

Client Churn Prediction in Retail Finance Institutions using Optimized XGBoost and Explainable AI Methods

Naina Kokate¹, Dhruv Parekh², Nisarg Parmar³ and Parth Patil⁴

¹⁻⁴Vishwakarma Institute of Technology, Pune, India

Email: naina.kokate@vit.edu, dhruv.parekh24@vit.edu, nisarg.parmar241@vit.edu, parth.patil241@vit.edu

Abstract— Customer churn prediction has developed into a field of study. Experimental literature in banking industry has shown that the costs of retaining customers are significantly lower than customer acquisition costs are. As a result, the paper provides the proposal of a machine-learning model that would predict a possible churner at the initial stage based on the demographic, behavioral, and transactional data. The research includes the processing of data, features selection, and the implementation of the different predictive models, such as “Logistic Regression (LR)”, “random forest (RF)”, and “gradient-boosting methods (GB)”. In an effort to increase the level of interpretability, the framework adopts “SHAP” analysis to explain the most pronounced features that can be attributed to customer churn. The suggested system has the potential to be a utility to the banks and non-banking financial corporations (NBFCs), since it will enable the identification of high-risk customer groups and allow understanding which factors are major contributors to churn. This will allow the financial institutions to curb the churn levels within the high-value and target customers

Index Terms— first term, second term, third term, fourth term, fifth term and sixth term.

I. INTRODUCTION

Loyalty of the customer is become the key issue of banking institutions in the modern fiscal environment. Acquisition initial cost is by far a huge cost compared to the cost of retaining existing client and hence churn prediction is a mandatory success factor that banks that seek long term success should integrate. The conventional churn management models used in the banking sector has, to a great extent, been based on the rule-based or statistical models that have failed to sufficiently reflect the non-linear association between the attributes of customers. Growth of customer information, on a geographic, transactional, and behavioral scale has allowed the use of machine learning and artificial intelligence to predictive analytics in the financial industry. The literature related to churn prediction in the past decade has largely concerned telecommunications, e-commerce, and services based on subscriptions, such as OTT systems. On the other hand, banking churn prediction research is limited, and most of the existing models are in interpretable, which restricts their usage to management decision-making activities. The black-box systems known as “Random Forest (RF)” and “Gradient Boosting (GB)” are usually the better options because they have better predictive accuracy but tend to be poorer in explainability. This is a trade-off between accuracy and explainability, which is a major obstacle to AI-driven solutions in such highly-regulated and controlled industries of the financial sector, where interpretability, accountability, and transparency are essential.

Recent developments in “Explainable Artificial Intelligence (XAI)” suggest methods that can allow interpreting machine-learning models without affecting accuracy or performance. “SHAP (SHapley Additive exPlanations)”, and “LIME (Local Interpretable Model-agnostic Explanations)” are approaches that expound on the input of every feature to an approximation. A banking situation can use the methods to understand the different meaningful drivers of churn e.g. “account balance”, “tenure and frequency of transactions” hence, transforming predictions into insights to be acted upon. In spite of the methodological advancement, there exists a gap in the research literature as to integration of the predictive accuracy and its practical implementation in the banking sector. Besides, a lack of research investigating the use of churn models to inform data-driven retention strategies that balance commercial intentions and regulatory requirements persists

II. LITERATURE REVIEW

The earlier churn prediction models used old statistical methods such as LR and “decision trees(DT)”. These models offered interpretability but had a limited predictive capability. Buckinx and Van den Peol[2] seminal work demonstrated that LR could serve as a base for financial churn forecasting. significant rise in customer data, more ML techniques began to perform better traditional models such as LR and Trees. In the banking industry, Zhu and Liu [1] conducted comparative analysis of RF and LR. This study shows ensemble models have a higher accuracy as compared to traditional models because of their nonlinear relationships among various customer attribute columns such as age, balance and tenure of the customer. This study places importance on demographical and transactional metrics like age of customer, account balance amount and number of products held. One of the key challenges is data imbalance, where the number of non-churning customers exceeds the churners. Gkonis and Tsalakos[3] specifically addressed this issue in the banking domain using DNN architecture. Their work combining oversampling with optimization strategies such as ‘Radam’ improves recall and stability for minority(churn) classes. Sikri et al introduces a Ratio-based balancing technique that improved the classification accuracy compared to the traditional oversampling and under sampling methods. Both of these works conclusively prove that preprocessing rebalancing and feature normalization influences model performance in churn prediction models. Across multiple domains, ensemble models have shown better results. Han Zhu and Liu[1] showed that Random Forest and “XGBoost” have a better performance as compared to simple classifiers. In addition, Khattak[2] proposed BiLSTM – CNN composite model simultaneously captures sequential and spatial dependencies. This hybrid model achieved high performance of 81% by using convolutional feature extraction. These advancements indicate the growing preference for hybrid and ensemble architectures that maintain performance with interpretability.

