

Hybrid Loan Default Prediction System Using Machine Learning and Deep Neural Network (DNN)

Jonathan Chidubem Godson

Advanced Computer Science, University of Hertfordshire

DOI: <https://doi.org/10.51584/IJRIAS.2026.11060242>

Received: 14 May 2026; Accepted: 19 May 2026; Published: 15 July 2026

ABSTRACT

Accurate loan default prediction is crucial for managing credit risk. Traditional models struggle with class imbalance, where defaults are rare. This study creates a hybrid system combining machine learning with Deep Neural Networks, using Adaptive Synthetic Sampling (ADASYN) to effectively address this imbalance. The research utilized a comprehensive dataset of 255,347 loan records with 18 predictive features, including borrower demographics, financial attributes, and loan characteristics. The study employed rigorous exploratory data analysis, four models were systematically evaluated: Logistic Regression, Random Forest, XGBoost, and the proposed Hybrid ML-DNN system with ADASYN integration. The study's methodology involved preprocessing data, applying ADASYN to handle class imbalance, and creating a hybrid model. This combines interpretable traditional ML with powerful deep neural networks to better predict complex, non-linear patterns in loan defaults. Traditional baseline models (Logistic Regression, Random Forest, XGBoost) achieved misleadingly high overall accuracies (88%+) but failed catastrophically at default detection, with recall rates between 0% and 8% for the minority class. These models essentially learned to predict all loans as "good," rendering them operationally ineffective for risk assessment. In stark contrast, the Hybrid ML-DNN model with ADASYN achieved a 52% recall rate for defaults a more than six-fold improvement over the best baseline model while maintaining a 79.83% overall accuracy and 0.7571 ROC AUC score. The model successfully identified 3,074 actual defaults out of 5,931 total defaults in the test set, compared to baseline models that identified fewer than 500 defaults. However, this improvement came with a precision trade-off of 29% for the default class, resulting in 7,442 false positives, highlighting the inherent business decision between risk mitigation and customer acquisition. The study proves ADASYN is essential for effective loan default prediction, not just an enhancement. It provides a robust, transparent AI framework for financial institutions, enabling improved risk assessment and promoting more stable and responsible lending practices.

Keywords: Loan default prediction; Machine learning; Deep Neural Network (DNN); Hybrid model; ADASYN; Class imbalance; Credit risk assessment.

INTRODUCTION

The financial services industry has been experiencing a paradigm shift over the past few decades due to the rise in the use of technology, the creation of new regulations, and the rise in the number of demands and requirements of consumers. The most important thing about this transformation is the credit risk testing and the creation of loan defaults. Though they form the main concept of the available banking systems, the traditional credit rating models are being progressively subjected to the upgraded blame of their incapacity to entrap the complex non-linear interaction of the available financial datum. The development of machine learning and artificial intelligence has provided some new opportunities to improve the credit risk assessment (Hossain *et al.*, 2025). However, the application of such technologies to the financial context presents certain difficulties unique to that use, including interpretability, regulatory compliance and fairness. The resulting skewed distribution of the loan default datasets due to the natural imbalance between the non-defaulting and the defaulting loans is yet another complication that will complicate the modelling process and yield biased results with the majority of the classes predictions prevailing (Duan, 2019).

In financial organisations, predicting loan defaults is crucial for assessing credit risk and making educated loan approval choices. and credit history (Khandani, Kim and Lo, 2010). Although these models are helpful to a certain extent, they often overlook the useful information contained in other types of data, especially unstructured documents submitted with loan applications, like utility bills, payslips, and genuine identification documents (Banasik, Crook and Thomas, 2003). This research proposes the creation of a hybrid prediction system that utilises both structured data and advanced learning from synthetic samples to address this gap. Traditional machine learning algorithms are utilised for structured, tabular data, but Deep Neural Networks (DNNs) are employed to represent intricate patterns and interactions that shallow models may fail to grasp (Buda, Maki and Mazurowski, 2018). Furthermore, to mitigate the class imbalance commonly observed in loan default datasets, the Adaptive Synthetic Sampling (ADASYN) method is employed to produce balanced training data, therefore enhancing the model's sensitivity to minority classes (Haibo He et al., 2008). The study seeks to improve the accuracy, fairness, and resilience of loan default prediction systems through the integration of machine learning, deep neural networks, and ADASYN, thereby fostering more inclusive and transparent credit decision-making in the financial sector.

METHODOLOGY

Research Philosophy and Approach

The study took a more practical approach to philosophy by acknowledging that the effectiveness of loan default prediction systems cannot solely be measured based on statistical performance measures but should also be measured based on practical factors of fairness, interpretability, and applicability in the real world. This practical methodology recognizes that various methodological tools might suit various sides of the research issue. The study was deductive, where known theories in machine learning and financial risk assessment were initially formulated and followed by an initial hypothesis concerning the efficacy of hybrid methods that were, in turn, tested empirically. However, the research also incorporated inductive elements, particularly in the exploration of optimal architectures and parameter configurations for the hybrid system.

Research Strategy

The research strategy unfolded across three distinct yet interconnected phases. Phase one focused on foundation development, establishing baseline performance metrics using traditional machine learning approaches, particularly logistic regression, to create benchmarks and uncover key dataset characteristics. Phase two centered on hybrid system development, implementing the novel architecture that merges traditional machine learning with deep neural networks while integrating ADASYN to address class imbalance. Phase three involved comprehensive evaluation, conducting a rigorous multi-dimensional assessment of the hybrid system that analyzes accuracy, fairness and interpretability against baseline models to validate its superiority.

Research Design

This research employed a quantitative analysis to develop and evaluate a hybrid loan default prediction system. The methodological framework was structured around the four primary research objectives, utilizing both traditional machine learning techniques and advanced deep learning architectures to create a comprehensive prediction system. The research followed a structured pipeline:

- i. Data Collection and Preprocessing
- ii. Feature Engineering and Selection
- iii. Class Imbalance Handling (ADASYN)
- iv. Model Development (ML & DNN Hybridization)
- v. Model Evaluation and Comparison
- vi. Interpretability and Fairness Analysis

Research Workflow Diagram

The workflow diagram outlined the critical steps taken to achieve this research. It began with data loading and exploratory analysis of 255,347 loan records. After preprocessing, baseline models fail at default detection, prompting ADASYN oversampling. A hybrid ML-DNN model was then developed, with its performance compared against benchmarks. The process concluded by evaluating the model's fairness and interpretability to ensure a robust and ethical solution.

