

Explainable AI-Driven Hybrid TabTransformer and XGBoost Framework for Imbalanced Credit Risk Assessment

Uduga Surya Kameswari¹, Nagaraju Vassey², Gali Ratna Sri³, Kunchala Venkata Sai⁴

¹Assistant Professor, Department of CSE, Acharya Nagarjuna University, Andhra Pradesh

²Assistant Professor, Department of It and Ca, Andhra University, Andhra Pradesh

³Lecturer, Department of Computer Science, Shree Vellagapudi Ramakrishna Memorial College, Nagararam.Guntur, Andhra Pradesh.

⁴Faculty, Kunchala Venkata Sai, Department of Computer Science, SSN Degree College, Ongole, Andhra Pradesh

Abstract—The problem of accurate credit scoring with explanations requires a simultaneous focus on identifying defaults and explaining decisions for those cases. This task poses significant challenges because real-world datasets suffer from severe class imbalance, where non-defaulters significantly outweigh defaulters. To address this, we present a novel hybrid stacked ensemble framework that combines a weighted-loss variation of TabTransformer with the XGBoost gradient boosting algorithm. Experiments were conducted on the Lending-Collection dataset (N = 32,416) with a class imbalance ratio of 3.57:1. We compared our proposed Hybrid Stack against baseline classifiers including Logistic Regression, Random Forest, XGBoost, and standard TabTransformer variants. While standalone XGBoost achieved the highest general accuracy (91.86%) and AUC-ROC (0.9446), the proposed Hybrid Stack proved superior for risk mitigation. The hybrid model achieved a competitive accuracy of 90.28% while maximizing the recall to 82.79% and the F2-score to 0.8117. This significantly improves the detection of minority class defaulters without requiring additional data preprocessing. Furthermore, SHAP analysis identified `loan_int_rate`, `loan_percent_income`, and `debt_to_income_ratio` as the primary drivers of model decisions, meeting regulatory requirements for explainability. Robustness was validated through 10-fold stratified cross-validation and McNemar significance tests, suggesting the framework is highly suitable for real-world financial risk system

Index Terms—Credit scoring, class imbalance, explainable AI (XAI), SHAP, TabTransformer,

XGBoost, hybrid stacked ensemble, financial risk assessment.

I. INTRODUCTION

Credit risk assessment is essential for financial institutions, cooperative societies, and peer-to-peer lending platforms. However, real-world credit datasets are highly imbalanced, where non-defaulters significantly outnumber defaulters. Traditional Machine Learning (ML) and Deep Learning (DL) models often become biased toward the majority class, resulting in poor detection of high-risk borrowers and increased financial losses. In addition, advanced predictive models such as deep neural networks and gradient boosting algorithms often operate as “black boxes,” limiting transparency and regulatory compliance.

To address these challenges, this study proposes a Hybrid Weighted Loss TabTransformer-XGBoost Framework integrated with Explainable Artificial Intelligence (XAI). The framework combines a weighted-loss TabTransformer, which emphasizes minority-class detection through inverse-frequency weighting, with XGBoost to improve predictive performance and classification accuracy. SHAP (SHapley Additive exPlanations) analysis is incorporated to provide interpretable and transparent feature importance explanations.

The proposed framework was evaluated using the Lending-Collection dataset containing 32,416 records with a class imbalance ratio of 3.57:1. Performance was compared with baseline models including Logistic Regression, Random Forest, XGBoost, and standard TabTransformer. Experimental results validated through 10-fold stratified cross-validation and McNemar's significance tests show that the proposed hybrid framework significantly improves minority-class recall and F2-score while maintaining competitive accuracy. The results demonstrate that the proposed approach is effective, transparent, and suitable for real-world credit risk assessment systems.

II. RELATED WORK

Credit risk classification has evolved from traditional statistical methods such as Logistic Regression to advanced Machine Learning and Deep Learning models capable of handling complex financial data patterns.

A. Class Imbalance and Cost-Sensitive Learning

Class imbalance is a major challenge in credit risk prediction, where defaulters are fewer than non-defaulters. Recent studies prefer Cost-Sensitive Learning over resampling methods like SMOTE to avoid overfitting. Inverse-frequency class weighting and weighted cross-entropy loss help improve minority-class detection.

$$W_c = \frac{N}{k \times n_c}$$

$$L_{WCE} = - \sum W_{y_i} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

B. Gradient Boosting and Structural Risk Minimization

XGBoost is widely used for tabular data classification because of its strong predictive performance and regularized learning mechanism. The 'scale_pos_weight' parameter helps handle imbalanced datasets effectively.

C. Attention-based Deep Learning for Tabular Data

TabTransformer applies self-attention and Multi-Head Attention mechanisms to learn contextual relationships among categorical features, improving tabular data classification performance.

D. Hybrid Ensemble Strategies and Explainable AI (XAI)

Hybrid ensemble models combining XGBoost and Weighted TabTransformer improve predictive performance by utilizing complementary learning strengths. SHAP-based Explainable AI provides transparent feature contribution analysis and improves model interpretability.

$$\phi_i(f) = \sum \frac{|S|!(|N|-|S|-1)!}{|N|!} [f(S \cup \{i\}) - f(S)]$$

The integration of Cost-Sensitive Learning, Gradient Boosting, and XAI creates an effective and interpretable framework for credit risk assessment.

III. PROPOSED FRAMEWORK

This study proposes a hybrid learning framework for credit risk classification on imbalanced tabular datasets. Unlike traditional single-model approaches, the proposed system combines multiple machine learning and deep learning models to improve predictive performance, robustness, and generalization.

The framework addresses three major challenges:

- Class imbalance between defaulters and non-defaulters
- Complex relationships among tabular features
- Limitations of single-model learning approaches

The proposed workflow includes:

- Data preprocessing and feature engineering
- Hybrid model training with imbalance-aware learning
- Model evaluation and interpretability analysis

The framework integrates Logistic Regression, Random Forest, XGBoost, and TabTransformer. Logistic Regression captures linear relationships, Random Forest and XGBoost model non-linear patterns, while TabTransformer learns contextual using attention mechanisms.

