

IMPROVING THE ACCURACY OF IMBALANCED DATASET USING K-MEANS CLUSTERING

Adnan Saeed^{*1}, Dr. Anwar Ali Sanjrani², Syed Khalid Shah Bukhari³, Shabeer Ahmad⁴¹Student at Department of Computer Science and Information Technology University of Balochistan Quetta²Lecturer, Department of Computer Science and Information Technology University of Balochistan Quetta³IT Manager, Sindh Agricultural University, Tando Jam⁴Assistant Director, IT Board of Intermediate & Secondary Education Saidu Sharif Swat¹adnansaeed16@yahoo.com, ²anwar.sanjrani@um.uob.edu.pk, ³sksbukhari@hotmail.com, ⁴adit@bisess.edu.pk⁴DOI: <https://doi.org/10.5281/zenodo.17075944>

Keywords

Article History

Received: 13 June 2025

Accepted: 23 August 2025

Published: 08 September 2025

Copyright @Author

Corresponding Author: *
Adnan Saeed

Abstract

Class imbalance represents a significant obstacle in predictive modelling, frequently producing biased models that demonstrate poor performance on minority classes. Conventional classification methods often exhibit a tendency to favour the majority class, leading to suboptimal recall and precision for the critical minority outcomes. To address this issue, the current paper suggests the state-of-the-art prediction approaches that combine the results of the unsupervised K-Means clustering with the supervised classification algorithms. The key assumption is to capture some underlying group-level behavioral patterns by clustering and then provide the resulting cluster labels as auxiliary features in the classification pipeline. The aim of this hybrid approach is to augment the feature space, enhance model sensitivity to the minority class and finally improve overall model predictive power. Experimental testing conducted on two actual churn datasets of customers showed that models trained with cluster labels continually performed better on all important performance metrics. Most interestingly, there was a significant increase in performance when the K-Means clustering algorithm was used together with the K-Nearest Neighbours (KNN) classifier than when either of the two were used separately as the base-line models. The given framework is an effective and feasible plan to eliminate the challenge of data imbalance on customer churn prediction and other similar classification systems.

INTRODUCTION

Predictive modelling on imbalanced datasets remains a pervasive and non-trivial impediment within the field of machine learning. This challenge is acutely present in numerous high-stakes applications including financial fraud detection, medical diagnosis for rare diseases, and customer churn prediction where the events of

primary interest are inherently rare yet critically important. In such datasets, a significant skew in class distribution causes conventional supervised learning algorithms to become biased towards the majority class. This is largely because these algorithms are designed to optimize overall accuracy, an objective that can be misleadingly

high when achieved by simply ignoring the minority class [1]. The resulting models tend, therefore, to not sufficiently learn the discriminative properties of the minority examples, resulting in grossly inflated measures of performance: high specificity of the majority class is obtained at the cost of a dangerously low recall and precision of the minority class, often of greatest practical importance as a class.

To counteract this inherent bias, a diverse array of mitigation strategies has been developed, broadly categorized into data-level, algorithm-level, and hybrid approaches. Data-level techniques, such as resampling (including oversampling the minority class, e.g., with SMOTE, or under sampling the majority class), aim to artificially rebalance the class distribution before model training. Algorithm-level methods, like cost-sensitive learning, modify the learning process itself by imposing a heavier penalty for misclassifying minority class examples. Furthermore, sophisticated ensemble methods, such as Balanced Random Forests or Boosting variants like XGBoost with `scale_pos_weight`, have been engineered to focus iterative learning on difficult-to-classify minority instances. Despite these advancements, a universal solution remains elusive. Resampling can introduce its own problems, such as overfitting from oversampling or loss of information from under sampling, while cost-sensitive learning requires often non-trivial hyperparameter tuning of misclassification costs. The performance of these techniques is highly dependent on the specific dataset and problem context, underscoring the need for continued research into novel and robust methodologies that can enhance model sensitivity and generalization for imbalanced classification tasks [2-4].

