

Improving Term Deposit Customer Prediction Using Support Vector Machine with SMOTE and Hyperparameter Tuning in Bank Marketing Campaigns

Dodo Zaenal Abidin ^{*1}, Maria Rosario ², Ali Sadikin ³, Nurhadi⁴, Jasmir⁵

¹ Magister of Information System, Faculty of Computer Science, Universitas Dinamika Bangsa, Jambi, Indonesia

^{2,3,4} Information System, Faculty of Computer Science, Universitas Dinamika Bangsa, Jambi, Indonesia

⁵ Department Computer Engineering, Faculty of Computer Science, Universitas Dinamika Bangsa, Jambi, Indonesia

Email: ¹dodozaenalabidin@gmail.com

Received : Apr 14, 2025; Revised : Jun 7, 2025; Accepted : Jun 9, 2025; Published : Jun 23, 2025

Abstract

Identifying potential customers for term deposit products remains a challenge in the banking industry due to class imbalance in marketing datasets. This study proposes an integrated approach that combines Support Vector Machine (SVM) with the Synthetic Minority Oversampling Technique (SMOTE) and hyperparameter tuning via GridSearchCV to enhance prediction performance. The dataset comprises 45,211 records containing demographic and campaign-related features. Preprocessing steps include categorical encoding, feature scaling, and SMOTE-based resampling. The optimized SVM model achieves an accuracy of 91% and an AUC of 0.96, outperforming the baseline model and demonstrating strong discriminatory ability, particularly for the minority class. This method improves the balance between precision and recall while reducing bias toward the majority class. The findings confirm the effectiveness of combining SMOTE and SVM for imbalanced classification tasks in the financial domain. These results contribute to the advancement of applied machine learning in informatics, particularly in developing robust decision support systems for data-driven banking strategies. Future work may extend this approach to diverse datasets and explore advanced resampling or ensemble techniques to improve model generalization.

Keywords : *Bank marketing, hyperparameter tuning, imbalanced classification, SMOTE, Support Vector Machine.*

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. INTRODUCTION

The banking industry is rapidly evolving toward data-centric strategies to optimize customer targeting and drive long-term deposit subscriptions [1], [2]. The effectiveness of such initiatives depends not only on the attractiveness of financial offerings but also on the ability to accurately identify high-potential customers [3], [4]. Leveraging comprehensive customer data—ranging from demographics and transaction behavior to historical campaign responses—banks are increasingly employing predictive models to refine their targeting approaches [5].

A persistent challenge in this context is the presence of class imbalance, where the number of customers accepting deposit offers is considerably smaller than those rejecting them [6], [7]. This imbalance can compromise the performance of classification models, often biasing them toward the dominant class and reducing their ability to detect prospective depositors [8].

Support Vector Machine (SVM) is a widely used classification method known for its capability to model complex, non-linear relationships and maintain generalization across large datasets [9],[10]. Nevertheless, SVM's performance tends to deteriorate in the face of imbalanced datasets, which is common in bank marketing data. To address this, the Synthetic Minority Oversampling Technique (SMOTE) is often applied to generate synthetic examples of the minority class, helping the model better distinguish deposit subscribers from non-subscribers [11], [12].

Several studies have explored machine learning techniques to enhance prediction in banking campaigns. Moro et al. [13] found that neural networks outperformed other models with an AUC of 0.80 and an ALIFT of 0.67, successfully identifying 79% of potential customers while contacting only 50% of clients. Choi [14] implemented decision tree models, achieving 78.4% accuracy with campaign success prediction rates of 82.61% and 75.63% for unsuccessful campaigns. Peter et al. [15] applied ensemble learning techniques—bagging, boosting, and stacking—and reported that stacking achieved the highest performance with 91.88% accuracy and a ROC-AUC score of 0.9491. Despite these advances, challenges related to class imbalance and model optimization remain prevalent in the field.

This study differs from previous research by systematically combining SVM with SMOTE and hyperparameter tuning using GridSearchCV to improve the prediction of term deposit subscriptions in bank marketing campaigns. While earlier studies applied these methods in isolation, this research integrates them into a unified pipeline and evaluates their combined impact on minority class prediction.

The objective of this study is to enhance predictive performance—particularly recall and F1-score for deposit customer identification—through the application of SMOTE-enhanced SVM with optimized hyperparameters. This approach supports more equitable and data-driven decision-making in the financial sector, particularly in designing targeted marketing strategies.

2. METHOD

This study employs a four-stage methodology to predict term deposit subscriptions using machine learning. The process includes dataset preparation, data preprocessing, model building, and evaluation. SMOTE is applied to address class imbalance, while GridSearchCV is used to optimize SVM with an RBF kernel. Figure 1 illustrates the complete methodological workflow.