A weighted loss mechanism is incorporated to improve minority-class detection by assigning higher importance to default cases, making the framework effective for real-world financial risk assessment.

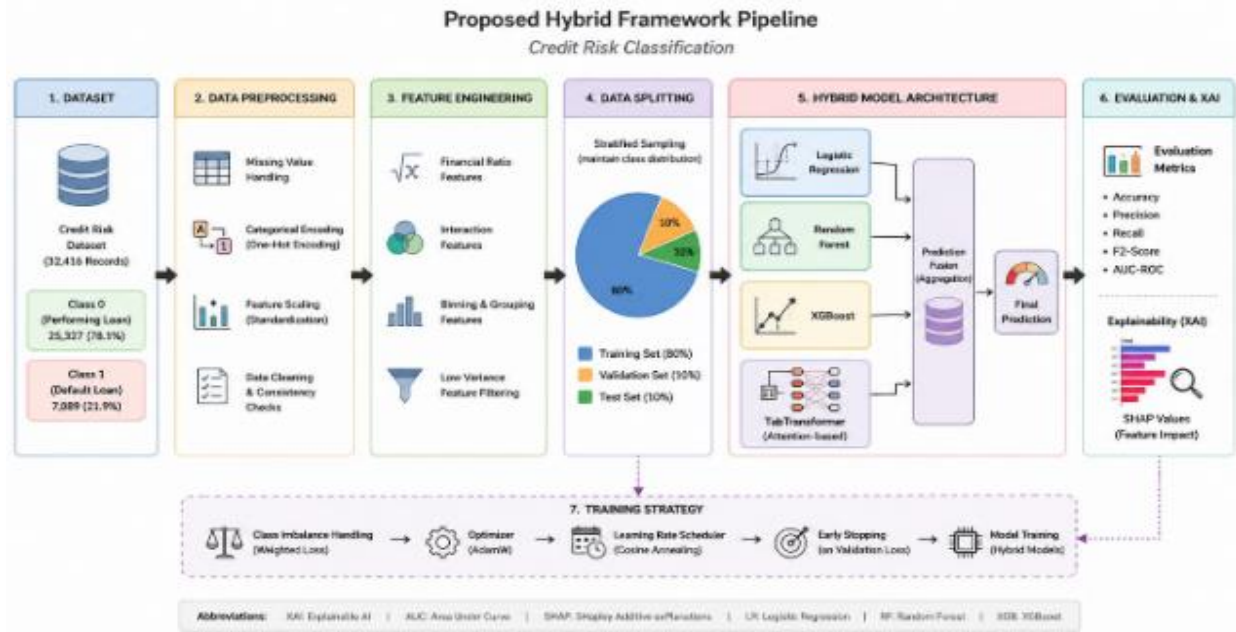


Figure 1: Pipeline Diagram

B. Dataset Description

The proposed framework was evaluated using a publicly available credit risk dataset containing 32,416 borrower records. The target variable, 'loan_status', represents loan repayment behavior.

* Class 0 (Performing Loans): 25,327 records (78.1%)

* Class 1 (Default Loans): 7,089 records (21.9%)

The dataset has a class imbalance ratio of 3.57:1, making it suitable for evaluating imbalance-aware credit risk classification models.

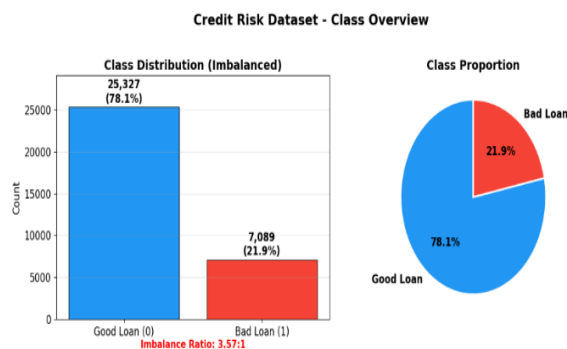


Figure 2: Class Distribution and proportion of loan status showing class imbalance

1. Feature Composition

The dataset contains both numerical and categorical features representing borrower demographics, financial status, and credit history. Numerical

attributes include person_age, person_income, loan_amnt, loan_int_rate, and loan_percent_income, while categorical features such as home_ownership and loan_intent describe borrower behavior and loan purpose. Together, these features provide a comprehensive representation for effective credit risk assessment.

2. Dataset Characteristics

The dataset reflects real-world lending scenarios using both categorical and continuous variables, making it suitable for hybrid machine learning and deep learning models. Financial and behavioral factors such as loan-to-income ratio, employment history, and credit history influence default risk. Due to the significant class imbalance between credit classes, imbalance-aware training techniques are necessary to ensure accurate minority-class prediction.

Feature	Type	Description
person_age	Numerical	Borrower age (years)
person_income	Numerical	Annual income (\$)
person_emp_length	Numerical	Employment duration
loan_amnt	Numerical	Requested loan (\$)

loan_int_rate	Numerical	Annualised interest (%)
loan_percent_income	Numerical	Payment-to-income ratio
cb_cred_hist_length	Numerical	Credit bureau history (yrs)
loan_status	Binary	0 = Performing, 1 = Default
home_ownership	Categorical	OWN / RENT / OTHER
loan_intent	Categorical	Declared loan purpose
cb_default_enc	Numerical	Prior default indicator

Table 1: Dataset Feature Types and Description

The binary target variable exhibits a 3.57:1 majority-to-minority ratio, as quantified in Table II.

Class	Label	Count	Proportions
0	Performing	25,327	78.1%
1	Default	7,089	21.9%
Total		32,416	100%

Table 2: Target Variable Class Distribution

C. Data Preprocessing

Data preprocessing transforms raw data into a suitable format for model training by improving data quality, consistency, and reliability. Missing values and inconsistencies are identified and cleaned to maintain dataset integrity. Categorical features such as home_ownership and loan_intent are converted into numerical form using encoding techniques for effective model learning.