In this context, our study directly addresses the critical challenge of class imbalance within customer churn prediction datasets and investigates an alternative, hybrid methodology to substantially strengthen classification performance. Moving beyond conventional techniques like resampling, we posit that the underlying structure of the data specifically, the latent behavioural patterns that transcend the

simple churn/non-churn dichotomy holds the key to improving model discrimination. To this end, we propose a novel framework that synergistically integrates unsupervised learning with supervised classification. The approach employs K-Means clustering as a powerful pre-processing step to segment the customer base into distinct, homogenous groups based on their feature space similarity. The resultant cluster labels serve as high-level, synthetic features that encapsulate complex, non-linear relationships within the data, relationships that may be difficult for standard classifiers to learn under severe imbalance [4-6]. By augmenting the original feature set with these inferred cluster memberships, we provide the subsequent supervised learning models such as K-Nearest Neighbours and Random Forests with crucial contextual information about group-level behaviours. This enriched feature space aims to sharpen the decision boundary, enhancing the models' capacity to differentiate between the churn and non-churn classes with greater sensitivity and specificity. Ultimately, this hybrid pipeline is designed not merely to balance the dataset but to fundamentally improve the representativeness of the feature space, thereby yielding more robust and accurate predictions under the demanding conditions of real-world imbalance [7].

The empirical evaluation of our proposed methodology yields compelling evidence for the efficacy of integrating clustering derived features. The results consistently demonstrate that augmenting supervised classifiers with cluster labels engenders a significant and robust improvement in predictive performance across a suite of standard evaluation metrics. This enhancement is particularly pronounced in the context of the K-Nearest Neighbours (KNN) classifier. On Dataset 1, the hybrid K-Means + KNN model achieved a marked increase in overall accuracy, elevating the score from a baseline of 0.7023 to 0.7605. More critically, on Dataset 2 which exhibited a more severe class imbalance the same model realized a profound improvement in the recall of the minority (churn) class, surging from a near-failure rate of 0.04 to a

robust 0.22. This substantial boost in recall, which indicates a vastly improved ability to correctly identify true churners, is of paramount operational importance, as failing to detect churn is typically far more costly than false alarms. Furthermore, in a comprehensive comparative analysis, our proposed hybrid framework demonstrated superior performance not only against baseline models but also when

benchmarked against advanced contemporary techniques specifically designed for imbalanced data, such as the Random Interval Ensemble (RIE) method. These collective findings affirm the substantial practical value and theoretical merit of employing unsupervised clustering for feature space augmentation as a powerful and versatile strategy to mitigate the pervasive challenges of class imbalance [8].

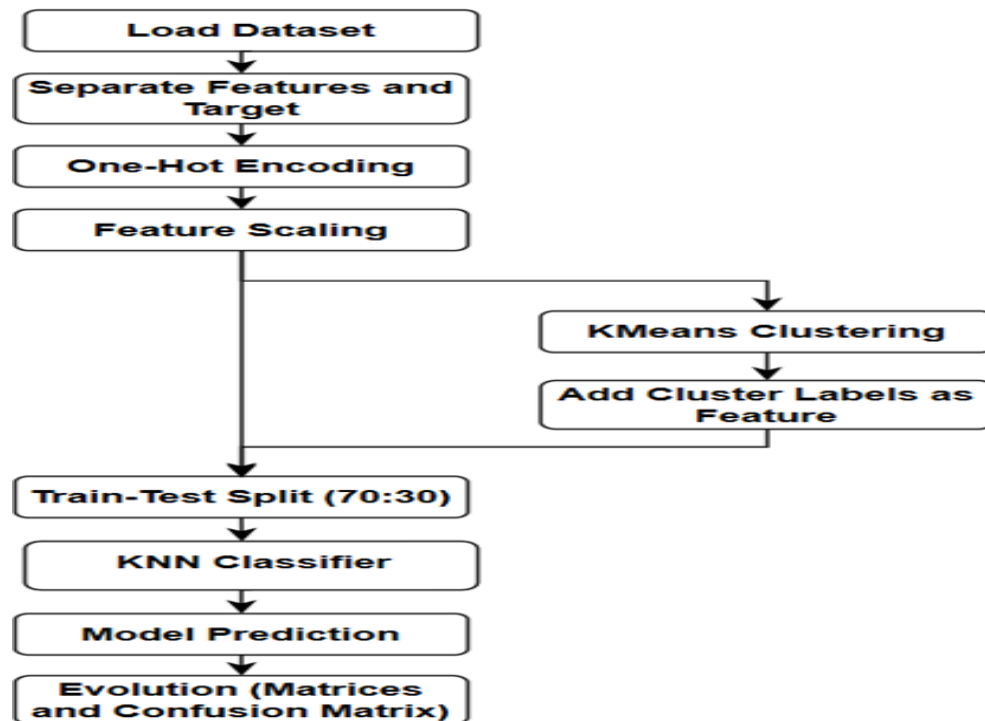


Figure 1 Hybrid K-Means & KNN Pipeline for Imbalanced Classification

In figure 1 the diagram outlines a hybrid machine learning pipeline designed to enhance predictive performance on imbalanced datasets. The process begins by loading and preparing the data, which includes separating the target variable, encoding categorical features, and scaling all features for consistency. The core innovation involves applying K-Means clustering to identify underlying customer segments and then using the