Beyond banking, cross-domain insights have provided insights for churn prediction models. In 2015 Renjith[23] proposed a retention framework for B2C e-commerce companies which combined churn prediction with personalized recommendation systems. This paper combined predictive models with actionable retention insights. Similarly, Gattermann-Itschert and Thonemann[5] showed that Random Forest churn prediction can be applied in B2B settings, where proactive interventions based on predicted churn probabilities reduced the number of churners and increased the overall profitability.

It contains detailed customer, account, transaction, loan, card, and service-feedback. Attributes at customer level will be unique customer identifier, first and last name, age, gender, address, city, contact number, and email address. Details of the account include the account type, account balance, date of opening the account and the last date on which an account was used. Transaction records include transaction id, date of transaction, the type of transaction, the amount of transaction, and balance in the account after every transaction. Branch association has been shown by the help of a branch identifier. Loan specifics are the loan ID, the loan amount, the interest rate, loan term, the date of loan approval or rejection, and the loan status. The card related details include card ID, type of credit card, credit limit, credit card balance, minimum amount required to be paid and payment due date, date of last credit card payment, and the rewards points earned. Also, customer feedback is recorded with the help of feedback ID, feedback type, the status of the resolution and the date of the resolution. There is also an indicator of anomaly in the dataset to determine irregular or suspicious patterns in the financial and transactional records.

III. METHODOLOGY

This study aims to develop a customer churn prediction model for retail banking, comprising of transactional, behavioral and credit-based data columns from a Kaggle Dataset. This particular study focuses on three algorithms LR, RF and “XGBoost” and compares their predictions on a preprocessed and feature-engineered

dataset. In addition, this study to make the results of the prediction more accessible to all stakeholders integrates “SHAP” and “LIME” along “Explainable AI (XAI)”. The entire process consists of data preprocessing and cleaning, feature construction and dataset construction, model development (using LR, random forest, XGBoost), model classification using standard classification metrics.

The Kaggle dataset contains customer-level financial information across a varied timeline. Each and every record corresponds to an individual customer and records account activity, product relationships and other relevant and indicators. The target variable is going to be binary -Churn (1 = churned, 0= retained). This is based on account closure, dormancy (based on transaction carried in a specified time frame) or transfers to a particular bank in a certain time-frame. The key raw attributes of the database are

Demographic data: Age, gender, customer tenure, location.

Account metrics: Average balance, account type, and account age.

Transactional data: Number and value of debit and credit transactions, monthly inflows/outflows, frequency of salary credits.

Compliance indicators: KYC/AML flags, account freeze history.

Behavioral text data: Transaction descriptions

The dataset contains a minor quantity of missing data (~ 2 to 3 %), skewed numeric distributions and imbalanced target class (around 12% churners) which was addressed during preprocessing.

TABLE I: MODIFIED DATASET COLUMNS

Attribute	Type	Description	Example
Customer ID	ID	Unique Customer Identifier	C1023
Avg Monthly Balance	Numeric	Mean Balance over 3months	45200
Balance Trend 3M	Numeric	3-months slope	5800
Debit Inflow Mean	Numeric	Monthly debit inflows	38000
Debit Outflow Mean	Numeric	Monthly debit outflows	35000
Salary credit flag	Binary	Salary inflow present	1/0
Product Count	Numeric	Number of products flow	3
Loan Count	Numeric	Active Loan products	1
Credit Utilization	Numeric	Credit Used	0.72
Credit Score	Numeric	Customer's credit score	695
KYC Issue Flag	Binary	KYC/AML issues	0
Account Freeze Flag	Binary	Freeze in last 6 months	0
Outflow Ratio	Numeric	Outflows	0.42
Transaction PCA1-PCA3	Numeric	Top 3 PCA components	0.12
Churn	Binary	Target Variable	0/1