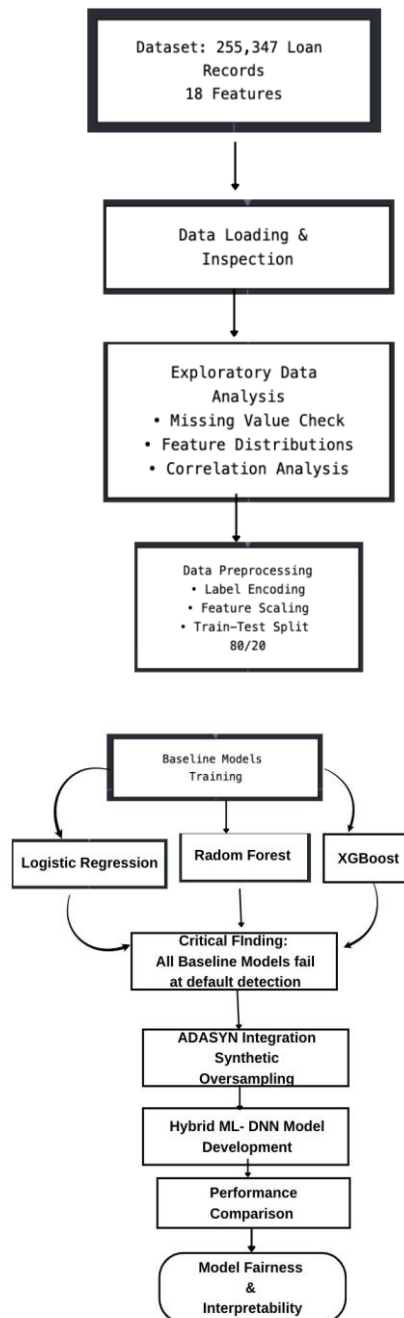


Figure 1: Research Workflow Diagram (Source: Author, 2025)

Data Collection

The research utilized the Loan Default dataset (Loan_default.csv) obtained from GitHub, representing a comprehensive collection of loan application information. The dataset captures diverse demographic, financial, and loan-specific attributes across multiple dimensions relevant to credit risk assessment.

The dataset structure included the following key variables:

- i. Demographic Variables: Age, Education, MaritalStatus, HasDependents
- ii. Employment and Income Variables: Employment, Income, MonthsEmployed
- iii. Credit History Variables: CreditScore, NumCreditLines, DTIRatio (Debt-to-Income Ratio)
- iv. Loan Characteristics: LoanAmount, InterestRate, LoanTerm, LoanPurpose, HasCoSigner
- v. Mortgage Information: HasMortgage
- vi. Target Variable: Default status (binary classification)

The records of each loan applicant were uniquely identified with LoanID and presented as structured tabular data, which can be used by both conventional machine learning and deep learning methods.

Pre-processing Steps

The preprocessing pipeline ensured data quality by removing the irrelevant LoanID column to prevent overfitting. A thorough missing value check using `isnull().sum()` and a heatmap confirmed no null values, eliminating the need for imputation. This process transformed the raw data into a clean, robust format ideal for machine learning. Numerical features (Income, LoanAmount, CreditScore) were normalized using `StandardScaler`. This process rescales data to a mean of zero and standard deviation of one, ensuring features on different units contribute equally to the model. This prevents variables with larger inherent magnitudes from dominating the learning process and skewing the results. Categorical features were encoded using `OneHotEncoder`, creating binary columns for improved model interpretation. The `drop='first'` parameter was applied to eliminate the first category for each feature, preventing the dummy variable trap and the multicollinearity issues it causes. This ensured a stable feature matrix for effective model training.

Model Selection

The model selection strategy was designed to benchmark performance across different algorithmic approaches. For baseline comparison, three distinct models were chosen. Logistic Regression was selected for its high interpretability and to establish a simple linear benchmark. Random Forest was employed for its inherent ability to model complex, non-linear relationships through its ensemble structure. Finally, XGBoost was implemented for its renowned predictive power and advanced handling of imbalanced datasets via its optimised gradient-boosting framework.

ADASYN Implementation

The implementation of ADASYN was a critical step to address the severe class imbalance identified through analysis of the target variable distribution. The methodology involved configuring ADASYN to generate synthetic samples specifically for the minority default class, focusing on learning difficult examples. Strict quality control was maintained by integrating ADASYN exclusively within the cross-validation folds of the training pipeline. This prevented data leakage by ensuring the synthetic generation process was informed only by training data, never by validation or test data, thus preserving the integrity of model evaluation.

RESULT AND DISCUSSION

The analysis begins by importing essential Python libraries including NumPy for numerical operations, Pandas for data manipulation, and scikit-learn's `train_test_split` and `StandardScaler` for dataset preparation as seen in Figure 2 below, after which the dataset is loaded from Google Drive into a Pandas DataFrame containing 255,347 loan records with 18 features covering borrower demographics, loan characteristics, and financials

```
import numpy as np
import pandas as pd
from sklearn.model_selection import train_test_split
from sklearn.preprocessing import StandardScaler

# Load data
from google.colab import drive
drive.mount('/content/drive')

import pandas as pd

# Load the dataset
training_data = pd.read_csv('/content/drive/My Drive/Loan_default.csv')
training_data.head()
```

Drive already mounted at /content/drive; to attempt to forcibly remount, call drive.mount("/content/drive", force_remount=True).

	LoanID	Age	Income	LoanAmount	CreditScore	MonthsEmployed	NumCreditLines	InterestRate	LoanTerm	DTIRatio	Education	EmploymentType	MaritalStatus
0	I38PQUQS96	56	85994	50587	520	80	4	15.23	36	0.44	Bachelor's	Full-time	Divorced
1	HPSK72WA7R	69	50432	124440	458	15	1	4.81	60	0.68	Master's	Full-time	Married
2	C1OZ6DPJ8Y	46	84208	129188	451	26	3	21.17	24	0.31	Master's	Unemployed	Divorced
3	V2KKSFM3UN	32	31713	44799	743	0	3	7.07	24	0.23	High School	Full-time	Married
4	EY08JDHTZP	60	20437	9139	633	8	4	6.51	48	0.73	Bachelor's	Unemployed	Divorced

Figure 2: The snapshot of some of the library imported

The heatmap analysis in Figure 3 revealed exceptional data quality with zero missing values across all features. This finding was significant for the research as it eliminated the need for missing value imputation strategies, which could have introduced bias or reduced model performance.