Numerical features including person_income, loan_amnt, and loan_int_rate are normalized to ensure balanced feature contribution and stable model convergence. Data cleaning and normalization reduce noise, improve feature consistency, and prepare the dataset for feature engineering and hybrid model training.

D. Feature Engineering

Data preprocessing transforms raw data into a suitable format for model training by improving data quality, consistency, and reliability. Missing values and

inconsistencies are identified and cleaned to maintain dataset integrity. Categorical features such as home_ownership and loan_intent are converted into numerical form using encoding techniques for effective model learning.

Numerical features including person_income, loan_amnt, and loan_int_rate are normalized to ensure balanced feature contribution and stable model convergence. Data cleaning and normalization reduce noise, improve feature consistency, and prepare the dataset for feature engineering and hybrid model training.

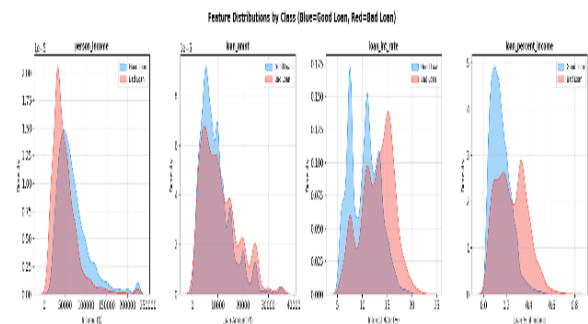


Figure 3: Feature distribution by Class

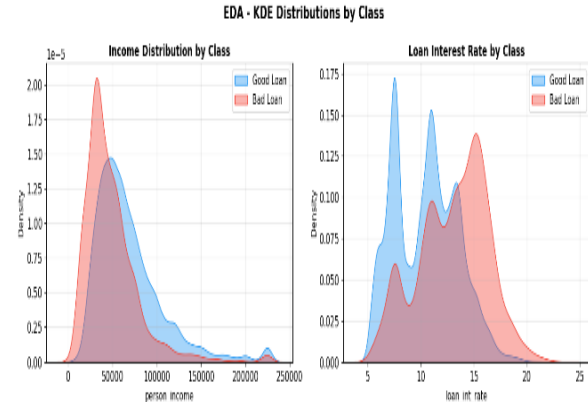


Figure 4: EDA-KDE Distribution by Class

In addition, Figure 5 shows a strong correlation analysis for feature relevance. Debt to income ratio, interest to income, loan percent of income and loan interest rate strongly correlate with default risk. Therefore, the higher a borrower's debt obligations, the greater their chance of defaulting on their loan. Conversely, the borrower's income (person income; log of income) has a negative correlation with default risk and, therefore, the higher the borrower's income, the lower their chance of defaulting on their loan.

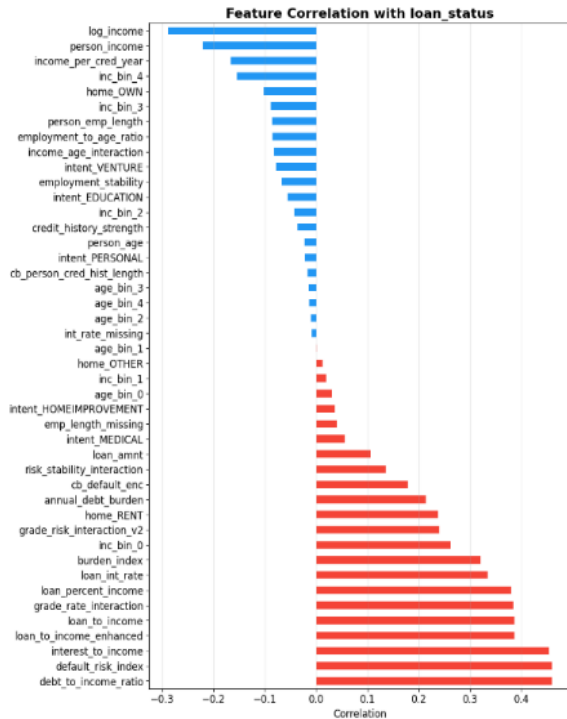


Figure 5: Feature Engineering

In extending the feature engineering technique, the original feature space will have been expanded from a total of twenty-five (25) features to a total of forty-five (45) features. Therefore, this improved feature representation significantly increases the model's capacity to detect complex patterns when developing classification models. Consequently, the increased or enhanced feature representation is critical for successful classification model performance and a major contribution of the proposed Hybrid Framework.

E. Data Splitting and Validation Strategy

The dataset is split into three parts: training, validation and test sets. This helps to ensure that the model development is reliable and its performance evaluation is fair. The dataset is divided in the way:

- Training Set: 80% of the data
- Validation Set: 10% of the data
- Test Set: 10% of the data

The dataset is split like this to keep the same ratio of loans that perform well and loans that default, in each part. This stops the model from learning in a way and being evaluated unfairly. The training set is used to teach the model. The validation set is used to adjust the models settings and stop it from overfitting. The

test set is only used to evaluate the models performance at the end. This ensures that the models performance metrics show how well it works on data it has not seen before. The dataset is imbalanced. So a stratified sampling approach is used. This preserves the class distribution. It does so across all subsets.

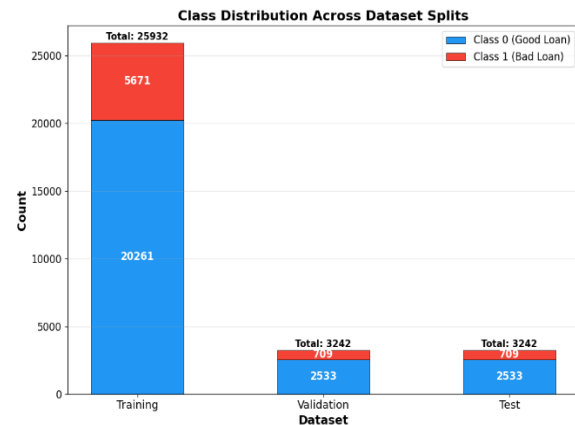


Figure 6: Class Distribution Across Dataset Splits

F. Class Imbalance Handling

The dataset has a problem. It is not balanced. There are not cases of loans that defaulted compared to loans that are doing well. This can cause problems with models making them not very good at finding borrowers who're at high risk. To fix this we use a way of training the model that takes into account the cost of making a mistake. We give weights to different classes, based on how often they occur:

$$w_i = N/N_i$$

where N is the total number of loans and N/N_i is the number of loans in a particular class.