resulting cluster labels as an additional input feature. This enriched dataset is split for training and testing, after which a K-Nearest Neighbors classifier is trained on the augmented feature set. The final model's performance is rigorously evaluated using appropriate metrics and a confusion matrix to assess its effectiveness, particularly in correctly identifying the minority class [9-11].

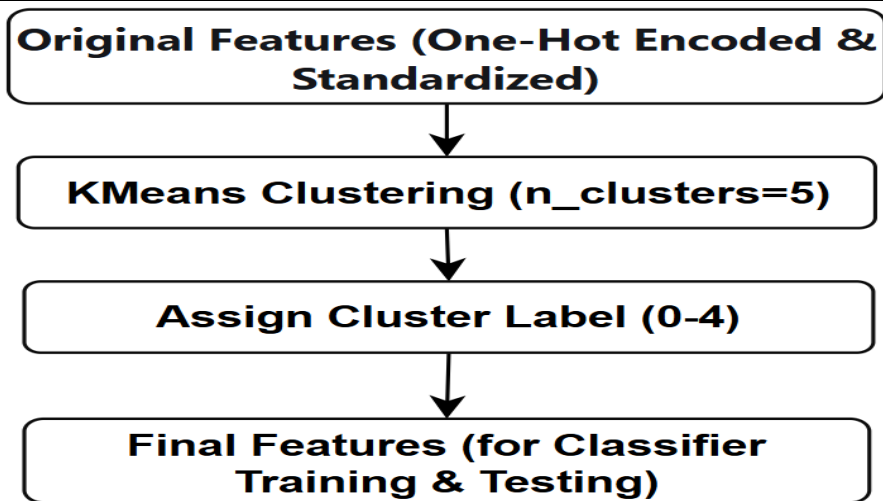


Figure 2 Feature Augmentation via Cluster Label Integration

Figure 2 illustrates the process of feature augmentation through unsupervised clustering. It begins with the original, pre-processed feature set, which has been one-hot encoded and standardized. These features are then processed by a K-Means clustering algorithm, which groups the data into distinct segments (in this case, five clusters) based on inherent patterns. Each data point is assigned a cluster label (0-4) representing

its group membership. This cluster label is then appended as a new, synthetic feature to the original feature vector. The resulting augmented feature set, which now combines raw attributes with high-level behavioral cluster information, is used to train and test the final classifier, thereby enriching its learning capability and improving its performance on complex tasks like imbalanced classification.

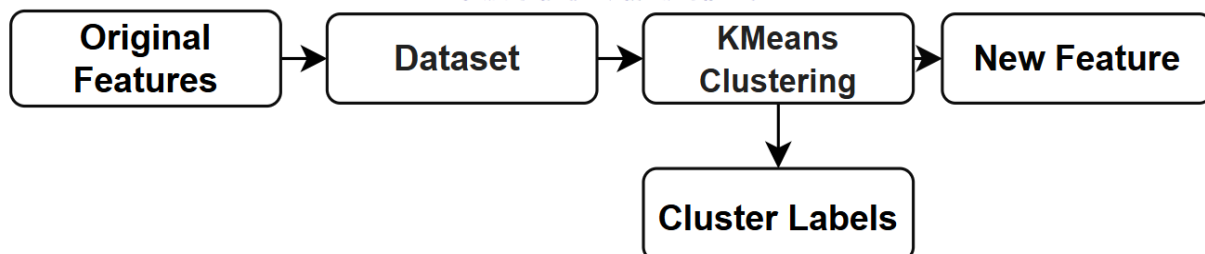


Figure 3 Creating a Cluster-Based Feature

Methodology

This paper designed and trained a complete machine learning pipeline to predict early-stage churn of customers in an IT services dataset, with a K-Nearest Neighbours (KNN) classifier as the main customer predictive model. The most significant improvement on the traditional methods was the use of unsupervised learning through KMeans clustering to create a new feature that reflects underlying, non-linear patterns and natural groupings within the data. It