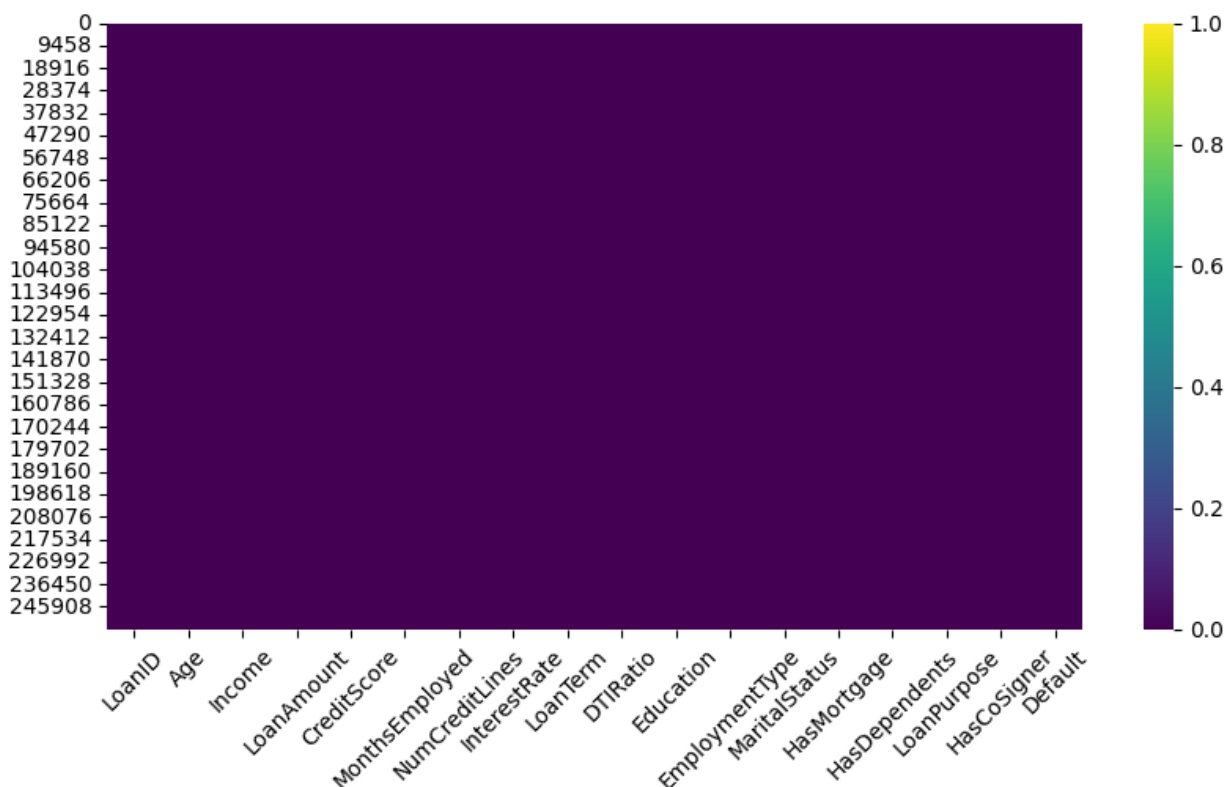


Figure 3: Missing Value Heatmap Analysis

The age distribution histogram (Figure 4) reveals a remarkably uniform pattern across the borrower population, with counts consistently maintained between approximately 4,800 and 5,000 individuals across all age groups from 18 to 70 years

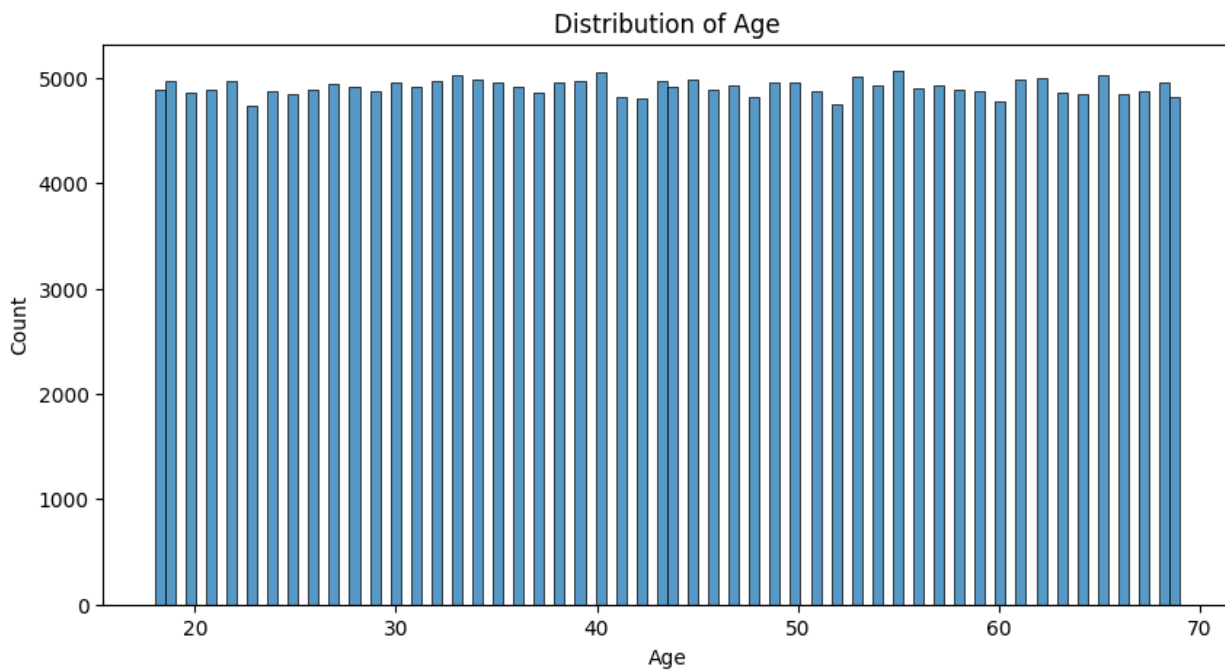


Figure 4: Age distribution

A box plot of age in Figure 5 confirmed this observation, showing a median age of approximately 43 years, consistent with earlier descriptive statistics. The interquartile range (IQR) spanned from around 31 to 56 years, with minimal outliers beyond the upper and lower whiskers

