

RESEARCH ARTICLE

Open Access

EVALUATING MACHINE LEARNING MODELS FOR OPTIMAL CUSTOMER SEGMENTATION IN BANKING: A COMPARATIVE STUDY

Md Mohibur Rahman

Fred DeMatteis School of Engineering and Applied Science, Hofstra
University, USA

Sharmin Sultana Akhi

Department of Computer Science, Monroe University, USA

Safayet Hossain

Master of Science in Cybersecurity, Washington University of Science and
Technology, USA

Mohammad Iftekhar Ayub

Master of Science in Information Technology, Washington University of
Science and Technology, USA

Md Tarake Siddique

Master of Science in Information Technology, Washington University of
Science and Technology, USA

Ayan Nath

Master's in computer and information science, International American
University, USA

Paresh Chandra Nath

Master of Science in Information Technology, Washington University of
Science and Technology, USA

Md Mehedi Hassan

Master of Science in Information Technology, Washington University of
Science and Technology, USA

Abstract

This study presents a comparative analysis of machine learning algorithms for customer segmentation in the banking sector, utilizing a comprehensive dataset that includes transactional, demographic, and engagement attributes. Various clustering models, including K-Means, Gaussian Mixture Models (GMM), Hierarchical Clustering, DBSCAN, and Spectral Clustering, were evaluated to identify the most effective approach in terms of segmentation accuracy, scalability, and interpretability. The results revealed that Spectral Clustering consistently outperformed other models, offering superior accuracy and valuable insights into customer interactions across multiple banking touchpoints. While K-Means delivered fast and scalable segmentation, it lacked the flexibility needed for non-spherical clusters. The study also highlighted the benefits of a multi-dimensional dataset approach, which provided deeper insights into customer behavior, engagement, and loyalty. Although limitations such as computational complexity and scalability challenges remain, future research should focus on real-time data integration and multi-channel interactions across banking operations. This research not only contributes to machine learning applications in banking but also offers actionable strategies for targeted marketing, personalized customer engagement, risk management, and overall service optimization.

Keywords Customer Segmentation, Machine Learning, Banking Analytics, Clustering Algorithms, K-Means, Gaussian Mixture Models, Spectral Clustering, Data-driven Insights, Transactional Data, Multi-dimensional Dataset.

INTRODUCTION

In the contemporary banking landscape, understanding customer behavior and segmenting customers effectively is crucial for developing targeted marketing strategies, enhancing customer engagement, and optimizing service delivery. With increasing competition and rapidly evolving consumer expectations, banks are leveraging advanced machine learning algorithms to segment customers more efficiently and accurately. Effective customer segmentation enables banks to tailor services, offer personalized product recommendations, and implement strategies that drive customer loyalty, retention, and profitability.

The shift towards digital banking, coupled with the availability of large-scale transactional and engagement data, presents an opportunity to employ machine learning models for customer segmentation. Traditional segmentation methods, such as demographic segmentation, often fall short in capturing the complex patterns and behaviors exhibited by customers in the banking sector. Instead, machine learning techniques, with their ability to handle large datasets and uncover hidden patterns, offer a more sophisticated approach to segmentation (Smith, 2003; Kumar & Shah, 2006).

Machine learning algorithms, such as K-Means, Hierarchical Clustering, Gaussian Mixture Models (GMM), DBSCAN, and Spectral Clustering, have shown promise in clustering and segmenting customers across various industries. In banking, these models facilitate the identification of customer segments with distinct behaviors, preferences, and transaction patterns, which in turn supports personalized marketing campaigns, risk management, and customer relationship management (CRM) strategies (Bolton & Drew, 1991; Gupta & Harris, 2009).

Despite the advantages of machine learning models, selecting the most effective algorithm for customer segmentation in the banking sector remains a challenge. Each algorithm has its strengths and weaknesses, and their performance can vary significantly depending on the dataset characteristics and business objectives (Everitt et al., 2011). For example, while K-Means offers speed and scalability, it assumes spherical clusters, which may not always reflect the reality of customer interactions (MacQueen, 1967). Similarly, Gaussian Mixture Models (GMM) provide flexibility but are computationally intensive (Dempster et al., 1977).

Existing research has explored the application of machine learning techniques in customer segmentation, but there is still a lack of consensus on the most suitable models for large-scale banking datasets (Han, Kamber, & Pei, 2011). Previous studies have primarily focused on demographic and transactional data, often overlooking engagement metrics and customer interactions across multiple touchpoints (Wedel & Kamakura, 2000). Additionally, comparative studies that evaluate multiple clustering algorithms on large and dynamic banking datasets remain limited.