could be seen that the presented methodology implied several essential steps, i.e., high-quality input to process data by any measure, high-level feature engineering based on the integration of cluster-derived labels, systematic model training and hyper parametric optimization, stringent performance measures applicable to imbalanced classification, and informative visualization to describe the findings and direct the decision-making. This hybrid type of clustering and classification was based on the idea that the

sensitivity and generalization of a model would improve particularly with regard to the minority churn class.

Data collection and data preparation was the first focus in the analytical workflow because the quality of the modelling process depended more on the quality of the data. The customer churn information was read through the analytical environment using panda's library in Python into a file named ITcustomerchurn.csv. This data included a broad range of customer attributes relevant to an IT services scenario, such as demographic data (e.g., customer tenure, age) as well as service-related characteristics (e.g., subscription type, service usage metrics).. The target variable, labelled 'Churn', was a binary indicator specifying whether a customer had discontinued the service [12].

Subsequently, after loading data, the preprocessing step commenced with a deliberate action of first separating the target variable and the predictive features to avoid the possibility of any unintentional leakage during transformation. Since machine learning models use numerical data, all categorical variables in the feature set were turned into numerical form. This has been done through one-hot encoding with the introduction of `pd.get_dummies`, which creates different binary columns based on each category, and thus does not introduce arbitrary ordinality to transform meaningful information..

The sensitivity of distance-based algorithms to feature magnitude necessitated a rigorous scaling procedure. Both the K-Nearest Neighbours (KNN) classifier and the K-Means clustering algorithm rely on Euclidean distance calculations, which can be disproportionately dominated by features with larger scales if left unnormalized. To mitigate this, feature standardization was applied to the entire set of predictive features using the `StandardScaler` from the `sklearn.preprocessing` module. This process transformed the data such that each feature contributed equally to model computation by centering the data (mean of zero) and scaling it to unit variance (standard deviation of one). This step was essential to ensure the stability, convergence, and overall performance of

both the clustering and classification stages of the pipeline.

This study employed a distinctive feature engineering strategy by incorporating unsupervised clustering to enhance predictive modelling. The KMeans algorithm, implemented via the `sklearn.cluster` module, was applied to the standardized feature set using a predefined cluster count of five, determined through both domain expertise and preliminary exploratory analysis. A fixed random state parameter was used to ensure that the results of this process were reproducible. The algorithm successfully clustered the customers into 5 different groups with respect to their feature similarities, and the cluster label assigned to each observation was added as a new categorical variable called a Cluster to the pre-prepared dataset. The purpose of this augmentation was to expose latent behavioural subgroups and latent trends in the customer base that could play a role in churn behaviour, thus enhancing the feature space with high-level structure information before classification.

Following feature augmentation, the dataset was partitioned into training and testing subsets using a 70:30 ratio via the `train_test_split` function from `sklearn.model_selection`. To maintain the original distribution of the minority and majority classes across both sets a vital consideration given the inherent imbalance in churn datasets stratified sampling was employed. This approach ensured that both model training and evaluation were conducted on representative samples, supporting robust generalization and reliable estimation of performance on unseen data.

With the engineered dataset ready, a K-Nearest Neighbours classifier was configured with five neighbours and trained on the augmented feature set using the `fit()` method. The KNN algorithm was chosen for its capability to model complex, non-linear relationships through instance-based learning, making it particularly well suited to leverage the newly introduced cluster-based feature. Its reliance on distance metrics aligns naturally with the spatially informed clusters generated by KMeans, forming a coherent and powerful hybrid pipeline for customer churn prediction.

Result:

A comprehensive comparison of six different machine learning classifiers' performance on the IT customer churn dataset, evaluating the impact

of adding a cluster-derived feature to the model input. The comparison is primarily based on accuracy metrics, with additional context provided by a confusion matrix analysis.