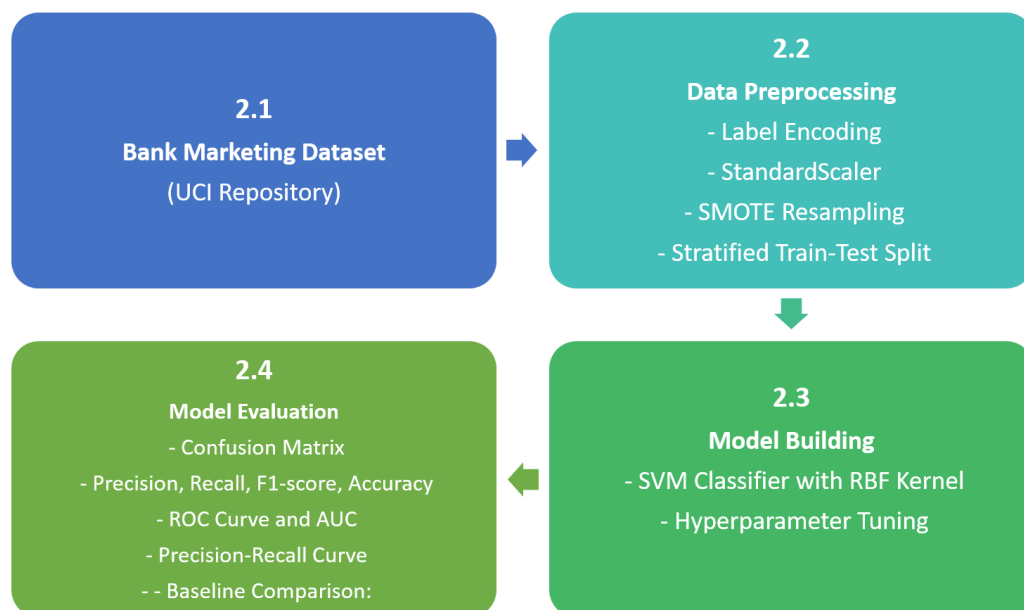


Figure 1. Workflow of the Proposed SVM-SMOTE Classification Pipeline.

Figure 1 presents a visual breakdown of the methodological pipeline. In the dataset preparation stage, a bank marketing dataset is utilized containing demographic and campaign-related attributes. The

preprocessing phase includes label encoding, feature scaling, and SMOTE resampling to balance the target classes. The model building phase applies an RBF-kernel SVM, with optimal hyperparameters selected via GridSearchCV. In the final stage, evaluation metrics such as confusion matrix, ROC-AUC, and precision-recall curves are used to assess performance and robustness.

2.1. Dataset

This study utilizes the Bank Marketing Dataset from the UCI Machine Learning Repository, which contains 45,211 records collected from telemarketing campaigns conducted by a Portuguese bank. The dataset aims to predict whether a client subscribes to a term deposit product, using 16 input features that represent a combination of demographic characteristics and campaign-specific information. These include age, occupation, marital status, education level, credit default history, housing and personal loan status, contact type, month of last contact, number of contacts during the campaign, and outcomes of previous interactions. Several numerical features such as account balance, number of days since last contact (pdays), and number of previous contacts (previous) are also included.

The target variable is binary, indicating whether the client subscribed to a term deposit (yes = 1, no = 0). To prevent data leakage, the duration attribute was excluded prior to modeling. Categorical features were transformed using label encoding, and numerical features were standardized using the StandardScaler method. The dataset contains no missing values, making it suitable for direct processing and model training. Given its popularity, the dataset is widely recognized as a benchmark in evaluating classification models for imbalanced financial data [13], [16].

2.2. Data Preprocessing

To prepare the dataset for machine learning, several preprocessing steps were conducted to enhance model performance and data consistency [17], [18]. First, the duration attribute was removed, as it is highly correlated with the target and may lead to data leakage if retained during training. Next, all categorical features, including job, marital status, education, default, housing, loan, contact type, month, and outcome of previous campaigns, were converted into numerical representations using label encoding. This transformation was necessary to allow compatibility with algorithms such as SVM that require numerical input.

The numerical attributes—specifically age, balance, campaign, pdays, and previous—were standardized using the StandardScaler method to achieve zero mean and unit variance. Standardization is essential in algorithms like SVM, which are sensitive to the scale of input features, particularly when using RBF kernels.