The data preprocessing was done to ensure data consistency, remove noise and handle missing values. All of these steps and prepare it for data input in the model. Python libraries such as (pandas, scikit-learn). The first step in this process was data cleaning under which all duplicate records removed using unique customer identifiers. The absent numeric values were imputed(replaced) using median. The missing categorical values labeled as “unknown”. The outliers in monetary columns were categorized as 1st and 99th percentile to keep within meaningful and ranges. The next step involved feature encoding and transformation. Certain categorical values such as account type were one-hot encoded. Binary indicators such as KYC were transformed as 0/1 indicators. The continuous variables such as (balances, transaction amounts) were standardized using z-score normalization. Next textual feature transaction was carried out. Transaction descriptions were processed using sentence – transformers to generate semantic embeddings. “Principal Component Analysis (PCA)” reduced these embeddings to the top 5 components. Next the overall lack of balance was addressed using “SMOTE (Synthetic Minority Oversampling Technique)” at the time of model training ensuring balanced approach between churned and retained classes.

$$z = \frac{x - \mu}{\sigma}$$

Z-score formula

Where:

- z— Standardized value of the numerical attribute after normalization
- x— Original value of the numerical feature
- μ — Mean (average) value of the feature computed over the dataset
- σ — Standard deviation of the feature

This processed dataset was used as the analytical foundation for all subsequent modelling. Three classification algorithms were implemented to model churn: LR, RF and “XGBoost” algorithm. A 70-30 train-test split ensured proportional representation of churn classes. Hyperparameter tuning was performed using 5-fold cross-validation optimizing F1-score to balance precision and recall value. The LR serves as a base for the model due to its ability to process linear data between independent relationships and churn probability.

$$P(y = 1 | x) = \frac{1}{1 + e^{-(\beta^0 + \sum_{j=1 \text{ to } p} \beta_j x_j)}}$$

LR Formula

Here,

- $P(y = 1 | x)$ — Probability that a customer will churn
- y — Target variable (1 = churned, 0 = retained)
- x — Input feature vector representing customer attributes
- β_0 — Intercept (bias term) of the model
- β_j — Coefficient associated with the j^{th} feature
- x_j — Value of the j^{th} feature
- p — Total number of input features
- e — Euler’s constant

The model is enhanced using log-likelihood (cross-entropy) loss:

$$L = -\frac{1}{n} \sum_{i=1}^n [y_i \log(\hat{p}_i) + (1 - y_i) \log(1 - \hat{p}_i)]$$

Where:

- L — Total log-loss (cross-entropy loss)
- n — Number of training samples
- y_i — Actual churn label of the i^{th} customer
- $\hat{p}_i$ — Predicted probability of churn for the i^{th} customer

Regularization (L2) was applied to prevent overfitting.

Random Forest technique handles non-linear data and recognizes complex relations between multiple features using ensemble learning.

$$\hat{y} = \text{mode} \{ h_1(x), h_2(x), \dots, h_T(x) \}$$

Random Forest Formula

Where:

- $\hat{y}$ — Final predicted class label
- $h_t(x)$ — Prediction produced by the t^{th} decision tree
- x — Input feature vector
- T — Total number of trees in the forest

mode — Majority vote among all tree predictions

Each and every tree is trained on a bootstrap sample with random feature selection at every split reducing the overall variance and overfitting.