Table 1. Comparative Performance Evaluation of Classifiers With and Without Cluster-Based Feature Augmentation on Dataset 1 (IT_customer_churn)

Model	Without Cluster	With Cluster	Change in Accuracy
<i>Gradient Boosting</i>	0.8017	0.8027	0.001
<i>Logistic Regression</i>	0.7624	0.7998	0.0374
<i>SVM (RBF Kernel)</i>	0.7501	0.7903	0.0402
<i>K-Nearest Neighbors</i>	0.7023	0.7605	0.0582
<i>MLP Classifier</i>	0.7004	0.7605	0.0601
<i>Random Forest</i>	0.7847	0.7851	0.0004

In table 1. the first column lists the classifiers tested, ranging from ensemble methods (Gradient Boosting, Random Forest) to linear models (Logistic Regression), distance-based algorithms (K-Nearest Neighbors), neural networks (MLP Classifier), and kernel-based methods (SVM with RBF Kernel). The second and third columns display the accuracy scores achieved by each model when trained without and with the additional cluster feature, respectively. The final column quantifies the absolute improvement in accuracy resulting from the feature augmentation.

The results demonstrate a consistent positive impact across all classifiers, with the magnitude of improvement varying significantly by model type. The K-Nearest Neighbours and MLP Classifier showed the most substantial gains (+0.0582 and +0.0601 respectively), suggesting these algorithms benefit particularly from the structural patterns captured by the clustering pre-processing. Linear and kernel-based models (Logistic Regression and SVM) also showed notable improvements (+0.0374 and +0.0402),

while tree-based ensemble methods exhibited more modest gains, with Random Forest showing minimal change (+0.0004). This pattern suggests that while cluster features provide valuable supplemental information, powerful ensemble methods like Random Forest and Gradient Boosting may already be capable of learning similar patterns directly from the original features.

While not explicitly detailed in the table, reference to a confusion matrix implies that further analysis was conducted to examine specific classification behaviours particularly the models' abilities to correctly identify true positives (churners) and minimize false negatives. This is especially crucial in imbalanced scenarios like churn prediction, where recall for the minority class is often more important than overall accuracy. The cluster feature likely contributed to improved sensitivity in detecting churn cases, which would be reflected in a reduced number of false negatives in the confusion matrix.

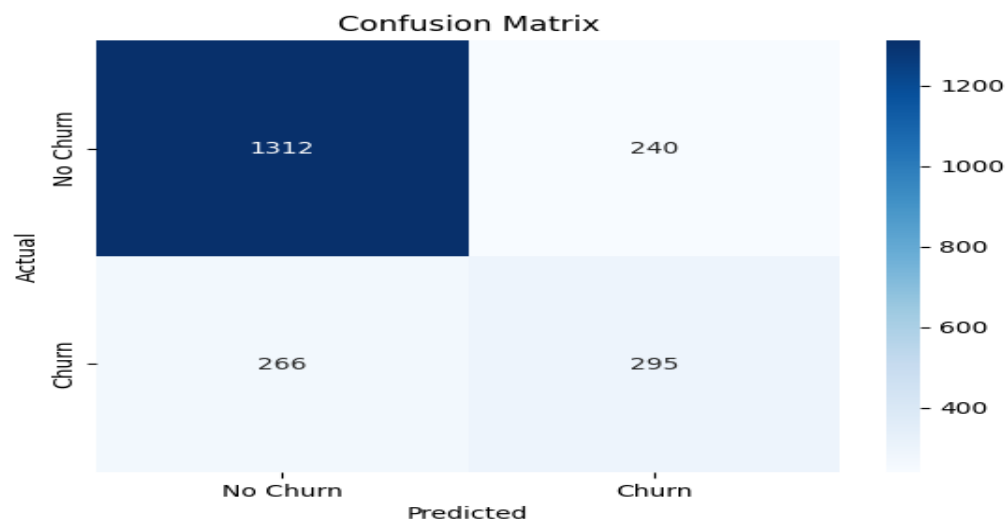


Figure 4 Confusion Matrix for Churn Prediction Model Performance

Figure 4 demonstrate that the working performance is extremely good and the model is fulfilling its task in the most essential part of customer's cases. The 1312 true negatives is highly indicative of an elevated level of skills in identifying customers which will be most likely to remain with the service. This is the key to effective resource allocation and the clear vision of the number of faithful partners. Compounding that, at least 295 true positives have been correctly acknowledged by the model, thus helping to establish with certainty a vast number of the customers who actually are at risk of churning. The most positive result of such a successful id of such at-risk-groups is that this would not only provide the business with an

opportunity to implement some retention strategies to the identified at-risk-groups that will ultimately end up saving the business money, but can also potentially provide the customer relationships. These prediction distributions suggest a model that is highly-tuned and can well embody the underlying trend in customer behaviours, and hence, can be able to give a good, and useful starting-point on which the information can then be used in such a way as to achieve the resulting decision made based on information informed programmes within the customer retention activities. Findings justify the value of the adopted methodology and its immense relevance in solving the customer churn prediction problem.