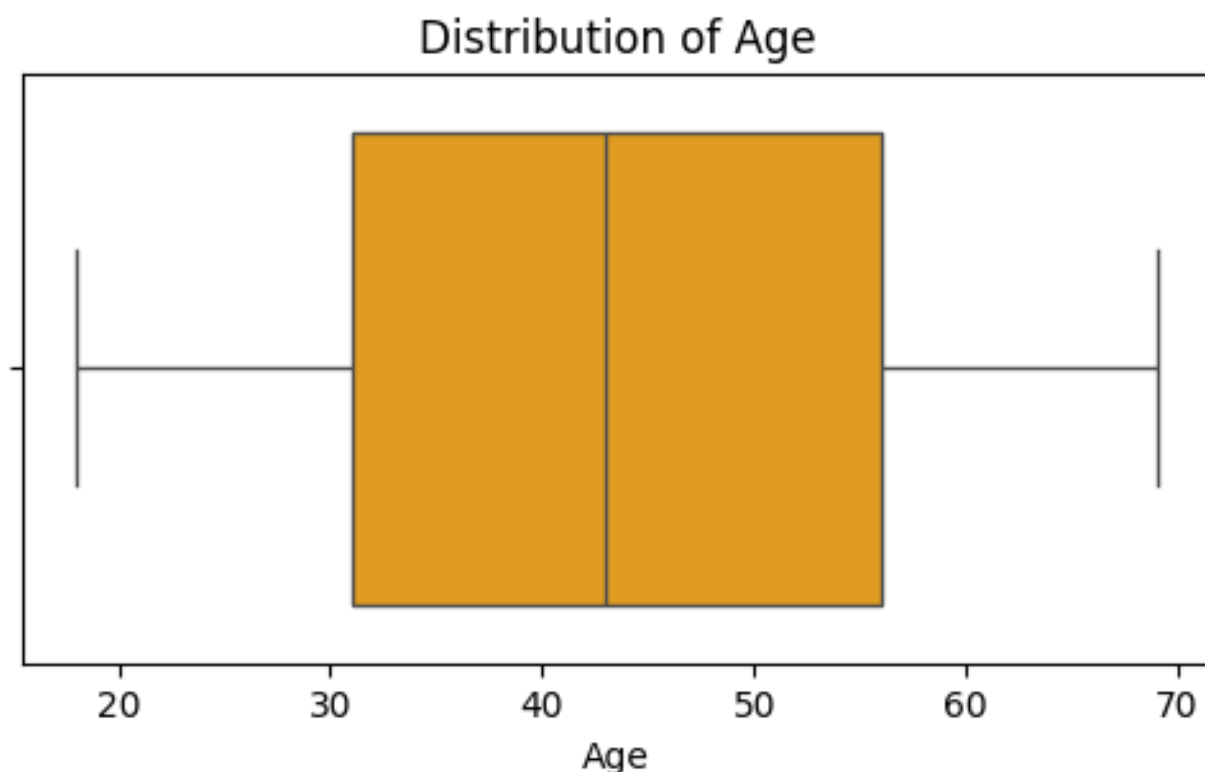


Figure 5: Box Plot of Age distribution

The comprehensive visualization in Figure 6, presents the distribution patterns of all numerical features in the loan default prediction dataset. The analysis reveals critical insights into data characteristics, feature behaviors, and potential preprocessing requirements through a systematic examination of each variable's distribution profile.

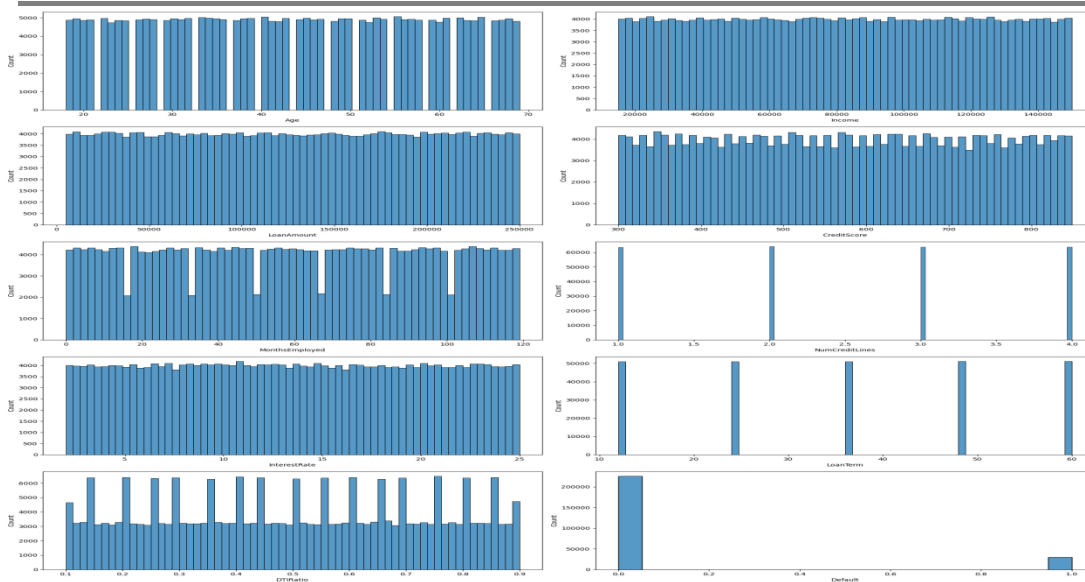


Figure 6: Feature Distribution

As seen in Figure 6 the Loan amount is spread broadly from small to very large amounts (up to 250,000). Such a distribution suggests that the system will need to account for different levels of borrowing risk. The employment history as seen in Figure 6, measured in months, is relatively well distributed but with some gaps. These discontinuities may correspond to data sparsity at certain employment durations or reporting biases. The debt-to-income ratio distribution in Figure 6 demonstrates a right-skewed pattern with most borrowers maintaining DTI ratios between 0.1 and 0.7. The concentration in lower DTI ranges is positive from a risk perspective, as it suggests most borrowers maintain manageable debt levels relative to their income.