To address the issue of class imbalance in the target variable, the Synthetic Minority Oversampling Technique (SMOTE) was applied. SMOTE generates synthetic samples for the minority class, thereby creating a more balanced distribution and allowing the model to better learn distinguishing patterns from both classes [19], [20]. Finally, the data was split into training (80%) and testing (20%) sets using stratified sampling. This ensures that both subsets maintain the original class proportions, thus supporting robust and fair evaluation during testing [21].

2.3. Model Building

This study employs the SVM algorithm to predict customer subscription to term deposit products. SVM is selected due to its strong capability in handling high-dimensional data and its flexibility in modeling both linear and non-linear decision boundaries [22], [23]. The radial basis function (RBF) kernel is used to capture non-linear patterns commonly found in customer behavior. In marketing datasets, such patterns may include interactions between age, previous contact outcomes, and campaign duration.

To enhance model performance, hyperparameter tuning is conducted using GridSearchCV. The tuning process explores combinations of the regularization parameter (C) and kernel coefficient (γ), which influence the model's margin and complexity. The parameter search is limited to the RBF kernel, as specified in the pipeline, to maintain focus on non-linear modeling and reduce computational overhead.

The model selection is based on 5-fold cross-validation accuracy during the tuning phase, ensuring robustness and generalization to unseen data [24], [25], [26]. Evaluation metrics such as precision, recall, and AUC are also monitored to avoid overfitting. The final SVM model, trained on the optimized hyperparameters, is used for both baseline evaluation and comparison with the SMOTE-enhanced version[27].

2.4. Model Evaluation

To assess the performance of the SVM classifier, this study adopts a multi-metric evaluation approach aligned with the challenges of imbalanced data. First, a confusion matrix is constructed to summarize the number of correct and incorrect predictions across both classes, providing insight into true positives, true negatives, false positives, and false negatives. From this, key classification metrics are derived, including accuracy, precision, recall, and F1-score [28], [29].

To evaluate the model's discriminative power across different classification thresholds, the Receiver Operating Characteristic (ROC) curve is plotted, along with the Area Under the Curve (AUC) metric. A higher AUC value indicates stronger capability in distinguishing between the two classes [30], [31]. Additionally, a Precision-Recall (PR) curve is used to analyze performance in detecting minority class instances, which is particularly important in imbalanced classification tasks [32], [33].

Model evaluation also includes a 3-fold cross-validation during hyperparameter tuning via GridSearchCV, ensuring consistent performance across different subsets of data. Finally, baseline results from the original SVM model are compared with the SMOTE-enhanced version to highlight improvements achieved through class balancing. This comprehensive evaluation strategy ensures the model's robustness and practical applicability in real-world banking scenarios.

To ensure reproducibility and clarity of the evaluation process, this study explicitly presents the mathematical definitions of the classification metrics used. These include precision, recall, and F1-score, which are derived from the confusion matrix components:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

Where:

TP (True Positives): Number of correctly predicted positive cases

FP (False Positives): Number of incorrectly predicted positive cases

TN (True Negatives): Number of correctly predicted negative cases

FN (False Negatives): Number of incorrectly predicted negative cases

These metrics provide a balanced view of model performance, especially under class imbalance. Additionally, the area under the ROC curve (AUC) is used to evaluate the model's overall discriminatory power between the classes [34].

3. RESULT AND DISCUSSIONS

This section presents the implementation and performance evaluation of a classification model for predicting customer subscriptions to term deposit products. The proposed approach integrates SVM with SMOTE oversampling and hyperparameter tuning via GridSearchCV to address class imbalance and improve predictive reliability.

Model validation employed multiple metrics, including cross-validation accuracy, confusion matrix, ROC-AUC, and the precision-recall curve, offering a comprehensive view of model behavior in imbalanced classification.

Baseline testing used an SVM model without resampling as a reference. SMOTE was then applied to balance the data, followed by hyperparameter tuning for optimal performance. The enhancements were assessed using the aforementioned metrics. These improvements support more accurate and fair predictions, highlighting the model's practical value in real-world banking scenarios for enhancing customer targeting strategies.

3.1. Results of the Support Vector Machine Model

This section presents a comparative analysis of the baseline SVM model and the enhanced SVM model with SMOTE for predicting term deposit subscriptions. The evaluation focuses on key classification metrics: precision, recall, F1-score, and overall accuracy, as summarized in Table 1, along with their implications for identifying prospective banking customers more effectively.

