.3784 seconds to generate explanations consuming a memory size of 1.2 MB. By demonstrating SHAP's capacity to record feature interactions and offer a more sophisticated comprehension of the relationships between important factors (such as income, credit history, and loan size) in loan approval decisions, the work advances the state of the art. This is important because it helps identify biases or unfair tendencies that more straightforward prediction tools could miss. By directly comparing the computational efficiency of LIME and SHAP, this study makes a novel contribution. It shows that while SHAP's in-depth explanations may be more appropriate for

regulatory compliance and auditing tasks, LIME's faster runtime makes it more suitable for real-time loan approval systems. The study also suggests investigating hybrid frameworks that combine the computing power of LIME with the in-depth feature analysis of SHAP, which could result in more reliable and transparent AI systems for financial decision-making. This study acknowledges the possibility that AI algorithms could maintain bias in judgments about loan acceptance. Institutions can identify instances in which sensitive factors (e.g., gender, race) are unfairly influencing model decisions and take corrective action to maintain compliance with fairness rules, such as the Fair Lending Act implemented at some jurisdictions. Explainable AI frameworks help financial institutions, regulatory agencies, and loan applicants to easily comprehend the rationale behind a loan's approval or denial. Their applicants can easily know why their applications are denied and what aspect of their profile needs improvement. Loan officers can verify the AI's judgments, especially in borderline cases, by using the justifications provided by LIME and SHAP. For instance, a loan officer may prioritize credit history when deciding if SHAP indicates that it is the most significant element in the application. Financial organizations can customize loan packages for specific consumers by identifying critical characteristics (such as income and credit history) that impact loan approvals. By highlighting patterns suggestive of fraud or low creditworthiness, LIME and SHAP can assist institutions in optimizing risk evaluations. Loan officers can look into suspicious patterns in real time by analyzing which variables are influencing high-risk judgments.

5 Conclusion

In loan approval systems, the ability to assess the eligibility of individuals and businesses for financial support is of utmost importance. This paper provides a comprehensive comparative analysis of LIME and SHAP within the context of loan approval systems. The analysis evaluated the performance, effectiveness, and interpretability of both XAI frameworks in the context of loan approval systems. The study utilized a loan dataset obtained from Kaggle and implemented LIME and SHAP frameworks to generate explanations for loan approval decisions using the Random Forest model. The results were then discussed and interpreted. Comparatively, LIME provides elaborate insights at the individual instance level while SHAP offers a holistic understanding of feature importance. LIME accommodates different ML approaches compared to SHAP. In terms of interpretability, LIME explanations are simpler and more easily understandable compared to SHAP. SHAP explanations are more complex due to the calculation of Shapley values and the consideration of feature interactions. While they may require a deeper understanding to interpret fully, SHAP explanations reveal hidden connections between features and unveil intricate relationships that contribute to the loan approval decision. The application of LIME and SHAP in loan approval systems has significant implications. It enhances transparency by providing interpretable explanations for individual loan instances and feature contributions, thus enabling stakeholders to understand the decision-making process and detect potential biases. Further research should explore datasets from other sources, as well as other application domains. Advances should also be made in investigating the possibility of combining LIME and SHAP into a hybrid framework that leverages the strength of both techniques.