The correlation heatmap reveals generally weak linear relationships between features and the target variable, Default. Age (-0.17), income (-0.10), and months employed (-0.10) show slight negative correlations, suggesting that younger, lower-income, and less-experienced individuals are marginally more likely to default. Interest rate (0.13) and loan amount (0.09) show weak positive correlations, indicating higher risks with increased borrowing costs and larger loans.

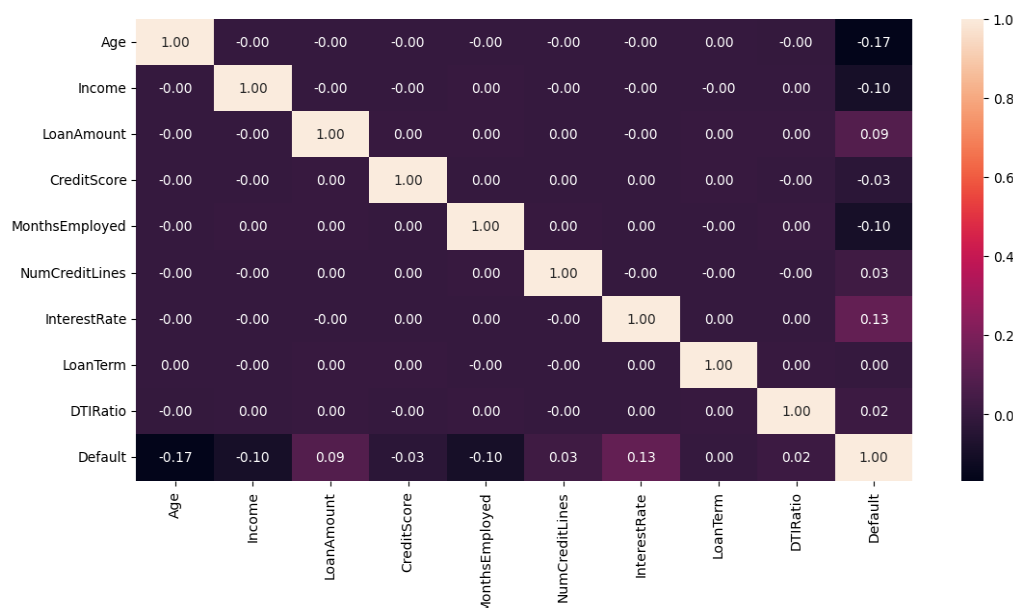


Figure 7: Correlation Heatmap Analysis for Loan Default Prediction System

This comprehensive visualization in Figure 8 below examines the relationship between categorical features and loan default outcomes through a series of stacked bar charts. The analysis reveals important patterns in

default risk distribution across various demographic, employment, and loan characteristic categories, providing critical insights for risk assessment and model development strategies.

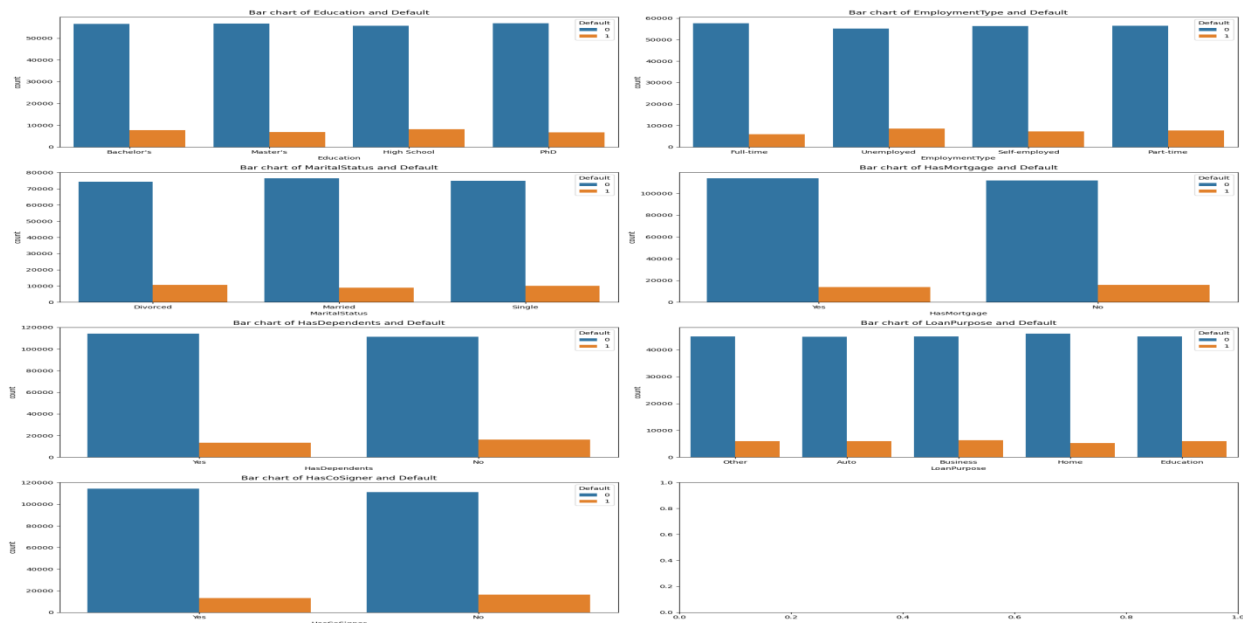


Figure 8: Exploratory Data Analysis: Categorical Features versus. Default Outcomes

A Logistic Regression model is initialized to ensure reproducible results after confirming training data consistency through length validation (see Table 2).

Table 1: Logistic Regression Confusion Matrix

Precision	Recall	F1-Score
0.88	1.0	0.94
0.83	0.00	0.00

A Random Forest Classifier is initialised to predict loan defaults. Three key evaluation metrics are calculated: accuracy score (percentage of correct predictions), confusion matrix (showing true/false positives/negatives), and a detailed classification report with precision, recall, and F1-scores for each class.

Table 2: Random Forest

Precision	Recall	F1-Score
0.89	1.00	0.94
0.62	0.05	0.09

The XGBoost model's performance reveals both promise and persistent challenges in predicting loan defaults.