XGBoost enhances prediction accuracy through sequential gradient boosting with regularization, efficiently capturing the non-linear dependencies and the feature interactions

$$L(t) = \sum_{i=1 \text{ to } n} [g_i f_t(x_i) + (1/2) h_i f_t(x_i)^2] + \Omega(f_t)$$

XGBoost formula

Where:

- $L^{(t)}$ — Objective function value at boosting iteration t
- n — Number of training samples
- x_i — Feature vector of the i^{th} customer
- $f_t(x_i)$ — Prediction score from the newly added tree

- g_i — First-order gradient of the loss function
- h_i — Second-order gradient (Hessian)
- $\Omega(f_t)$ — Regularization term controlling model complexity

The model was evaluated using test dataset and cross-validation metrics focusing on accuracy (overall correctness of the results), precision (proportion of true churn among the predicted churn), recall (ability to detect churners), f1-score (harmonic mean of precision and recall), ROC-AUC (discrimination capability).

IV. ADVANTAGES

The proposed bank churn prediction system combines interpretability, predictive accuracy, and scalability. These three dimensions were not treated together in the prior studies. This system moves away from the traditional black-box models and integrates LR, RF and “XGBoost” to achieve a balance between explainability and performance. The use of interpretable models allows the banks to understand each and every customer attribute and metric affects the churn probability of that particular customer. This ensures transparency and compliance. The multi-model comparison strengthens result reliability and minimizes overfitting thus making the system more stable across multiple customer profiles.

The key difference between this study and other comparable studies carried out recently is the comprehensive data preprocessing and feature engineering pipeline. The dataset was systematic cleaned, encoded, scaled and balanced through SMOTE to remove class imbalance which is the common problem in the previously carried out studies. The inclusion of bank specific features like balance-to-salary ratio and product tenure index, further enhances the model’s discriminatory power. The result has improved recall and fairness towards the minority (churn) class. This is so that the system not only predicts churn precisely but also it is efficient in identifying a high-risk customer. Reproducibility is also promoted by the structured pre-processing workflow, where careful validation and benchmarking of models in future studies can be carried out with ease and efficiency. With its lightweight model, it has a high degree of scalability and cost efficiency meaning it can easily be deployed even at small scale, unlike resource-intensive deep-learning alternatives. The analysis model, which consisted of accuracy, precision, recall, F -1 score and ROC -AUC, presented a balanced and realistic evaluation of the model performance. Banks can use these insights to be proactive to target high-revenue/high-risk customers, resource marketing and improve customer retention and revenue stability.

V. CHALLENGES

A. Scalability and Adaptability

The system has been built on a dataset of a small scale and its usefulness in practice, in high-volume operational environments has not been empirically tested. The financial institutions across the world have millions of clients and operate in high frequency transactions and therefore the model must maintain the same performance even when the load is quite high. In addition, the model needs to be agile- it has to adapt to new financial product, customer behaviors and changes in regulatory demands without affecting the stability or accuracy. Therefore, it is not only a contingent characteristic but a necessary condition of a long-lasting reliability of continuous evolution.

B. Dataset Generalization

The model has been trained on a dataset that is selected to fit one banking setting, and although it has been shown to exhibit high performance in that particular setting, it is unclear whether it can be extended to other financial institutions. Distributional changes brought about by differences in transaction organization, customer characteristics and regulatory specifications among the banks can undermine predictive accuracy. Based on this, a thorough analysis on a wide array of external data is necessary in order to determine the stability of the model, and to verify its relevance in a wider range of banking contexts.

C. Feature Stability Over Time

The time-varying behavior of customers in terms of financial activity and changing regulatory regimes in modern credit-risk modelling are the cause of what is more generally known as feature drift, the progressive loss of predictive power of model inputs over time. It has been proved empirically that any minor change in the distribution of the important variables can trigger quantifiable reductions in the out-of-sample accuracy, thus undermining the validity of risk measurements at the operational level. One possible solution to reduce this degradation is to have a strong governance policy that includes periodic retraining of the underlying algorithms, periodic performance measurement, and feature-selection adjustment processes that can automatically re-weight

or replace outdated predictors. All these practices can guarantee that the models are kept in line with the current economic environment and regulatory demands and, therefore, keep predicting with a high degree of fidelity in real-life applications.

VI. NEED FOR THIS SOLUTION

There is existing critical need for financial institutions across the world to predict their churn rates to avoid high customer attrition rates. While there has been significant research in churn prediction across varied industries including the banking sector but there was a comprehensive lack towards explaining these insights to the relevant teams. This study fills that gap by presenting a model which integrates the accuracy, interpretability of ML models and explainability of the XAI tools such SHAP and LIME. With this model using XAI to display and present the results it is expected to fill the void between results generation and communication.