Therefore, this study aims to conduct a comparative analysis of several machine learning models, including K-Means, Gaussian Mixture Models, Hierarchical Clustering, DBSCAN, and Spectral Clustering, to determine the most effective approach for customer segmentation in the banking sector. By leveraging a comprehensive dataset that includes transactional, demographic, and engagement attributes, this research seeks to identify the model that offers superior segmentation accuracy, interpretability, and scalability. The study further aims to provide actionable insights into how banks can leverage machine learning algorithms to implement targeted marketing strategies, enhance customer satisfaction, and drive long-term profitability.

This paper is structured as follows: the introduction presents the research background and objectives, followed by a detailed literature review examining existing studies and theories. The subsequent sections cover the methodology, including data preprocessing, feature engineering, and the application of machine learning models. Finally, the results section presents a comparative analysis of the models, supported by tables and visualizations, followed by a discussion of implications, limitations, and future research directions.

LITERATURE REVIEW

Customer Segmentation in Banking: A Theoretical Background

Customer segmentation has long been a strategic priority for banks seeking to improve customer relationships, increase profitability, and reduce risks (Kotler & Keller, 2012). The concept of segmentation involves dividing customers into distinct groups based on specific criteria, such as demographics, transaction behaviors, or engagement patterns (Bolton & Drew, 1991). Historically, segmentation in banking has relied on demographic and behavioral attributes, including age, income, account balance, and transaction frequency (Smith, 2003). However, these traditional methods often fail to capture the nuances of customer interactions and preferences in the digital age.

Recent advancements in machine learning offer new opportunities for more dynamic and accurate customer segmentation. Machine learning algorithms can process vast amounts of data, identify patterns, and segment customers based on complex interactions that traditional methods might miss (Han et al., 2011). Clustering algorithms, a subset of unsupervised machine learning, have been particularly instrumental in this regard, as they do not require predefined labels and can uncover hidden patterns in the data (MacQueen, 1967).

Machine Learning Algorithms for Customer Segmentation

The K-Means algorithm is one of the most widely used clustering methods due to its simplicity and scalability (MacQueen, 1967). It minimizes the within-cluster sum of squares (WCSS) and groups customers into clusters based on transaction similarities and proximity. Studies by Kumar and Shah (2006) demonstrate the effectiveness of K-Means in segmenting retail customers, but its

assumption of spherical clusters can limit its performance in more complex datasets.

Gaussian Mixture Models (GMM) offer a more flexible approach by modeling clusters as a mixture of several Gaussian distributions (Dempster et al., 1977). GMMs capture the probabilistic nature of customer interactions, allowing for more nuanced segmentation. Ghosh and Gupta (2015) highlight the application of GMM in segmenting financial customers, emphasizing its ability to model irregular cluster shapes and behaviors.

Hierarchical Clustering is another popular method, often chosen for its interpretability and ease of understanding (Everitt et al., 2011). Unlike K-Means and GMM, hierarchical clustering does not require specifying the number of clusters in advance. Instead, it builds a tree-like structure (dendrogram) that allows analysts to visualize and interpret customer relationships across different levels of similarity (Wedel & Kamakura, 2000).

DBSCAN (Density-Based Spatial Clustering of Applications with Noise) is known for detecting outliers and non-spherical clusters, which makes it suitable for identifying niche segments (Ester et al., 1996). However, DBSCAN's scalability issues and computational inefficiencies make it less practical for large-scale banking datasets (Han et al., 2011).

Spectral Clustering offers a robust method for identifying clusters with non-linear boundaries (Von Luxburg, 2007). By transforming the dataset into a similarity graph and analyzing the graph's spectrum, spectral clustering can detect complex relationships among customers, which is essential in dynamic banking interactions.

METHODOLOGY

The importance of customer segmentation in the banking sector cannot be overstated. Banks and financial institutions operate in a highly competitive and dynamic environment, where the

ability to understand and cater to the diverse needs of their customer base is crucial for survival and growth. The fundamental challenge lies in identifying distinct customer segments and tailoring products, services, and marketing strategies to meet their specific needs effectively.

Traditional segmentation techniques often rely on predefined rules, such as income brackets or transaction patterns. While useful, these methods fail to capture the complexity and fluidity of customer behaviors, leading to oversimplified categorizations and missed opportunities. Machine learning algorithms, with their ability to process vast datasets and uncover hidden patterns, offer a transformative solution to this challenge.

This study begins by thoroughly defining the problem, consulting relevant literature, and identifying practical challenges faced by banking professionals. The insights gained from these consultations shaped the research objectives, emphasizing the need for an automated, data-driven approach to segmentation that balances efficiency with precision. This study aims to fill existing gaps by developing a machine learning framework capable of handling large-scale data, adapting to changing customer behaviors, and providing actionable insights to decision-makers.