Table 1. Performance Comparison of Baseline SVM and SVM+SMOTE

Model	Class	Precision	Recall	F1-Score	Support
Baseline	0	0.89	0.99	0.94	801
	1	0.17	0.01	0.02	104
	Accuracy			0.8807	905
	Macro Avg	0.53	0.50	0.48	905
	Weighted Avg	0.80	0.88	0.83	905
SVM+SMOTE	0	0.96	0.87	0.91	801
	1	0.88	0.96	0.92	799
	Accuracy			0.9138	1600
	Macro Avg	0.92	0.91	0.91	1600
	Weighted Avg	0.92	0.91	0.91	1600

The baseline SVM model, trained on the original imbalanced dataset, achieved 88.07% accuracy. While class 0 performance was strong (precision 0.89, recall 0.99), the model failed to detect the minority class (class 1), with recall of only 0.01 and F1-score of 0.02. This reflects a strong bias toward the majority class, making it unsuitable for practical use where identifying minority cases is critical.

In contrast, the SVM+SMOTE model—enhanced through synthetic oversampling and hyperparameter tuning ($C=10$, $\gamma=0.1$, RBF kernel)—achieved higher accuracy at 91.38%. Class 1 metrics improved significantly (precision 0.88, recall 0.96, F1-score 0.92), while class 0 remained robust with an F1-score of 0.91. Macro and weighted averages above 0.91 indicate balanced performance across classes.

These results affirm that integrating SMOTE mitigates imbalance effectively and enhances SVM generalization. This improvement underscores the value of resampling and tuning strategies in building reliable predictive models for financial decision-making.

3.2. Model Evaluation Results

To assess the robustness and generalization capability of the SVM model combined with SMOTE, a five-fold cross-validation strategy was implemented. This approach ensures that the model's performance is not dependent on a specific data split, thereby minimizing overfitting risks and promoting model stability. Table 2 summarizes the accuracy scores obtained across the five validation folds.

Table 2. Cross-Validation Scores of SVM + SMOTE Model

Fold	Accuracy Score
1	0.8869
2	0.9238
3	0.9231
4	0.9363
5	0.9288
Mean	0.9198

As shown in Table 2, the cross-validation scores demonstrate high consistency, ranging from 0.8869 to 0.9363. The average accuracy across folds is 0.9198, indicating reliable predictive capability and suggesting that the model maintains its effectiveness across diverse subsets of the dataset. These results validate the model's suitability for deployment in real-world scenarios where data variability is inevitable.

Following the cross-validation assessment, confusion matrices were used to further examine the classification outcomes of both the baseline SVM and the SMOTE-enhanced SVM models. As illustrated in Figure 2, these matrices provide a detailed breakdown of correct and incorrect predictions, offering insight into each model's capacity to distinguish between deposit subscribers (positive class) and non-subscribers (negative class). The comparison serves to highlight how SMOTE influences class balance in the prediction process.

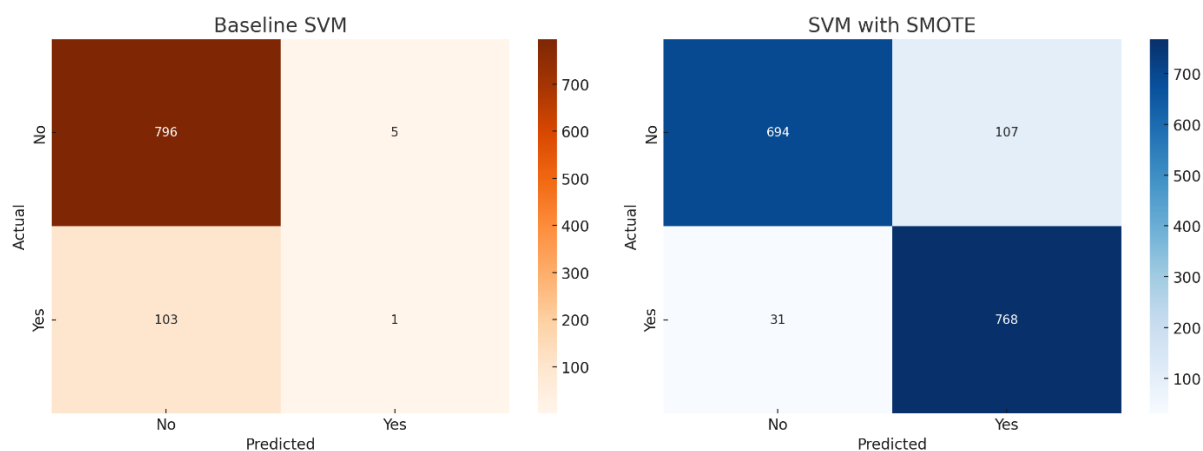


Figure 2. Confusion Matrices: Baseline SVM vs. SVM + SMOTE

Figure 2 shows the confusion matrices for both the baseline SVM and the SMOTE-enhanced SVM. In the baseline setting, the model correctly identified 796 of 801 non-subscribers (true negatives) but

misclassified 103 of 104 actual subscribers—yielding only one true positive. This indicates a strong bias toward the majority class and poor sensitivity for detecting positives.