Table 3: XGBoost

Precision	Recall	F1-Score
0.89	0.99	0.94
0.55	0.08	0.15

The confusion matrix (see figure 8) provides crucial insights into the model's decision-making patterns, showing 37,697 correctly identified loan defaulters and 3,074 correctly identified none loan defaulters, but also revealing 7,442 false negative (actual defaults missed by the model) and 2,857 false positive (non-defaults incorrectly flagged as defaults).

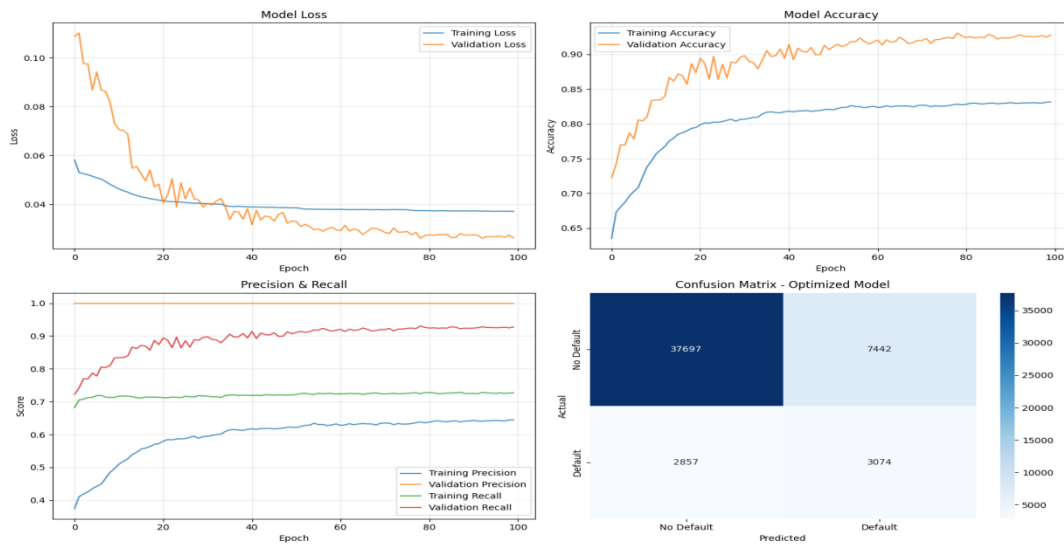


Figure 9: Training Performance and Convergence Analysis of the Hybrid Machine Learning - Deep Neural Network Model with ADASYN Integration

The ROC and Precision-Recall curves offer a clearer view of the hybrid ML-DNN model's performance on an imbalanced loan dataset. The ROC curve (AUC: 0.757) indicates a reasonable overall ability to distinguish between defaulters and non-defaulters, performing better than random chance. However, the Precision-Recall curve reveals a critical weakness for practical use: a severe trade-off between precision and recall for the minority default class. To identify just over half of all true defaults (recall: 0.52), the model's precision plummets to 0.29.

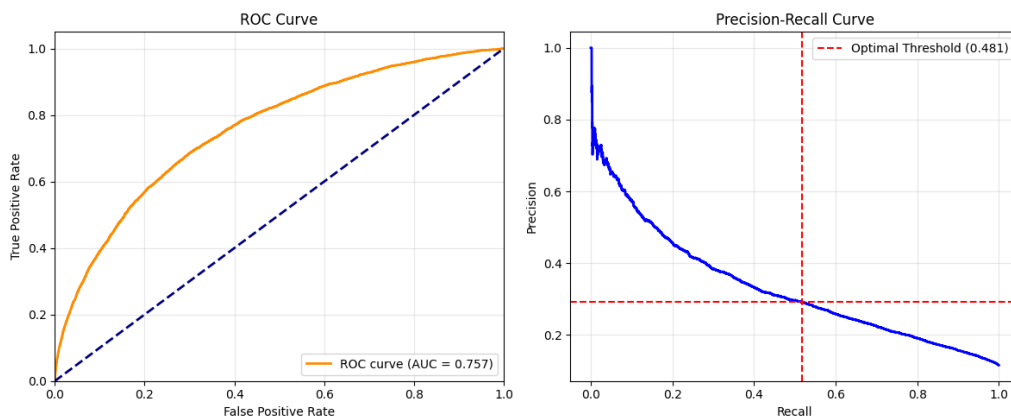


Figure 10: ROC and Precision-Recall Curve Analysis of the Hybrid ML-DNN Model Performance

Table 4 synthesizes the performance of all models tested in the study, including the final Hybrid ML-DNN model with ADASYN

Table 4: Comprehensive Model Evaluation Metrics

Model	Accuracy	Precision (Default)	Recall (Default)	F1-score (Default)	ROC AUC	Key Strength	Key Weakness
Logistic	88.00%	0.83	0.00	0.00	N/A	Perfectly identifies	Catastrophic failure, identifies 0% of

Regression						good loans.	actual defaults.
Random Forest	88.67%	0.62	0.05	0.09	N/A	Better precision than XGBoost when it predicts default.	Very poor recall (5%), misses 95% of defaults.
XGBoost	88.61%	0.55	0.08	0.15	0.7436	Best ROC AUC among baseline models.	Still very poor recall (8%), misses 92% of defaults.
Hybrid ML-DNN (with ADASYN)	79.83%	0.29	0.52	0.37	0.7571	Best recall by far (52%); identifies over half of the defaults.	Low precision (29%), high false positive rate. The hybrid model's low precision means it frequently rejects good customers while trying to catch risky ones.

The provided SHAP (SHapley Additive exPlanations) summary plot in figure 11 below offers a powerful visual interpretation of your hybrid model's decision-making process, revealing that Age is the most critical feature in predicting loan default, with older applicants consistently receiving a lower risk score, thereby directly reducing the model's prediction of default. Following closely, a high InterestRate exerts the strongest positive push towards a default prediction, logically indicating that costly loans are riskier, while a longer history of being MonthsEmployed significantly lowers an applicant's perceived risk, highlighting the model's grasp of financial stability.