References

- [1] OECD. Artificial Intelligence, Machine Learning and Big Data in Finance: Opportunities, Challenges, and Implications for Policy Makers, <https://www.oecd.org/finance/artificial-intelligence-machine-learning-big-data-in-finance.htm>, 2021.
- [2] V. Hassija, V. Chamola, and A. Mahapatra. Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence. *Cogn Comput* 16, 45–74, 2024.
- [3] Guidotti, R., Monreale, A., Ruggieri, S., Turini, F., Giannotti, F., & Pedreschi, D. (2018). A Survey of Methods for Explaining Black Box Models. *ACM Computing Surveys*, 51(5), 1–42.
- [4] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You?": Explaining the Predictions of Any Classifier. *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics*, pp. 97 - 101
- [5] Erasmo Purificato, Flavio Lorenzo, Francesca Fallucchi & Ernesto William De Luca (2022): The Use of Responsible Artificial Intelligence Techniques in the Context of Loan Approval Processes, *International Journal of Human-Computer Interaction*, DOI: 10.1080/10447318.2022.2081284
- [6] Diwate, Y., Prashant, Rana, P. S., and Chavan, P. A. (2023). Loan approval prediction using machine learning. *International Research Journal of Modernization in Engineering Technology and Science*, doi: 10.56726/irjmets39658
- [7] Caruana, R., Lou, Y., Gehrke, J., Koch, P. L., Sturm, M., & Elhadad, N. (2015). *Intelligible Models for HealthCare*. <https://doi.org/10.1145/2783258.2788613>
- [8] Madhav, A.V.S., and Tyagi, A.K. (2023). Explainable Artificial Intelligence (XAI): Connecting Artificial Decision-Making and Human Trust in Autonomous Vehicles. In: Singh, P.K., Wierzechoń, S.T., Tanwar, S., Rodrigues, J.J.P.C., Ganzha, M. (eds) *Proceedings of Third International Conference on Computing, Communications, and Cyber-Security. Lecture Notes in Networks and Systems*, vol 421. Springer, Singapore. https://doi.org/10.1007/978-981-19-1142-2_10
- [9] Wachter, S., Mittelstadt, B., and Floridi, L. (2017). Transparent, explainable, and accountable AI for robotics. *Science Robotics*. Doi: 10.1126/scirobotics.aan6080
- [10] Valina, L., Teixeira, B., Reis, A., Vale, Z., and Pinto, T. (2024) Explainable Artificial Intelligence for Deep Synthetic Data Generation Models. *IEEE Conference on Artificial Intelligence (CAI)*, Singapore, pp. 555-556, doi: 10.1109/CAI59869.2024.00109.

- [11] Baesens, B., Van Gestel, T., Viaene, S., Stepanova, M., Suykens, J. K., & Vanthienen, J. (2003). Benchmarking state-of-the-art classification algorithms for credit scoring. *Journal of the Operational Research Society*, 54(6), 627–635. Doi: <https://doi.org/10.1057/palgrave.jors.2601545>
- [12] Wang, F., and Rudin, C. (2015). Falling Rule Lists. In *International Conference on Artificial Intelligence and Statistics*. pp. 1013–1022.
- [13] Al Islam et al. (2023) Al Islam, F., Saha, A., Bristy, E. J., Islam, Md. R., Afzal, R. and Ridita, S.A. (2023). Lime-based Explainable AI Models for Predicting Disease from Patient's Symptoms. *International Conference on Computing Communication and Networking Technologies*. DOI:10.1109/ICCCNT56998.2023.10307223
- [14] Lundberg, S. M., & Lee, S. (2017). A unified approach to interpreting model predictions. In *Neural Information Processing Systems* 30, pp. 4768–4777.
- [15] El Shaw R., Sherif, Y., Al-Mallah, M.H., and Sakr, S. (2020). Interpretability in healthcare: A comparative study of local machine learning interpretability techniques. *Computational Intelligence*, 37(8). DOI:10.1111/coin.12410
- [16] Aldughayfiq, B., Ashfaq, F., Jhanjhi, N.J. and Humayun, M. (2023). Explainable AI for Retinoblastoma Diagnosis: Interpreting Deep Learning Models with LIME and SHAP, in *Diagnostics* 13(11). <https://doi.org/10.3390/diagnostics13111932>
- [17] Veerappa, M., Anneken, M., Burkart, N., and Huber, M. (2021). Validation of XAI Explanations for Multivariate Time Series Classification in the Maritime Domain in *Journal of Computational Science*. Vol. 58. Doi: 10.1016/j.jocs.2021.101539
- [18] Tyagi S. (2022). Analyzing Machine Learning Models for Credit Scoring with Explainable AI and Optimizing Investment Decisions. *American International Journal of Business Management (AIJBM)* ISSN- 2379-106X Vol. 5 (01).
- [19] Gramegna, A., and Giudici, P. (2021). SHAP and LIME: An Evaluation of Discriminative Power in Credit Risk. *Frontiers in Artificial Intelligence*. Vol. 4.
- [20] Ariza-Garzon, M.J., Arroyo, J., Caparrini, A., and Segovia-Vargas, M., 2020. Explainability of a Machine Learning Granting Scoring Model in Peer-to-Peer Lending. *IEEE Access*. Doi: 10.1109/ACCESS.2020.2984412.
- [21] Chen, Y., Calabrese, R., Martin-Barragan, B. (2023). Interpretable Machine Learning for Imbalanced Credit Scoring Datasets, *European Journal of Operational Research* 312(1), 357-372.