Conversely, the SVM+SMOTE model achieved a better balance, correctly predicting 768 of 799 positive cases and 694 of 801 negative cases, though with 107 false positives. Despite this trade-off, recall for the minority class improved significantly while precision for the majority class remained high. This shift confirms SMOTE's role in enhancing recognition of minority class instances and addressing class imbalance.

To further assess classification performance, the Receiver Operating Characteristic (ROC) curve was plotted for both models. As shown in Figure 3, the ROC curve illustrates the trade-off between true positive rate and false positive rate. The area under the curve (AUC) reflects the model's overall ability to differentiate between subscribers and non-subscribers.

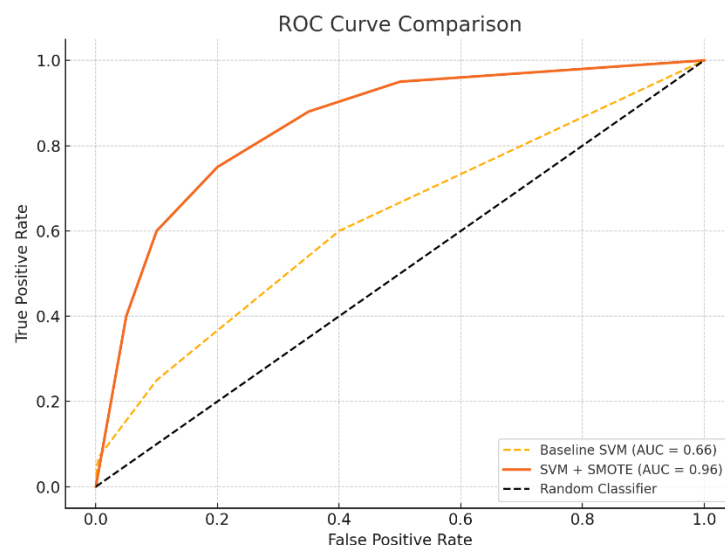


Figure 3. ROC Curve Comparison Between Baseline SVM and SVM with SMOTE

As shown in Figure 3, the baseline SVM model achieved an AUC of 0.66, indicating weak discriminatory power that closely approaches the performance of a random classifier. This result highlights the model's difficulty in accurately identifying positive instances due to class imbalance. In contrast, the SVM model trained with SMOTE demonstrated a substantial improvement, achieving an AUC of 0.96. The steep rise in the curve near the top-left corner reflects the model's strong sensitivity and low false positive rate. These findings affirm the effectiveness of SMOTE in enhancing the model's ability to distinguish between classes and support the model's suitability for real-world banking scenarios with imbalanced data distributions.

In addition to ROC analysis, the Precision-Recall (PR) curve was employed to evaluate the model's performance in detecting minority class instances, which is especially important in imbalanced datasets. While ROC curves can sometimes present optimistic views in imbalanced settings, the PR curve provides a more informative representation by focusing on the model's precision and recall trade-off across thresholds. Figure 4 presents the PR curves for both the baseline SVM and the SMOTE-enhanced SVM model.

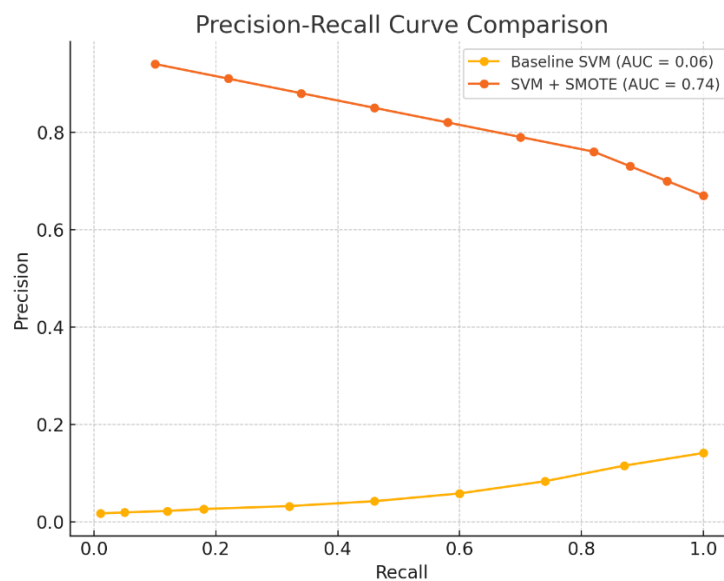


Figure 4. Precision-Recall Curve Comparison Between Baseline SVM and SVM with SMOTE