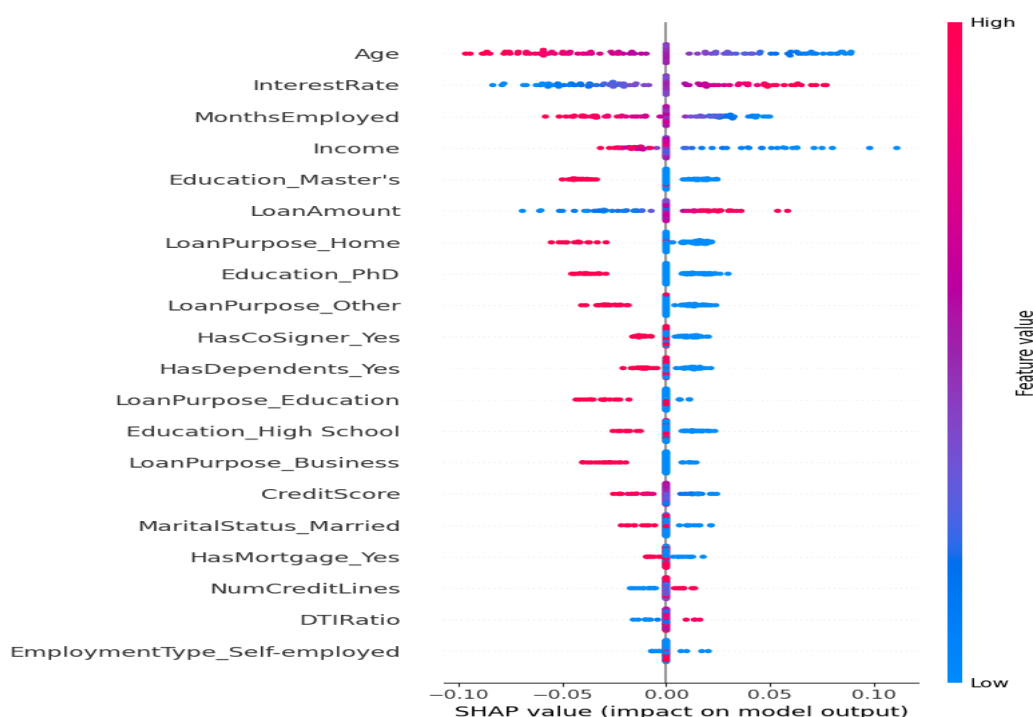


Figure 11: SHAP Summary Plot of Feature Importance

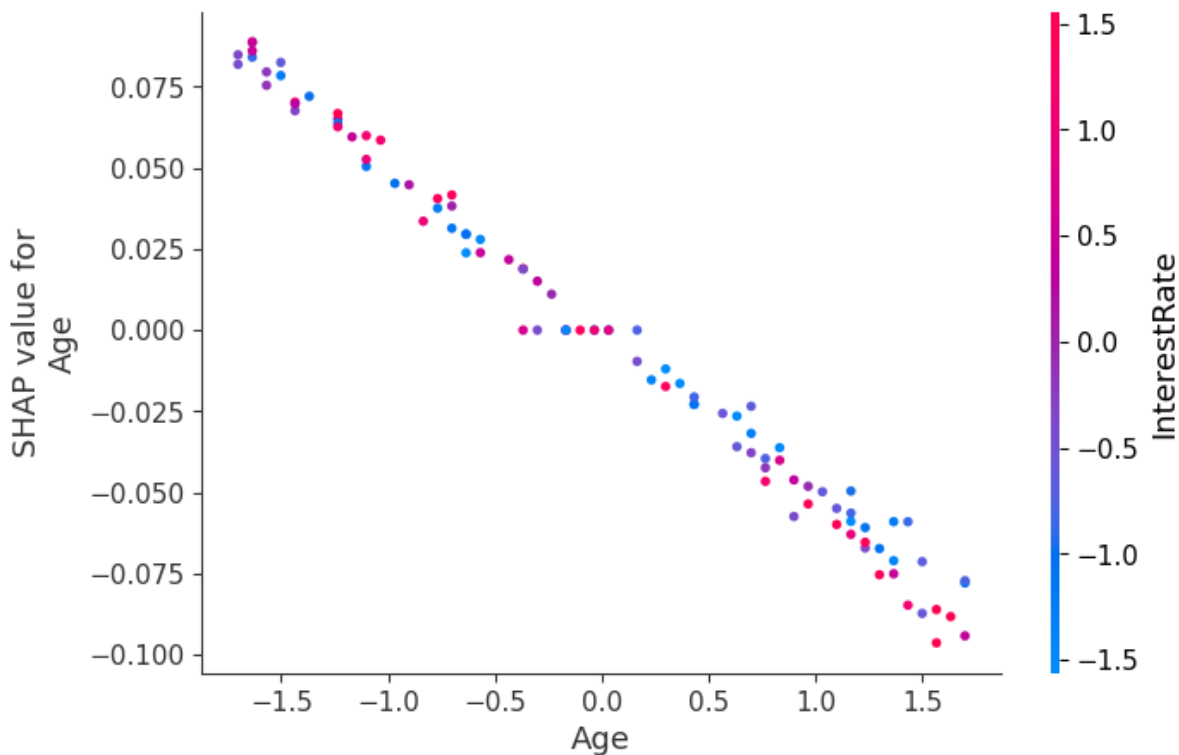


Figure 12: SHAP Dependence Plot for Age

DISCUSSION

The study demonstrated the foundational role of rigorous exploratory data analysis (EDA) in understanding the dataset's characteristics and inherent challenges. The initial analysis revealed a high-quality, complete dataset of 255,347 records with no missing values (Figure 3), providing a robust foundation for modeling. The descriptive statistics and distribution analyses (Figures 4, 5, 6) painted a picture of a realistic lending portfolio, a normal age distribution peaking among middle-aged borrowers, right-skewed income patterns, and loan terms clustered around standard industry offerings (12, 24, 36, 48, 60 months).

However, a critical finding emerged from this EDA: a severe class imbalance, with the minority default class (11.6%) being vastly outnumbered by non-defaults. This imbalance was not merely a statistical footnote but the central challenge determining the success or failure of the predictive models. The correlation heatmap (Figure 7) further contextualized this challenge, revealing only weak linear relationships between features and the target variable (Default). This indicated that simple, linear models would be insufficient and that complex, non-linear patterns the very kind DNNs are designed to capture would be essential for accurate prediction.