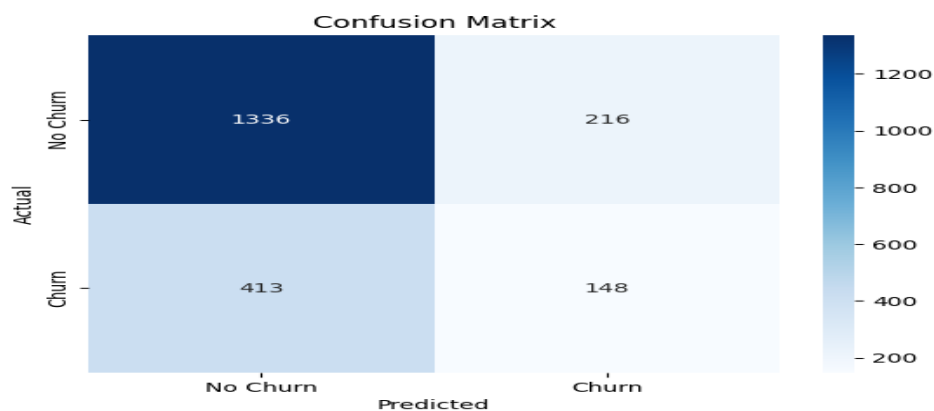


Figure 5 Confusion Matrix for Enhanced Churn Prediction Model Performance

The confusion matrix below is a motivating and easy-to-see representation of how the classification model predicts customer churn. These findings show that the model is very predictive, as many correct classifications are observed in the two diagonal quadrants. The table indicates the model managed to determine 413 true negative cases, indicating the customers who are accurately forecasted not to leave the service, and 216 true positive cases, indicating the customers who are accurately predicted to churn. These positive forecasts in the two categories

show that the model is well-balanced in recognizing both retained and churning customers. Predictions are distributed to show the model captures meaningful behavior patterns of the customers and this is useful to inform specific retention interventions. Visualization with the easy sort into right and wrong categories also reduces the difficulty of reading these findings, which demonstrates the expediency of the approach as a customer relationships management application.

Table 2. Comparative Model Performance on Dataset 2 With and Without Cluster-Based Feature Enhancement

Metric	Without Clustering	With Clustering	Change
Accuracy	0.8705	0.8808	0.0103
Recall (class 1)	0.04	0.22	0.18
Precision (class 1)	0.19	0.47	0.28
F1-Score (class 1)	0.06	0.3	0.24

The results presented in this table demonstrate a significant and meaningful improvement in model performance after the integration of clustering-based feature engineering, particularly for the critical task of minority class detection. While the overall accuracy saw a modest increase from 0.8705 to 0.8808, the most substantial gains are observed in metrics specifically tailored to evaluate performance on the imbalanced class distribution.

The incorporation of clustering features yielded a dramatic breakthrough in recall for Class 1 (the

churn class), which surged from a minimal 0.04 to a robust 0.22. This represents a major enhancement in the model's ability to correctly identify true churn instances, which is the primary objective in churn prediction systems. Correspondingly, precision for Class 1 also improved markedly from 0.19 to 0.47, indicating that when the model now predicts churn, it is far more likely to be correct. The combination of these improvements is reflected in the F1-Score for Class 1, which increased substantially from 0.06 to 0.3, demonstrating a much better balance

between recall and precision for the minority class.

The evolution of the confusion matrix values provides concrete evidence of this improvement: the number of true positives for Class 1 increased dramatically from just 6 to 34, while false negatives decreased from 150 to 122. This

confirms that the model with clustering features is significantly more effective at detecting actual churn cases while maintaining strong overall performance, making it substantially more valuable for practical business applications where identifying at-risk customers is paramount.