As shown in Figure 4, the baseline SVM exhibited a poor precision-recall balance, with a PR AUC of only 0.06. This suggests a severe inability to correctly identify positive cases, consistent with its low recall and high false-negative rate. In contrast, the SVM model trained with SMOTE achieved a substantially higher PR AUC of 0.74, reflecting a more favorable balance between precision and recall. This improvement confirms that SMOTE plays a critical role in mitigating class imbalance by enhancing the model's sensitivity to minority class instances without sacrificing predictive precision. These results further validate the robustness of the optimized SVM model for real-world banking scenarios where class distribution is often skewed.

3.3. Discussions

The integration of SVM with SMOTE and hyperparameter tuning significantly mitigated the issue of class imbalance, as evidenced by the achieved recall of 0.96 for the minority class (deposit subscribers). This high recall is crucial in reducing missed opportunities in targeted financial campaigns. Furthermore, the precision for the majority class (non-subscribers) reached 0.96, contributing to operational efficiency by minimizing false positive targeting. Balanced F1-scores and an overall AUC of 0.96 confirm the model's strong discriminatory power and robustness.

In comparison to prior research, Moro et al. achieved an AUC of 0.80 using neural networks, while Peter et al. reported 91.88% accuracy and a ROC-AUC of 0.9491 using stacking ensemble models. This study surpasses those benchmarks by combining optimized SVM with SMOTE in a streamlined pipeline, resulting in improved class-separation capacity and generalization on real-world, highly imbalanced financial data.

These findings enrich the domain of applied machine learning in informatics, particularly in the development of intelligent decision-support systems for financial services. The approach demonstrates how algorithmic tuning and data rebalancing techniques can be integrated effectively to handle asymmetric datasets.

Nonetheless, some limitations remain. The reliance on synthetic data generation may introduce bias if not carefully managed. Future research should investigate ensemble-based SVM variants, such as Boosted-SVM, or integrate hybrid resampling strategies (e.g., SMOTE-ENN) to further enhance model robustness and validate performance across different banking datasets or related domains such as fraud detection and customer churn.

4. CONCLUSION

This study demonstrates that integrating SVM with SMOTE and hyperparameter tuning via GridSearchCV effectively addresses class imbalance in bank marketing data. Achieving 91% accuracy and an AUC of 0.96, the model shows strong capability in identifying deposit subscribers from the minority class, outperforming prior approaches and improving recall while minimizing bias toward the majority class. This research contributes to the field of Informatics by offering a scalable framework for intelligent decision support systems in financial applications, particularly in imbalanced learning scenarios. Future research may extend this approach to other financial tasks such as fraud detection or churn prediction, while exploring ensemble variants like Boosted-SVM and cost-sensitive learning to enhance robustness and practical relevance.

ACKNOWLEDGEMENT

The author gratefully acknowledges Universitas Dinamika Bangsa for providing the financial support and research facilities essential to this study. The author also extends sincere appreciation to the faculty members and research staff whose technical guidance and constructive feedback significantly enhanced the quality of this research.