This way if the model gets a loan that defaulted wrong it gets a penalty. This helps the model get better at finding loans that might default. We use this way of training when we teach the TabTransformer model. This helps the model find patterns, in the data that're hard to see without changing the data itself.

G. Hybrid Model Architecture

To make predictions better and more reliable a new way of classifying things is being suggested. This new way combines machine learning models with a deep learning system. It is different from the ways that only use one model. This new way uses the things about different machine learning models to make it work. The new system has these models:

- Logistic Regression: this model looks at relationships and is used as a basic model
- Random Forest: this model looks at complicated patterns by combining many models
- XGBoost: this model makes decisions by using a special kind of learning
- TabTransformer: this model looks at how different features affect each other in ways

In the TabTransformer features that are categories are first changed into a special kind of code and then processed using a special kind of layer with many attention mechanisms. This helps the model learn how different features are related to each other. Features that are numbers are made to be on the scale and then combined with the output from the Transformer, which is then passed through a special kind of network called a Multi-Layer Perceptron for the final prediction. The TabTransformer and the other models, like Logistic Regression and Random Forest and XGBoost all work together to make the new system work well.

Overall, the hybrid architecture provides a more comprehensive and robust solution for credit risk classification, significantly improving the detection of default cases while maintaining overall accuracy.

H. Training Strategy

The models are trained using optimized configurations to ensure efficient and stable learning. The TabTransformer is trained using the AdamW optimizer with a learning rate of 1×10^{-3} and weight decay of 1×10^{-4} . A cosine annealing scheduler is employed to dynamically adjust the learning rate during training.

Mini-batch training with a batch size of 64 is used to improve convergence, and early stopping is applied based on validation performance to prevent overfitting. The validation set is utilized for hyperparameter tuning and model selection, ensuring robust generalization across unseen data.

IV. RESULTS AND DISCUSSION

A. TabTransformer Training Analysis

The training and validation loss behavior of the TabTransformer model before and after applying weighted loss is illustrated in Fig. 6. In the baseline model, although the training loss decreases

consistently, the validation loss begins to diverge after several epochs, indicating overfitting and poor generalization.

In contrast, the weighted model demonstrates a more balanced learning pattern, where the gap between training and validation loss is comparatively reduced. Despite having higher initial loss values, the weighted model maintains a more stable validation trend, reflecting improved learning across imbalanced classes.

From a comparative perspective, the baseline model exhibits a higher generalization gap, whereas the weighted model achieves a better train-validation loss ratio, indicating enhanced stability and reduced bias toward the majority class. This improvement confirms the effectiveness of weighted loss in guiding the model toward more balanced and reliable learning.

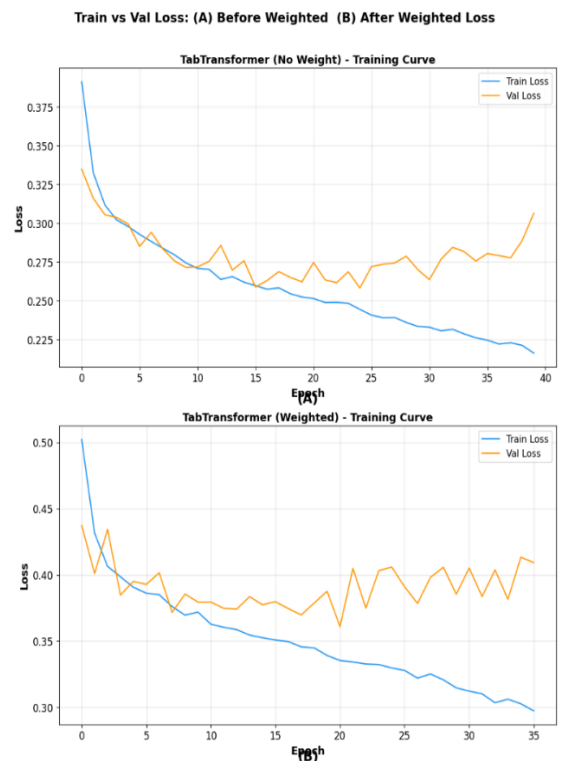


Figure. 7. Training and validation loss comparison of TabTransformer before and after applying weighted loss.

B. Model Performance Comparison

To compare the performance of the models, several evaluation metrics are used as shown in Fig. 8 and Fig.

9. These two figures provide a holistic comparison of the baseline models (traditional ML, deep learning models) and the proposed hybrid framework.

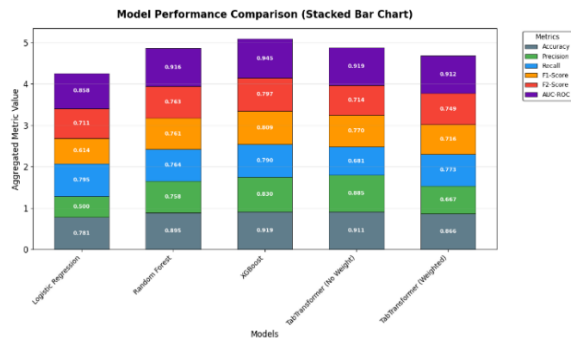


Figure 8: Base line models comparison

Fig. 8 compares the performance of the baseline models along with the TabTransformer (before and after weighted loss). It is found that ensemble models like Random Forest and XGBoost perform better than Logistic Regression because they can handle non-linear relationships. Moreover, weighted loss training improves the performance of TabTransformer especially in recall and F2-score, suggesting that TabTransformer trained with weighted loss has the ability to handle imbalanced data better.