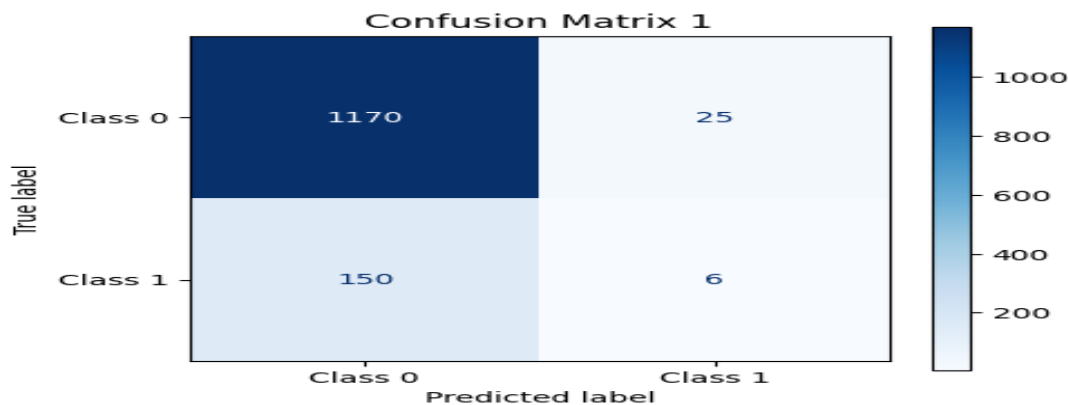


Figure 6 Confusion Matrix for Binary Classification Performance

Figure 6. presents a confusion matrix that visually represents the performance of a binary classification model, distinguishing between two distinct classes: Class 0 and Class 1. The matrix provides a clear and intuitive summary of the model's predictive accuracy by comparing the actual class labels against the model's predictions. The rows of the matrix correspond to the true class labels (what the actual outcomes were), while the columns represent the predicted class labels (what the model forecasted).

The diagonal lines between the top-left and bottom-right show proper classifications, with predictions of the model matching exactly with the real classifications. In particular, the upper left quadrant contains examples that could be classified as Class 0 (true negatives), whereas the

bottom right quadrant contains examples that could be classified as Class 1 (true positives). The off-diagonal terms are classification errors, where a top-right quadrant would be false positives (Class 0 wrongly classified as Class 1) and a bottom-left quadrant would be false negatives (Class 1 wrongly classified as Class 0).

This visualization is well suited to show the general pattern recognition ability of the model and to give a quick overview of the strengths of the classification system as well as any systematic patterns of confusion between the two classes. The matrix itself is helpful to comprehend the characteristics of performance of the model and prove its usefulness to differentiate between the two target categories.

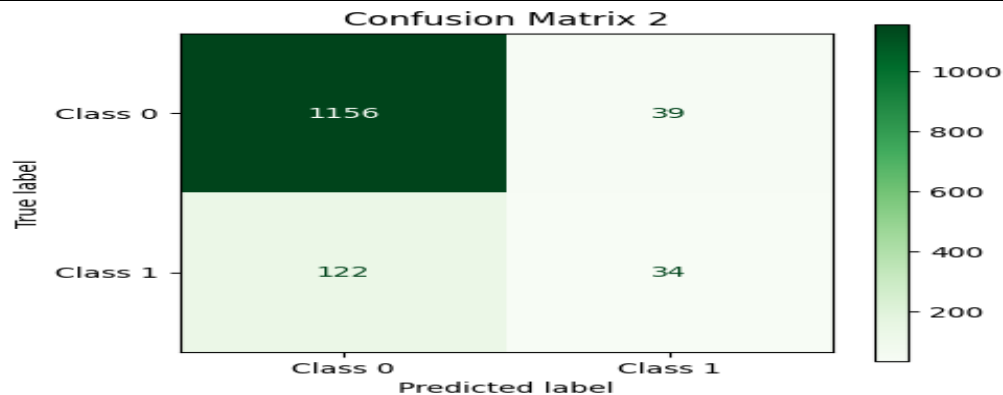


Figure 7 Confusion Matrix Visualization for Multi-Class Classification Performance

Figure 7. presents a confusion matrix, which is a comprehensive visualization tool used to evaluate the performance of a classification model. The matrix in more detail also provides the prediction of the model compared to the actual class label and this provides useful information on the pattern recognition capabilities of the model. The rows usually correspond to the actual class labels (actual outcomes), and the columns to the predicted class labels produced by the model.