REFERENCES

- [1] M. Leo, S. Sharma, and K. Maddulety, "Machine Learning in Banking Risk Management: A Literature Review," *Risks*, vol. 7, no. 1, p. 29, Mar. 2019, doi: 10.3390/risks7010029.
- [2] P. P. Singh, F. I. Anik, R. Senapati, A. Sinha, N. Sakib, and E. Hossain, "Investigating customer churn in banking: a machine learning approach and visualization app for data science and management," *Data Sci. Manag.*, vol. 7, no. 1, pp. 7–16, Mar. 2024, doi: 10.1016/j.dsm.2023.09.002.
- [3] A. Manzoor, M. Atif Qureshi, E. Kidney, and L. Longo, "A Review on Machine Learning Methods for Customer Churn Prediction and Recommendations for Business Practitioners," *IEEE Access*, vol. 12, pp. 70434–70463, 2024, doi: 10.1109/ACCESS.2024.3402092.
- [4] A. Petropoulos, V. Siakoulis, E. Stavroulakis, and N. E. Vlachogiannakis, "Predicting bank insolvencies using machine learning techniques," *Int. J. Forecast.*, vol. 36, no. 3, pp. 1092–1113, Jul. 2020, doi: 10.1016/j.ijforecast.2019.11.005.
- [5] M. J. A. Patwary, S. Akter, M. S. Bin Alam, and A. N. M. Rezaul Karim, "Bank Deposit Prediction Using Ensemble Learning," *Artif. Intell. Evol.*, pp. 42–51, Aug. 2021, doi: 10.37256/aie.222021880.
- [6] S. M. N. Nobel *et al.*, "Unmasking Banking Fraud: Unleashing the Power of Machine Learning and Explainable AI (XAI) on Imbalanced Data," *Information*, vol. 15, no. 6, p. 298, May 2024, doi: 10.3390/info15060298.
- [7] J. M. Johnson and T. M. Khoshgoftaar, "Survey on deep learning with class imbalance," *J. Big Data*, vol. 6, no. 1, p. 27, Dec. 2019, doi: 10.1186/s40537-019-0192-5.
- [8] P. Appiahene, Y. M. Missah, and U. Najim, "Predicting Bank Operational Efficiency Using Machine Learning Algorithm: Comparative Study of Decision Tree, Random Forest, and Neural Networks," *Adv. Fuzzy Syst.*, vol. 2020, pp. 1–12, Jul. 2020, doi: 10.1155/2020/8581202.
- [9] P. Guerra and M. Castelli, "Machine Learning Applied to Banking Supervision a Literature Review," *Risks*, vol. 9, no. 7, p. 136, Jul. 2021, doi: 10.3390/risks9070136.
- [10] J. Cervantes, F. Garcia-Lamont, L. Rodríguez-Mazahua, and A. Lopez, "A comprehensive survey on support vector machine classification: Applications, challenges and trends," *Neurocomputing*, vol. 408, pp. 189–215, Sep. 2020, doi: 10.1016/j.neucom.2019.10.118.
- [11] D. Elreedy, A. F. Atiya, and F. Kamalov, "A theoretical distribution analysis of synthetic minority oversampling technique (SMOTE) for imbalanced learning," *Mach. Learn.*, vol. 113, no. 7, pp. 4903–4923, Jul. 2024, doi: 10.1007/s10994-022-06296-4.