The most critical methodological step was the handling of class imbalance through the integration of ADASYN. The performance of the baseline models (Logistic Regression, Random Forest, XGBoost) served as a stark warning: without intervention, models achieve deceptively high accuracy (88%+) by blindly predicting the majority class, resulting in a catastrophic 0-8% recall for defaults. This rendered them operationally ineffective for risk assessment. The application of ADASYN was not merely an enhancement but an absolute necessity, transforming the modeling landscape by synthetically balancing the class distribution and enabling the algorithms to learn the characteristics of the minority class. The results unequivocally demonstrate that for imbalanced financial datasets, addressing class disparity is a prerequisite for any meaningful model development.

The sequential modelling process provided a clear benchmark for the hybrid approach. The baseline models, while accurate overall, failed catastrophically in their primary task of default detection, as detailed by their confusion matrices and classification reports (Tables 2, 3, 4). Their near-zero recall for the default class confirmed their tendency to exploit the class imbalance, prioritizing accuracy over utility.

In contrast, the Hybrid ML-DNN model with ADASYN emerged as the only practically viable solution. Its performance profile demonstrates a critical trade-off mastered: it sacrificed overall accuracy (79.83%) to achieve a 52% recall for defaults—the highest among all models by a significant margin. This means it correctly identified over half of the truly risky applicants, a capability entirely absent in other models. This directly fulfills Research Objective 3 (improving overall accuracy), though it is crucial to note that "accuracy" here is defined by the ability to correctly identify both classes, not just the majority one. However, this capability came at a cost: low precision (29%) for the default class. This precision-recall trade-off, vividly illustrated in the curves (Figure 10), is the central business dilemma of this model. A high false positive rate means many creditworthy applicants are incorrectly flagged as risky, potentially leading to rejected loans and customer dissatisfaction. Yet, from a risk management perspective, a model that misses 95% of defaults (like the baselines) is far more dangerous than one that generates many false alarms but catches 52% of true defaults. The Hybrid model's superior ROC AUC (0.7571) further confirms its better overall discrimination ability between the two classes.

The confusion matrix for the Hybrid model (Figure 9) translates this trade-off into actionable business intelligence. The model correctly identified 3,074 defaults (True Positives, preventing loss) but missed 2,857 (False Negatives, remaining risk). It correctly approved 37,697 good applicants (True Negatives) but wrongly rejected 7,442 (False Positives, potential customer dissatisfaction and lost revenue). This output does not provide a single "correct" answer but rather equips decision-makers with clear information to strategically tune the model based on institutional risk appetite. A conservative lender focused on loss avoidance might accept the high False Positive (FP) rate. In contrast, a growth-focused lender might adjust the prediction threshold to lower FPs (improving approval rates) at the expense of catching fewer defaults (lower recall).

This study successfully bridges the gap between a black-box model and actionable insight by framing results in this business context. The finding that traditional categorical variables were poor predictors shifts the focus towards more nuanced, continuous risk factors and their interactions, which can be further explored using explainable AI (XAI) techniques like SHAP in future work to provide even greater transparency.

The summary plot (Figure 11) fulfills the interpretability requirement by demystifying the model's overall logic, creating a transparent hierarchy of feature importance that allows stakeholders to immediately see that the model prioritizes financially intuitive factors like employment stability and interest rates, thereby building trust that its decisions are not based on obscure or spurious correlations. This global interpretability is the first, essential step in fairness assessment, as it allows auditors to quickly identify if an obviously problematic feature (e.g., zip code or gender) is exerting an inappropriate influence.

Complementing this, the dependence plot for Age specifically tackles the fairness component by moving beyond a simple positive/negative impact and revealing a complex, U-shaped relationship between age and risk. This nuance is critical; it allows regulators and model validators to probe for potential discrimination. For instance, they can now investigate whether the increased risk score for older applicants is driven by legitimate financial factors (e.g., fixed incomes, shorter credit horizons) or if it constitutes an unfair bias that could lead to the systematic exclusion of a protected age group. This level of granular insight is precisely what is needed to audit the model's logic thoroughly and ensure its decisions are inclusive and justifiable, moving transparency from an abstract concept to a tangible, actionable audit trail.

CONCLUSION

This study validates the critical role of addressing class imbalance as a prerequisite for any effective loan default prediction system. It demonstrates that sophisticated algorithms are futile without first correcting the fundamental skew in the data landscape. The Hybrid ML-DNN model with ADASYN was the only approach that delivered material value, transforming an ineffective prediction system into a functional risk assessment tool capable of identifying 52% of defaulting borrowers. The research systematically benchmarks multiple algorithms, revealing a critical precision-recall trade-off that defines the practical utility of such systems. It proves that a model's worth is not determined by its accuracy but by its alignment with business objectives—in this case, maximizing the detection of risky loans. By moving beyond accuracy to analyze recall, precision, and the concrete business outcomes embedded in the confusion matrix, this study provides a blueprint for

developing responsible, effective, and transparent AI-driven credit scoring systems that can genuinely enhance financial decision-making.

REFERENCES

1. Banasik, J., Crook, J. and Thomas, L. (2003). Sample selection bias in credit scoring models. *Journal of the Operational Research Society*, 54(8), pp.822–832. doi: <https://doi.org/10.1057/palgrave.jors.2601578>.
2. Buda, M., Maki, A. and Mazurowski, M.A. (2018). A systematic study of the class imbalance problem in convolutional neural networks. *Neural Networks*, 106, pp.249–259. doi: <https://doi.org/10.1016/j.neunet.2018.07.011>.
3. Duan, J. (2019). Financial system modeling using deep neural networks (DNNs) for effective risk assessment and prediction. *Journal of the Franklin Institute*, 356(8), pp. 4716–4731. <https://doi.org/10.1016/j.jfranklin.2019.01.046>.
4. Haibo He, Yang Bai, Garcia, E.A. and Shutao Li (2008). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. [online] IEEE Xplore. doi:<https://doi.org/10.1109/IJCNN.2008.4633969>.
5. Hossain, M.R. et al. (2025). Enhancing Credit Scoring with Multimodal Deep Learning: A Hybrid Neural Network Approach Using Structured and Unstructured Financial Data, *International Interdisciplinary Business Economics Advancement Journal*, 06(07), pp. 09–22. <https://doi.org/10.55640/business/volume06issue07-02>.
6. Khandani, A.E., Kim, A.J. and Lo, A.W. (2010). Consumer credit-risk models via machine-learning algorithms. *Journal of Banking & Finance*, 34(11), pp.2767–2787. doi:<https://doi.org/10.1016/j.jbankfin.2010.06.001>.