VII. RESULTS AND DISCUSSIONS

The performance of the three classification models LR, RF, and Tuned XGBoost was evaluated using accuracy, precision, recall, F1-score, and ROC–AUC on the test dataset. The comparative performance is summarized in Table 3.

TABLE III: COMPARATIVE MODEL PERFORMANCE

Model	Accuracy	ROC–AUC	Precision	Recall	F1–Score
LR	0.51	0.514	0.49	0.52	0.51
Random Forest	0.54	0.559	0.53	0.50	0.51
XGBoost (Initial)	0.813	0.894	0.80	0.82	0.81
Tuned XGBoost (GridSearchCV)	0.845	0.913	0.83	0.85	0.84

A. Logistic Regression

The accuracy of the LR model was 0.51 and ROC-AUC was 0.514 which is near random classification. The level of precision (0.49) and recall (0.52) were small, which means that it has a low discriminative ability.

[256,260
230,254]

shows balanced misclassification across both classes. This suggests that linear relationships between input features and churn behavior are weak, and the model failed to capture the non-linear dependencies present in customer transaction and behavioral patterns.

B. Random Forest

The RF classifier marginally improved accuracy to 0.544 with an ROC–AUC of 0.56, though its training accuracy of 1.0 revealed overfitting. The confusion matrix

[303,213
243,241]

indicates a slight bias toward predicting the majority class (non-churn). Despite of using ensemble averaging, Random Forest had a problem with the high dimensionality and class imbalance of the dataset. This led to limiting its generalization capability.

C. Initial XGBoost Model

The initial XGBoost model performed substantially better than the baseline algorithms, achieving 0.813 accuracy and 0.894 ROC–AUC.

[417,99
88,396]

The model shows an equal distribution of churners and non-churners. This is because the model has the ability to take non-linear interaction and explain the relationship of features with the use of gradient boosting. However, the use of default parameters led to slight variations in the recall and precision which is a pointer to the fact that further optimization and general performance improvement can be achieved.

D. Tuned XGBoost Model (Optimized via Grid SearchCV)

The tuned “XGBoost” model was the most successful in all the evaluation metrics after hyperparameter optimization with the help of the GridSearchCV. The model showed an accuracy of 0.845, an ROC-AUC of

0.912 and balanced precision (0.83) and recall (0.85). The F1-score is 0.84 which means that the number of false positives is balanced with the number of false negatives. The associated confusion matrix

$$\begin{bmatrix} 433,83 \\ 72,412 \end{bmatrix}$$

has a high predictive reliability, making it detect most cases of churn as well as non-churn within the dataset.

The hyperparameter optimization procedure of parameters including; learning rate, maximum depth of the tree as well as regularization variables assisted in reducing overfitting and enhanced generalization. The tuned “XGBoost” model had about 65 percent higher accuracy and 55 percent higher ROC-AUC than the baseline LR, which proves it to be the best model in the possibility to predict complex behavioral and financial studies in the data set.

E. Comparative Discussion

The fact that the performance of LR and RF and finally of XGBoost gradually improved indicates that non-linear modelling and gradient-based optimization are significant in prediction of churn. The inadequacy of linear models to account multi-dimensional behavioural data is highlighted by bad performance of the LR. RF enhanced variance reduction but did not have parameter fine-tuning, which resulted in poor generalization. On the contrary, XGBoost that incorporated iterative boosting and regularization was able to reflect complex interaction between features eventually leading to the markedly higher predictive power.

The ROC-AUC value was high and equal to 0.912 indicating a high degree of discriminative power of the model together with the near-equal and equal precision and recall that indicates that the model is also consistent with both classes. This ratio is important in banking use cases because false negatives (underrecognized churners) can result in losing a customer and false positives (mistakenly flagged customers) can result in retention spending that is not necessary.

VIII. FUTURE SCOPE

Although the suggested churn predictive framework can be accurately represented with a historical banking data, a number of improvements can be made to make this framework applicable and efficient in the context of an actual banking setting.