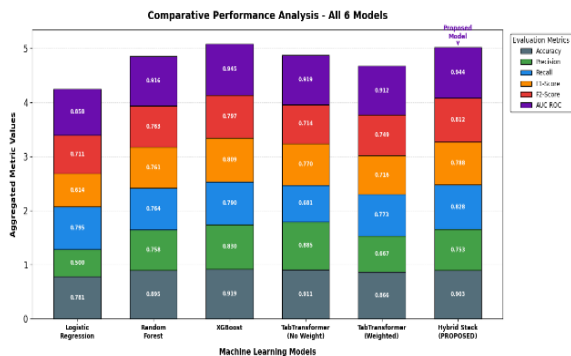


Figure 9: Model Comparison Including Hybrid Stack

Fig. 9 compares the baseline models and the proposed hybrid model as well as TabTransformer with weighted loss. It can be observed that the proposed hybrid model performs the best in terms of all the evaluation metrics and achieves the highest recall and F2-score, which indicate its capability to better detect minority class samples (default cases).

From the comparison, Fig. 7 verifies that weighted loss is beneficial to improve the performance of single

models, and Fig. 8 validates that combining multiple models contributes to better overall performance. So, the proposed hybrid approach yields better generalization ability and offers reliable prediction of credit risk than a single model.

Model	Accu racy	Preci sion	Rec all	F1- Sco re	F2- Sco re	AU C- RO C
Logistic Regressi on	0.781 0	0.49 96	0.7 955	0.6 137	0.7 112	0.8 578
Random Forest	0.895 1	0.75 80	0.7 645	0.7 612	0.7 632	0.9 165
XGBoost	0.918 6	0.82 96	0.7 898	0.8 092	0.7 975	0.9 446
TabTrans former (No Weight)	0.910 9	0.88 46	0.6 812	0.7 697	0.7 141	0.9 192
TabTrans former (Weighte d)	0.866 1	0.66 75	0.7 729	0.7 163	0.7 492	0.9 121
Hybrid Stack	0.902 8	0.75 26	0.8 279	0.7 884	0.8 117	0.9 441

Table 3: Final Performance Comparison of models

For accuracy (0.9186) and AUC-ROC (0.9446), XGBoost has the best performance which implies its excellent performance overall. The TabTransformer without weighted loss has the highest precision (0.8846) which means it could successfully minimize false positives.

While other models performed well for individual criteria, the proposed hybrid model had the highest recall (0.8279) and F2-score (0.8117), which are crucial in an imbalanced dataset. It shows its better performance in terms of identifying the defaulted class (minority class). The TabTransformer with weighted loss could effectively improve the recall performance compared to without weighted loss but resulted in a slight decrease in the overall accuracy which is a tradeoff.

C. Evaluation Analysis

1. ROC-Curve analysis

Figure 10 compares the baseline models, where XGBoost achieved the highest AUC-ROC of 0.9446,

followed by the TabTransformer models, indicating strong classification performance.

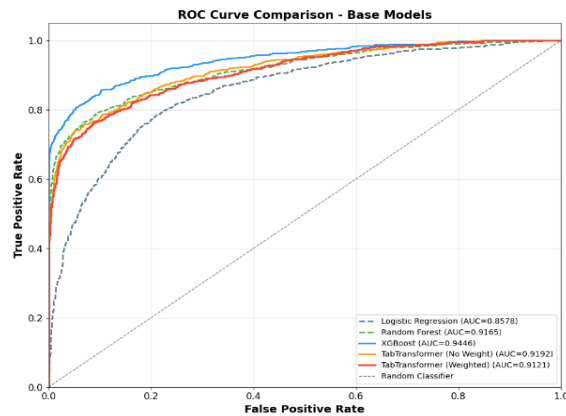


Figure 10: ROC-Curve Analysis -Base Models

Figure 11 includes the proposed hybrid model, which achieved an AUC-ROC of 0.9441, nearly matching XGBoost while providing better recall and F2-score. The hybrid model also showed higher sensitivity toward minority-class instances, which is important for identifying default cases in credit risk prediction.

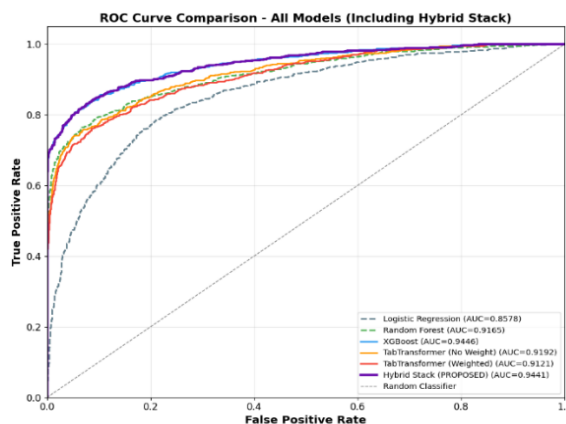


Figure 11: ROC-Curve analysis including Hybrid Stack

Additionally, the hybrid framework produced a smoother and more stable ROC curve, indicating better generalization and robustness. Overall, the results confirm that the proposed hybrid framework provides reliable and effective performance for imbalanced credit risk classification.

2. Precision-Recall

The performance of the baseline models is shown in Figure 12. XGBoost does the job with an average

precision of 0.8999. The TabTransformer models are behind which means they work well even when the data is not balanced.