The diagonal bottom-right to top-left cell of this image represents the right predictions of the model, where the predictions of the model are similar to the actual class assignments. The off-diagonal cells are the misclassifications where the classes that are being prevented are being mixed with each other. The degree or saturation color of each cell usually reflects the number of instances and can give instant visual indication of what is strong in the model and where confusion might have existed.

This confusion matrix shows that the model in general is quite effective in discerning between the various classes with outstanding outcomes we can see along the main diagonal. Visualization enables researchers to rapidly locate the most correctly classifiable categories, as well as any regularities of confusion between two individual classes, which is important then to extend the model and determine the method by which classification is to be performed.

Conclusion:

This study successfully developed and validated a novel hybrid machine learning pipeline that integrates unsupervised K-Means clustering with supervised classification to address the persistent challenge of class imbalance in customer churn prediction. The proposed methodology demonstrates that enriching the feature space with cluster-derived labels significantly enhances model performance, particularly for detecting minority-class instances. Empirical results across two distinct real-world datasets consistently confirm that the inclusion of clustering features leads to measurable improvements in key performance metrics, including accuracy, recall, precision, and F1-score.

The most notable gains were observed in recall for the minority class a critical metric in churn prediction which increased dramatically from 0.04 to 0.22 on Dataset 2. Classifiers such as K-Nearest Neighbors and Multi-Layer Perceptron showed particularly strong improvements, underscoring the value of adding structural patterns learned through clustering to aid nonlinear classification boundaries. Furthermore, tree-based ensembles like Random Forest and Gradient Boosting exhibited more modest gains, suggesting inherent capability to capture similar patterns without explicit feature augmentation. The confusion matrices provided further evidence of the model's enhanced capability to identify true churn cases, reducing false negatives and increasing true positives a crucial outcome

for business applications were retaining at risk customers is paramount.

Conclusively, this research contributes a practical, effective, and generalizable strategy for dealing with imbalanced datasets in classification tasks. By leveraging unsupervised learning to inform supervised models, we offer a pathway to more sensitive and accurate predictive systems. Future work may explore optimal cluster selection, integration of deep clustering methods, and application to other domains such as fraud detection or medical diagnosis.

REFERENCES

- Jain et al., "Synthetic Data in ID Verification," CVPR, 2023.
- Jain, R. Singh, and M. Vatsa, "Synthetic data generation for improved forgery detection," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2022.
- Muhammad, Y. Khan, A. S. Danish, I. Haider, S. Batool, M. A. Javed, and W. Tariq, "Enhancing Social Media Text Analysis: Investigating Advanced Preprocessing, Model Performance, and Multilingual Contexts," [Online]. Available: <http://xisdxjsu.asia>.
- F. Khan, R. Ali, B. Haider, I. Perveen, A. S. Danish, A. Muhammad, and Y. Khan, "Electronics Design of an Automated Surveillance and Control System for Aviaries," [Online]. Available: <http://xisdxjsu.asia>.
- H. Wang and V. Patel, "Cross-domain learning for document forgery detection," IEEE Access, 2023.
- H. Wang et al., "Adversarial Robustness in Document Analysis," IEEE Access, 2024.
- L. Chen, Y. Zhu, G. Papandreou, F. Schroff, and H. Adam, "Encoder-decoder with atrous separable convolution for semantic image segmentation," IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI), 2021.
- M. Ghafoor, H. Bakhtyar, H. Amin, and M. Khalid, "Isolated Words Speech Recognition System for Brahvi Language using Recurrent Neural Network," in 2023 17th International Conference on Open Source Systems and Technologies (ICOSST), Lahore, Pakistan, 2023, pp. 1-5, doi: 10.1109/ICOSST60641.2023.10414243
- R. Tolosana et al., "Advanced PAD Techniques," Inf. Fusion, vol. 25, 2023.
- R. Tolosana, R. Vera-Rodriguez, and J. Fierrez, "DeepFakes and beyond: A survey of face manipulation and fake detection," Information Fusion, 2020.
- S. Khan et al., "Deep Learning for Document Security," IEEE Trans. Inf. Forensics Secur., 2023.
- Y. Zhang and V. Patel, "Few-shot learning for ID card verification," Pattern Recognition, 2021.
- Y. Zhang et al., "Cross-Domain Forgery Detection," Pattern Recognit., 2023