First, the model can be incorporated with real-time core banking systems to be allowed to monitor the behavior of the customers continuously. Therefore, processing live transactional and interaction data streams enables the system to issue early churn risks indicators and facilitate automatic retention processes that enable banks act proactively instead of reactively.

Second, when the predictive power of the model is required to be enhanced, one can do it by expanding the feature space. Other behavioral and experience signals that can be used in future work include mobile and internet banking usage habits, log of customer complaints, call-center conversations and sentiment analysis based on textual information. Such unstructured and multi-source data can be included and could offer a more extensive insight into customer engagement and dissatisfaction.

Third, another potential line of research is the implementation of complex machine learning methods. Deep learning structure, ensemble stacking approach and hybridization can be considered to identify complicated non-linear correlations in customer behavior. Moreover, with AutoML pipelines, it is possible to support the best model selection and hyperparameter optimization in addition to enhancing reproductivity and resilience. Moreover, the system can be improved by scaling cloud-based deployment to the Amazon Web Services, Microsoft Azure, or Google Cloud. Migration to clouds would allow effectively processing a huge amount of customer data, quicker model inference, and better access to its functions on the organizational units. The possibility of the API based deployment can also enable easy integration with other system like Customer Relationship Management (CRM) systems and marketing automation software.

Lastly, the future research can be devoted to the creation of the customer retention strategy engine based on the churn predictions. This type of engine would be able to suggest customized retention measures based on rule-based or machine learning-driven methods, depending on the risk characteristics of individual customers. Simulation or cost-benefits analysis models can be used to assess the financial efficiency of various retention strategies, thus resolving the issue of disconnecting predictive analytics and actionable business decisions.

IX. CONCLUSION

This paper explores customer churn in the banking industry and the strategies will follow the behavioral, financial and institutional approaches. The authors do not intuitively turn churn prediction into an algorithmic

solution, instead placing it in a wider socio-economic framework that brings in trust relationships, regulatory concerns and customer engagement. The evaluation of variables (e.g., the use of credit and transaction trends, the presence of violations in compliance and the general state of the customer-bank relationship) references proves that churn often indicates significant alterations in the financial status and engagement of a client. The main implication, therefore, is that behavioral and regulatory indicators should be incorporated in the predictive models in order to be more specific in capturing the relations that spur customer defection.

ACKNOWLEDGMENT

We would like to express our gratitude to Dr. Sandip Shinde, Head of the “Department of Computer Engineering” “Vishwakarma Institute of Technology” for providing us this invaluable opportunity to make and present this project. His academic and administrative help guided us to complete this study.

We are deeply thankful to Prof. Naina S Kokate, “Vishwakarma Institute of Technology”, whose guidance and invaluable contribution helped us to complete this study. Her support to us throughout helped us understand institutional expectations. We would also like to extend our heartfelt thanks to Prof. Dheeraj Jadhav, “Vishwakarma Institute of Technology”, for his guidance and suggestions on regular basis throughout the timeline of this project