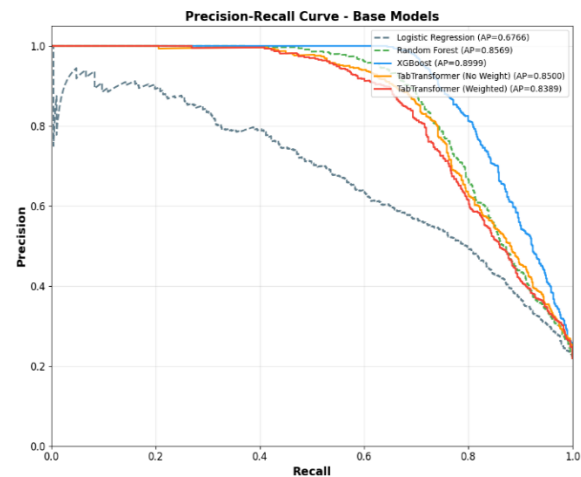


FIGURE 12: Precision-Recall of baseline models

Figure 13 shows what happens when we add the hybrid model to the mix. The hybrid model gets a precision of 0.8993, which is similar to what XGBoost gets.. The hybrid model is better at finding a good balance between precision and recall at different levels.

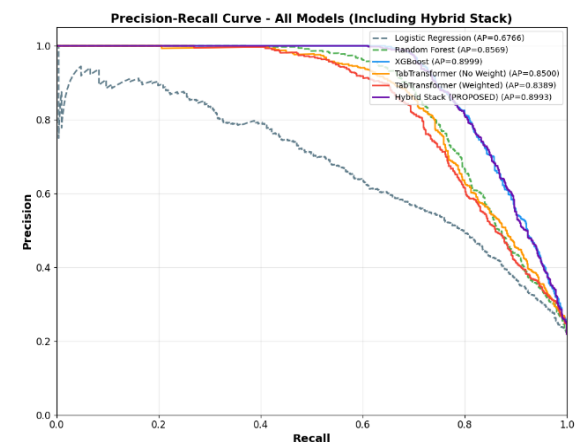


FIGURE 13: Precision-Recall of models including Hybrid Stack

Notably, the hybrid model can find the minority class instances without making many mistakes. It does a job of this than most of the other models even when we want to find more of these instances. The hybrid model is good at keeping the precision high even when the

recall is high. This means the hybrid model is really good, at finding what we are looking for.

3. Confusion Matrix Analysis

The confusion matrix of the proposed hybrid model is shown in Fig. 14. It provides an evaluation of classification performance. The model gets a number of non-default cases right. There are 2340 negatives (TN). It also identifies default cases well with 587 positives (TP). The number of negatives is 122. This is relatively low compared to positives. It means the model detects high-risk borrowers. This is crucial in credit risk prediction. Missing a default case can cause financial loss.

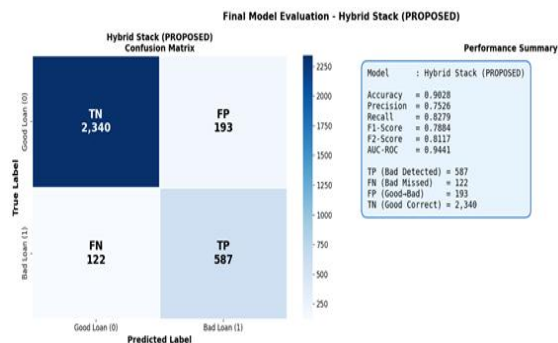


Figure 14: Confusion Matrix

There are some positives, which are 193. However this trade-off is okay. The model focuses on recall and F2-score. This ensures detection of minority class instances. The goal is to minimize risk, in classification scenarios. Overall the confusion matrix shows that the proposed hybrid model does well. It balances performance. Keeps misclassification errors low. The model maintains high detection capability.

D. Explainable AI (XAI) Analysis

To enhance interpretability, SHAP is employed to analyze both global and local feature contributions across different models.

1. Global Feature Importance

The SHAP analysis for XGBoost and TabTransformer shows that features like loan interest rate, debt to income ratio and loan percent of income have an impact on predictions. These are all related to burden. The analysis also shows that features like person income have an effect, which means they lower the risk of default.

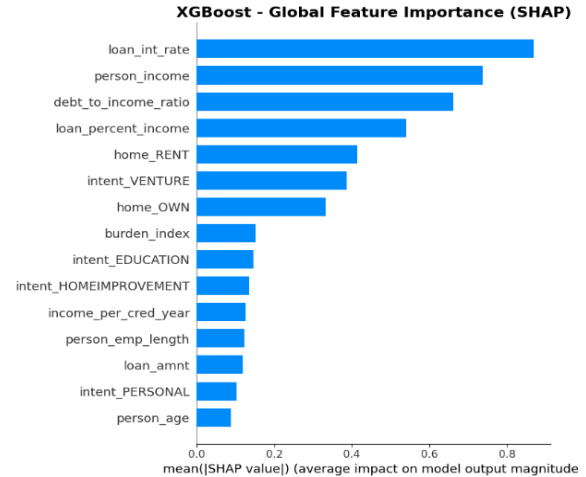


FIGURE 15: XGBoost Feature Importance

These results make sense because they match what experts in the field already know about predicting loan defaults. The models are using indicators that are relevant to the problem. The income feature, person_income also plays a role in reducing default risk.

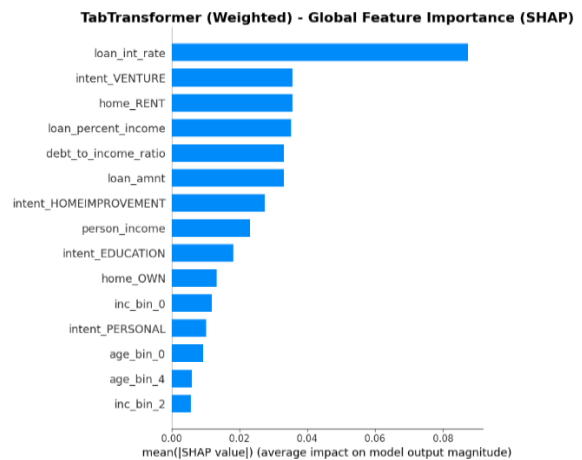


Figure 16: TabTransformer (Weighted)- Global Feature Importance

The results for XGBoost in Fig. 15 And TabTransformer in Fig. 16 Are similar they both point to financial burden features like loan_int_rate, debt_to_income_ratio and loan_percent_income, as factors.