Authors' Profiles



Isaac Terngu Adom holds a MSc. in Computer Science from the University of Bonn, Germany, and a BSc. in Computer Science with First Class Honors from Ahmadu Bello University, Zaria Nigeria. He is currently working as a University Lecturer and Researcher at the Department of Mathematics and Computer Science, Benue State University, Makurdi, Nigeria. His research interest is in Data Science, Data Mining, Artificial Intelligence, Knowledge Representation, and Explainable Artificial Intelligence. A member of the Nigeria Computer Society (NCS), Isaac has many years of industry, teaching, and research experience where he has worked on many projects, taught Computer Science courses, and currently supervising undergraduate research projects. An assiduous academic who is committed to imparting and creating knowledge, Isaac has published papers in reputable journals and presented at conferences. Isaac has six (6) accolades from different affiliations.



Christiana O. Julius holds a Bachelor of Science degree in Computer Science from Benue State University, Makurdi, Nigeria. She is a research assistant in the Department of Mathematics and Computer Science, at Benue State University. Her research interest is in Artificial Intelligence, Natural Language Processing, Human-Computer Interaction and Explainable Artificial Intelligence. Christiana carries out lab classes for undergraduate computer science students as well as other assigned administrative responsibilities. She is creative and innovative in solving complex problems using programming languages and Artificial Intelligence tools and packages.



Stephen Akuma is a Senior Lecturer of Computer Science in the Department of Mathematics/Computer Science at Benue State University. Stephen holds a PhD in Computing and a Master's degree in Software Development (Distinction) from Coventry University, United Kingdom. He also obtained a Bachelor's degree in Computer Science (2.1) from Benue State University. He is an experienced research scientist and university lecturer for over 15 years with a demonstrated history of working in data science, human-computer interaction, feedback indicators, eye gaze-enhanced interaction, user modeling and personalization. He has been able to continuously take ideas from conception, and development, to deployment in lab-based and large-scale environments. He has published in leading scientific proceedings and journals.



Samera Uga Otor is a distinguished academic and researcher, currently serving as a Senior Lecturer at the Benue State University, Makurdi, where she has been involved in teaching and research since 2010. She holds a Ph.D. in Computer Science from Obafemi Awolowo University, Ile-Ife, specializing in cybersecurity and optimization. Dr. Otor has authored several peer-reviewed publications, focusing on network security, intrusion detection models, mobile networks, Artificial intelligence, Data Science and IoT. Additionally, she actively contributes to academic service through roles in program development, student supervision, and committee memberships. She is also a member of professional organizations like Nigeria Women in Information Technology, Nigeria Computer Society and Computer Professional Registration Council of Nigeria.

How to cite this paper: Isaac Terngu Adom, Christiana O. Julius, Stephen Akuma, Samera U. Otor, "Comparative Analysis of Explainable AI Frameworks (LIME and SHAP) in Loan Approval Systems", International Journal of Information Engineering and Electronic Business(IJIEEB), Vol.17, No.6, pp. 60-70, 2025. DOI:10.5815/ijieeb.2025.06.05