REFERENCES

- [1] Zhu, H. (2024). Bank Customer Churn Prediction with Machine Learning Methods. *Advances in Economics, Management and Political Sciences*, 69(1), 23–29. <https://doi.org/10.54254/2754-1169/69/20230773>
- [2] Khattak, A., Mehak, Z., Ahmad, H. et al. Customer churn prediction using composite deep learning technique.
- [3] Gkonis, Vasileios & Tsakalos, Ioannis. (2024). Deep Dive Into Churn Prediction in the Banking Sector: The Challenge of Hyperparameter Selection and Imbalanced Learning. *Journal of Forecasting*, 44. 281-296. 10.1002/for.3194.
- [4] (Enhancing Customer Churn Prediction using Machine Learning and Deep Learning approaches with principal component analysis, 2023)
- [5] Gattermann-Itschert, Theresa & Thonemann, Ulrich. (2022). Proactive customer retention management in a non-contractual B2B setting based on churn prediction with random forests.
- [6] P A, J., & N, A. (2023). Bank Customer Churn Prediction. *Indian Journal of Data Mining*, 2(2), 1–5.
- [7] Xiping, S., & Yi, W. (2024). Research on Integrated Learning Algorithm Model of Bank Customer Churn Prediction. *International Journal of Database Management Systems*, 16(1/2/3/4/5), 01–13.
- [8] Ashraf, R. (2024). Bank Customer Churn Prediction Using Machine Learning Framework. *Journal of Applied Finance & Banking*, 65–109
- [9] Chen, P., Liu, N., & Wang, B. (2023). Evaluation of Customer Behaviour with Machine Learning for Churn Prediction: The Case of Bank Customer Churn in Europe. *Proceedings of the International Conference on Financial Innovation, FinTech and Information Technology, FFIT 2022, October 28-30, 2022, Shenzhen, China. Proceedings of the International Conference on Financial Innovation, FinTech and Information Technology, FFIT 2022, October 28-30, 2022, Shenzhen, China.*
- [10] Baghla, S., & Gupta, G. (2022). Performance Evaluation of Various Classification Techniques for Customer Churn Prediction in E-commerce. *Microprocessors and Microsystems*, 94, 104680.
- [11] Černevičienė, J., & Kabašinskas, A. (2024). Explainable artificial intelligence (XAI) in finance: a systematic literature review. *Artificial Intelligence Review*, 57(8).
- [12] Domingos, E., Ojeme, B., & Daramola, O. (2021). Experimental Analysis of Hyperparameters for Deep Learning-Based Churn Prediction in the Banking Sector. *Computation*, 9(3), 34.
- [13] Wu, H. (2022). A High-Performance Customer Churn Prediction System based on Self-Attention (Version 1).
- [14] Wen, X., Wang, Y., Ji, X., & Traoré, M. K. (2022). Three-stage churn management framework based on DCN with asymmetric loss. *Expert Systems with Applications*,
- [15] Arshad, U., Khan, G., Alarfaj, F. K., Halim, Z., & Anwar, S. (2024). Q-ensemble learning for customer churn prediction with blockchain-enabled data transparency. *Annals of Operations Research*, 353(2), 607–633.
- [16] Zheng, X., Zhang, Y., Zheng, Y., Luo, F., & Lu, X. (2021). Abnormal event detection by a weakly supervised temporal attention network. *CAAI Transactions on Intelligence Technology*, 7(3), 419–431.
- [17] Haddadi, S. J., Farshidvard, A., Silva, F. dos S., dos Reis, J. C., & da Silva Reis, M. (2024). Customer churn prediction in imbalanced datasets with resampling methods: A comparative study. *Expert Systems with Applications*, 246, 123086
- [18] Amin, A., Adnan, A., & Anwar, S. (2023). An adaptive learning approach for customer churn prediction in the telecommunication industry using evolutionary computation and Naïve Bayes. *Applied Soft Computing*, 137, 110103
- [19] Zhang, X., Guo, F., Chen, T., Pan, L., Beliakov, G., & Wu, J. (2023). A Brief Survey of Machine Learning and Deep Learning Techniques for E-Commerce Research. *Journal of Theoretical and Applied Electronic commerce Research*, 18(4), 2188–2216.
- [20] Credit Card Service Churn Prediction by Machine Learning Models. (2024). *JST: Smart Systems and Devices*, 34(1), 16–22.

- [21] Sai, C. A. L. V. S., Kumari, K. S., & Gupta, B. A. (2024). Churn Prediction Based on Customer Segmentation in Banking Industry using Machine Learning Techniques. 2024 International Conference on Automation and Computation (AUTOCOM), 388–393.
- [22] C, J. K., Anumala, A., E.P, S., & E.S, B. (2024). Bank Customer Churn and Extra Benefits Prediction Using Machine Learning Model. 2024 Second International Conference on Advances in Information Technology
- [23] Renjith, Shini. (2015). An Integrated Framework to Recommend Personalized Retention Actions to Control B2C E-Commerce Customer Churn. International Journal of Engineering Trends and Technology. 27. 152-157. 10.14445/22315381/IJETT-V27P227.
- [24] Buckinx, Wouter & Van den Poel, Dirk. (2005). Customer base analysis: Partial defection of behaviourally loyal clients in a non-contractual FMCG retail setting. European Journal of Operational Research. 164. 252-268. 10.1016/j.ejor.2003.12.010.