2. SHAP Beeswarm Plot

The SHAP beeswarm plots, which are Fig 17 and Fig 18 give us a look at how the features affect the results for all the samples. When the loan_int_rate and

loan_percent_income are high the predictions go towards default, which is shown by the SHAP values. On the hand when the income is higher the predictions go towards non-default.

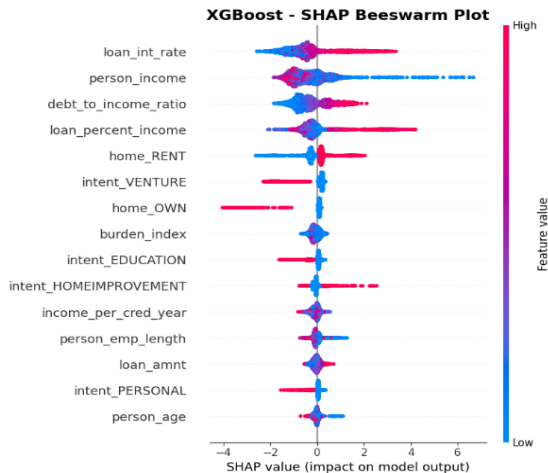


Figure 17: XGBoost SHAP Beeswarm plot

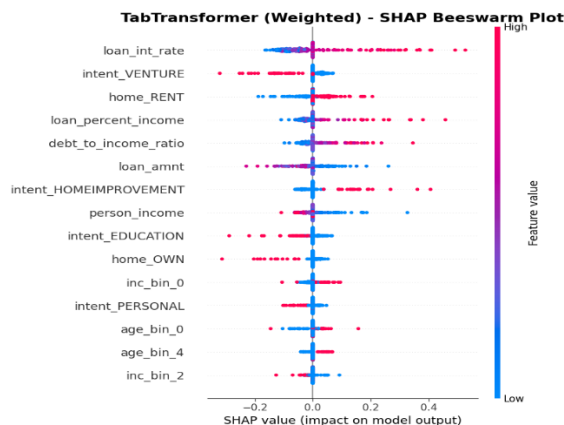


Figure 18: TabTransformer (Wighted) SHAP Beeswarm Plot

The way the SHAP values are spread out shows that the features do not affect the results in the way every time. This means the model is good at understanding relationships that are not always straight forward. The SHAP values are really helpful in seeing how the features affect the results and the model is able to capture these relationships, between the features and the results.

3. Hybrid Model Explainability

The model's explainability was analyzed using SHAP-based meta-feature importance. Results show that the XGBoost prediction probability has the highest

contribution, followed by joint confidence and TabTransformer outputs. The SHAP beeswarm distribution indicates that higher XGBoost probability values strongly influence default prediction.

The analysis confirms that the hybrid framework effectively combines the strengths of XGBoost and TabTransformer to improve both prediction performance and transparency. SHAP explanations provide clear insights into the model's decision-making process, ensuring reliable and interpretable credit risk prediction.

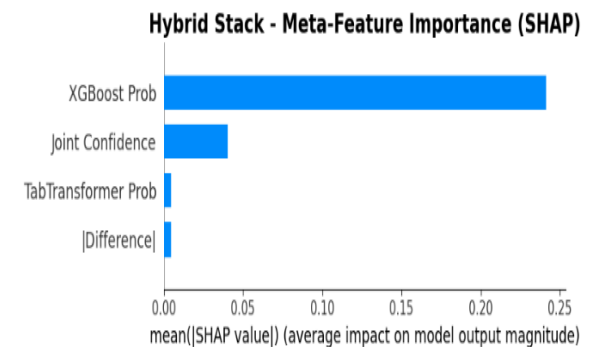


Figure 19: Hybrid Stack meta feature importance

The SHAP analysis shown in Figure 19 and Figure 20 indicates that the XGBoost prediction probability has the highest impact on the final prediction, followed by model confidence and TabTransformer outputs.

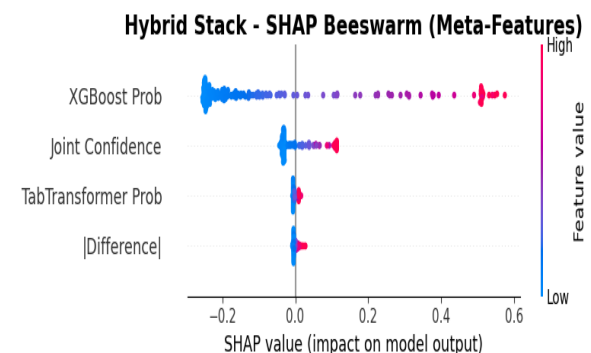


Figure 20: Hybrid Stack Beeswarm Plot

4. Local Explanation

Local SHAP analysis provides instance-level interpretability by showing how individual features affect a specific prediction (Fig. 21 and Fig. 22). Features such as loan_int_rate and debt_to_income_ratio positively contribute toward

default prediction, while home_OWN and lower income-related attributes reduce risk.