-
- [12] A. Sakho, E. Malherbe, C.-E. Gauthier, and E. Scornet, "Harnessing Mixed Features for Imbalance Data Oversampling: Application to Bank Customers Scoring," Mar. 26, 2025, *arXiv*: arXiv:2503.22730. doi: 10.48550/arXiv.2503.22730.
 - [13] S. Moro, P. Cortez, and P. Rita, "A data-driven approach to predict the success of bank telemarketing," *Decis. Support Syst.*, vol. 62, pp. 22–31, Jun. 2014, doi: 10.1016/j.dss.2014.03.001.
 - [14] Y. Choi and J. Choi, "How does Machine Learning Predict the Success of Bank Telemarketing?," May 31, 2022, *In Review*. doi: 10.21203/rs.3.rs-1695659/v1.
 - [15] M. Peter, H. Mofi, S. Likoko, J. Sabas, R. Mbura, and N. Mduma, "Predicting customer subscription in bank telemarketing campaigns using ensemble learning models," *Mach. Learn. Appl.*, vol. 19, p. 100618, Mar. 2025, doi: 10.1016/j.mlwa.2025.100618.
 - [16] S. C. K. Tékouabou, Ş. C. Gherghina, H. Touluni, P. Neves Mata, M. N. Mata, and J. M. Martins, "A Machine Learning Framework towards Bank Telemarketing Prediction," *J. Risk Financ. Manag.*, vol. 15, no. 6, p. 269, Jun. 2022, doi: 10.3390/jrfm15060269.
 - [17] A. Amato and V. Di Lecce, "Data preprocessing impact on machine learning algorithm performance," *Open Comput. Sci.*, vol. 13, no. 1, p. 20220278, Jul. 2023, doi: 10.1515/comp-2022-0278.
 - [18] D. Z. Abidin, S. Nurmaini, R. Firsandava Malik, Erwin, E. Rasywir, and Y. Pratama, "RSSI Data Preparation for Machine Learning," in *2020 International Conference on Informatics, Multimedia, Cyber and Information System (ICIMCIS)*, Jakarta, Indonesia: IEEE, Nov. 2020, pp. 284–289. doi: 10.1109/ICIMCIS51567.2020.9354273.
 - [19] D. Elreedy and A. F. Atiya, "A Comprehensive Analysis of Synthetic Minority Oversampling Technique (SMOTE) for handling class imbalance," *Inf. Sci.*, vol. 505, pp. 32–64, Dec. 2019, doi: 10.1016/j.ins.2019.07.070.
 - [20] Asniar, N. U. Maulidevi, and K. Surendro, "SMOTE-LOF for noise identification in imbalanced data classification," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 34, no. 6, pp. 3413–3423, Jun. 2022, doi: 10.1016/j.jksuci.2021.01.014.
 - [21] V. R. Joseph, "Optimal ratio for data splitting," *Stat. Anal. Data Min. ASA Data Sci. J.*, vol. 15, no. 4, pp. 531–538, Aug. 2022, doi: 10.1002/sam.11583.
 - [22] V. Djuricic, L. Kascelan, S. Rogic, and B. Melovic, "Bank CRM Optimization Using Predictive Classification Based on the Support Vector Machine Method," *Appl. Artif. Intell.*, vol. 34, no. 12, pp. 941–955, Oct. 2020, doi: 10.1080/08839514.2020.1790248.
 - [23] X. Tang and Y. Zhu, "Enhancing bank marketing strategies with ensemble learning: Empirical analysis," *PLOS ONE*, vol. 19, no. 1, p. e0294759, Jan. 2024, doi: 10.1371/journal.pone.0294759.
 - [24] Y. Ali, E. Awwad, M. Al-Razgan, and A. Maarouf, "Hyperparameter Search for Machine Learning Algorithms for Optimizing the Computational Complexity," *Processes*, vol. 11, no. 2, p. 349, Jan. 2023, doi: 10.3390/pr11020349.
 - [25] J. Allgaier and R. Pryss, "Cross-Validation Visualized: A Narrative Guide to Advanced Methods," *Mach. Learn. Knowl. Extr.*, vol. 6, no. 2, pp. 1378–1388, Jun. 2024, doi: 10.3390/make6020065.
 - [26] Department of Computer Science and Informatics, University of Energy and Natural Resources, Sunyani, Ghana, I. K. Nti, O. Nyarko-Boateng, and J. Aning, "Performance of Machine Learning Algorithms with Different K Values in K-fold CrossValidation," *Int. J. Inf. Technol. Comput. Sci.*, vol. 13, no. 6, pp. 61–71, Dec. 2021, doi: 10.5815/ijitcs.2021.06.05.
 - [27] U. H. Laksono and R. R. Suryono, "Sentiment Analysis Of Online Dating Apps Using Support Vector Machine And Naïve Bayes Algorithms," *J. Tek. Inform. Jutif*, vol. 6, no. 1, pp. 229–238, Feb. 2025, doi: 10.52436/1.jutif.2025.6.1.2105.
 - [28] I. Markoulidakis and G. Markoulidakis, "Probabilistic Confusion Matrix: A Novel Method for Machine Learning Algorithm Generalized Performance Analysis," *Technologies*, vol. 12, no. 7, p. 113, Jul. 2024, doi: 10.3390/technologies12070113.
 - [29] J. Li, H. Sun, and J. Li, "Beyond confusion matrix: learning from multiple annotators with awareness of instance features," *Mach. Learn.*, vol. 112, no. 3, pp. 1053–1075, Mar. 2023, doi: 10.1007/s10994-022-06211-x.
-

-
- [30] A. M. Carrington *et al.*, “A new concordant partial AUC and partial c statistic for imbalanced data in the evaluation of machine learning algorithms,” *BMC Med. Inform. Decis. Mak.*, vol. 20, no. 1, p. 4, Dec. 2020, doi: 10.1186/s12911-019-1014-6.
- [31] R. Kannan and V. Vasanthi, “Machine Learning Algorithms with ROC Curve for Predicting and Diagnosing the Heart Disease,” in *Soft Computing and Medical Bioinformatics*, in SpringerBriefs in Applied Sciences and Technology. , Singapore: Springer Singapore, 2019, pp. 63–72. doi: 10.1007/978-981-13-0059-2_8.
- [32] Q. M. Zhou, L. Zhe, R. J. Brooke, M. M. Hudson, and Y. Yuan, “A relationship between the incremental values of area under the ROC curve and of area under the precision-recall curve,” *Diagn. Progn. Res.*, vol. 5, no. 1, p. 13, Dec. 2021, doi: 10.1186/s41512-021-00102-w.
- [33] T. J. Firdaus, J. Indra, S. A. P. Lestari, and H. Hikmayanti, “Sentiment Analysis Of The Sambara Application Using The Support Vector Machine Algorithm,” *J. Tek. Inform. Jutif*, vol. 5, no. 4, pp. 1183–1192, Sep. 2024, doi: 10.52436/1.jutif.2024.5.4.2673.
- [34] I. Markoulidakis and G. Markoulidakis, “Probabilistic Confusion Matrix: A Novel Method for Machine Learning Algorithm Generalized Performance Analysis,” *Technologies*, vol. 12, no. 7, p. 113, Jul. 2024, doi: 10.3390/technologies12070113.