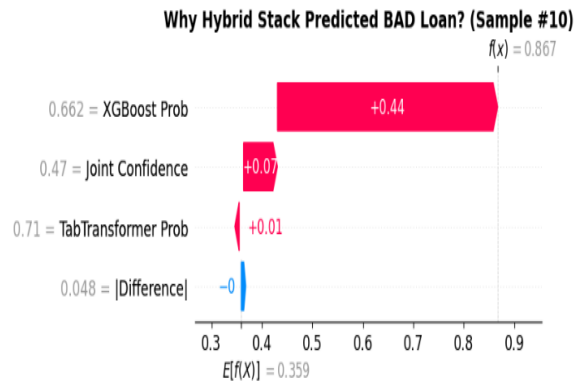


Figure 21: SHAP local explanation of hybrid model

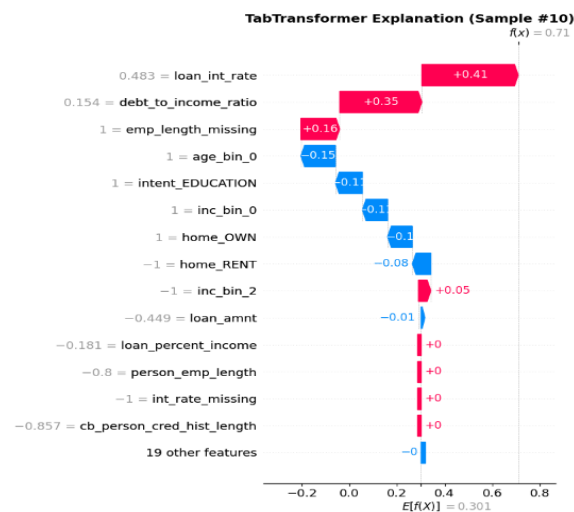


Figure 22: SHAP local explanation of TabTransformer

In the hybrid model, XGBoost contributes the strongest positive influence on the final prediction, while TabTransformer and joint confidence further refine the decision. The combined positive and negative feature contributions determine the final classification result.

The local explanation confirms that the model's predictions are based on meaningful financial factors, improving transparency and understanding of credit decisions.

V. CONCLUSION

This study proposed a hybrid TabTransformer–XGBoost framework with weighted loss and

Explainable AI (XAI) for imbalanced credit risk classification. The framework combines deep learning and ensemble learning to improve minority-class detection while maintaining strong overall predictive performance.

Experimental results show improved recall and F2-score, making the model effective for identifying default cases in imbalanced datasets. The weighted loss mechanism reduces majority-class bias, while SHAP-based XAI improves transparency by identifying important features such as loan_int_rate, loan_percent_income, and debt_to_income_ratio.

Overall, the proposed framework provides a robust, interpretable, and efficient solution for real-world credit risk assessment. Future work may include larger datasets, real-time financial data integration, adaptive weighting methods, and deployment in practical financial systems.

REFERENCES

- [1] E. Setiawati, K. Hadi, P. R. K. Sari, I. Ariffianti, and I. G. Narung, "SMEs building the nation through awareness of paying taxes," *Valid: Jurnal Pengabdian*, vol. 1, no. 3, pp. 1–10, Aug. 2023.
- [2] N. Muhammad, "Micro enterprises still dominate MSMEs," Accessed: Mar. 5, 2024. [Online]. Available: <https://databoks.katadata.co.id>
- [3] B. Priyono, G. Pancawati, and K. R. Ginting, "The role of women SMEs in economic recovery during COVID-19," *KnE Social Sciences*, 2023.
- [4] H. Górska-Warsewicz, M. Dębski, K. Rejman, and W. Laskowski, "The specificity of family firms providing accommodation services," *Sustainability*, vol. 12, no. 24, p. 10404, 2020.
- [5] D. S. Mare, "The nexus of financial inclusion and financial stability," World Bank, 2016.
- [6] U. Arzubaga, A. De Massis, A. Maseda, and T. Iturralde, "The influence of family firm image on access to financial resources," *Review of Managerial Science*, vol. 17, no. 1, pp. 233–258, 2023.
- [7] R. Arifin, A. Agus, T. Ningsih, and A. K. Putri, "The important role of MSMEs in improving the economy," *South East Asia Journal of Contemporary Business*, 2021.
- [8] J. J. Macha, Y. L. Chong, and I. C. Chen, "Smallholder farmers' intention to adopt microfinance services," *International Journal of*

- Business Innovation Research, vol. 19, no. 3, p. 304, 2019.
- [9] A. Moro, M. Fink, and D. Maresch, "Reduction in information asymmetry and credit access for SMEs," *Journal of Financial Research*, vol. 38, no. 1, pp. 121–143, 2015.
- [10] T. I. Eldomiaty, "Determinants of corporate capital structure," *International Journal of Commerce and Management*, vol. 17, no. 1, pp. 25–43, 2008.
- [11] P. Plawiak et al., "DGHNL: Deep genetic hierarchical network for credit scoring," *Information Sciences*, vol. 516, pp. 401–418, 2020.
- [12] M. M. Ahmad, A. I. Hunjra, and D. Taskin, "Do asymmetric information and leverage affect investment decisions?" *Quarterly Review of Economics and Finance*, vol. 87, pp. 337–345, 2023.
- [13] S. Shi, R. Tse, W. Luo, S. D'Addona, and G. Pau, "Machine learning driven credit risk: A systematic review," *Neural Computing and Applications*, vol. 34, no. 17, pp. 14327–14339, 2022.
- [14] M. Mahbobi, S. Kimiagari, and M. Vasudevan, "Credit risk classification using deep learning," *Annals of Operations Research*, 2023.
- [15] N. 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.
- [16] V. Kanaparthi, "Credit risk prediction using ensemble machine learning," in *Proc. ICICT*, 2023, pp. 41–47.
- [17] Y. Zhong and H. Wang, "Internet financial credit scoring models using deep forest," *IEEE Access*, vol. 11, pp. 8689–8700, 2023.
- [18] A. Vaswani et al., "Attention is all you need," in *Proc. NeurIPS*, 2017.
- [19] X. Huang, A. Khetan, M. Cvitkovic, and Z. Karnin, "TabTransformer: Tabular data modeling using contextual embeddings," 2020.
- [20] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. ACM SIGKDD*, 2016, pp. 785–794.
- [21] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [22] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE TKDE*, vol. 21, no. 9, pp. 1263–1284, 2009.
- [23] N. V. Chawla et al., "SMOTE: Synthetic minority over-sampling technique," *JAIR*, vol. 16, pp. 321–357, 2002.
- [24] J. Davis and M. Goadrich, "The relationship between precision-recall and ROC curves," in *Proc. ICML*, 2006.
- [25] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," 2017.
- [26] S. M. Lundberg et al., "Explainable AI for trees," *Nature Machine Intelligence*, vol. 2, pp. 56–67, 2020.
- [27] M. Sahakyan, Z. Aung, and T. Rahwan, "Explainable AI for tabular data," *IEEE Access*, vol. 9, pp. 135392–135422, 2021.