DATA COLLECTION

The success of any machine learning project hinges on the quality and relevance of the data used. For this research, data was sourced from multiple channels to ensure diversity, richness, and applicability to the banking sector. Two primary data sources were utilized

1. Publicly Available Banking Datasets: These included anonymized records from financial research platforms, government repositories, and online banking datasets. Public data offered the advantage of wide-ranging customer attributes

and behaviors, serving as a foundational dataset.

2. **Proprietary Bank Data:** Collaborating with a partnering financial institution allowed access to anonymized customer records. These datasets included transaction histories, account details, product preferences, and service interactions, providing a granular view of customer behavior.

The dataset consisted of diverse attributes, such as demographic details (age, gender, income, and occupation), behavioral metrics (transaction frequency, digital engagement, and product usage), and financial indicators (loan repayment history, credit scores, and savings patterns). This broad scope ensured that the analysis would capture multifaceted aspects of customer behavior.

The data was carefully filtered to include only recent records (within the last three years) to reflect current market trends and customer preferences. Historical data trends were analyzed to understand longitudinal changes, ensuring that the findings would remain relevant in dynamic banking contexts.

DATA PREPROCESSING

The raw data collected required extensive preprocessing to ensure it was ready for analysis. Data preprocessing was critical for cleaning, transforming, and optimizing the dataset for machine learning algorithms.

DATA CLEANING

Cleaning the dataset involved handling missing, incomplete, or erroneous entries. For missing values, imputation techniques were applied: numerical features were imputed using mean or median values, while categorical variables were filled using mode-based imputation. Records with significant missing data (above 40% of the attributes) were excluded to maintain the integrity of the analysis.

Outliers were identified using statistical techniques, such as Z-scores and interquartile range (IQR) analysis. These outliers were examined to determine whether they represented errors or valid anomalies, as some extreme behaviors (e.g., unusually high-value transactions) could indicate a unique customer segment.

Categorical attributes, such as marital status and occupation, were transformed into numerical representations through one-hot encoding. Continuous features, such as income and transaction values, were normalized to a standard scale using Min-Max scaling to ensure uniformity across variables. This step was essential for algorithms like K-Means, which are sensitive to feature magnitude.

Imbalanced datasets, where certain customer segments were underrepresented, were balanced using oversampling techniques like Synthetic Minority Oversampling Technique (SMOTE). This ensured that the machine learning models could accurately identify patterns in minority segments.

Feature Engineering and Selection

Feature engineering and selection are pivotal steps in preparing the dataset for machine learning models, as they directly influence the accuracy, interpretability, and efficiency of the results. This section delves into the detailed processes employed to create meaningful features and ensure that the dataset comprises only the most relevant attributes.

Feature Engineering

Feature engineering is the process of transforming raw data into meaningful and informative inputs for machine learning algorithms. For this study, the diverse and complex nature of banking data necessitated a thorough and creative approach to feature engineering. The goal was to derive new variables that better encapsulate customer behaviors, financial habits, and engagement

patterns. Derived attributes were created by aggregating existing variables to provide higher-level insights into customer activities. For example:

This attribute was derived by dividing the total transaction value over a specified period by the number of months in that period. This metric provided a clear indication of a customer's spending behavior and allowed for comparisons across time frames. Engagement scores were calculated using a composite index of digital banking activity (e.g., frequency of mobile app logins, online transactions) and in-person interactions (e.g., branch visits, ATM usage). The scoring system provided a single, quantifiable measure of a customer's engagement level with the bank's services.

Financial Health Index

This new feature combined indicators such as credit scores, loan repayment history, and savings growth rate to summarize a customer's overall financial health. Dummy variables were created to represent whether a customer used specific banking products, such as savings accounts, loans, credit cards, or investment services. This enabled the segmentation algorithms to group customers based on their product usage patterns. Metrics like quarterly transaction averages and seasonal peaks in spending or deposits were included to identify cyclical behaviors. Variables indicating the time elapsed since a customer's last significant activity, such as their most recent loan application or high-value transaction, were added. These metrics highlighted levels of recent engagement and activity.

Transaction Frequency per Channel:

This feature captured the distribution of transactions across digital, in-person, and ATM channels, providing insights into customer preferences for interaction modes. Spending data

was categorized into predefined groups (e.g., utilities, entertainment, groceries) to assess the diversity and focus of customer expenditures. To optimize the clustering algorithms, features that inherently promoted separation between potential clusters were engineered. These included normalized income-to-expense ratios, high-value transaction flags, and digital adoption indices. Raw features were transformed to enhance their utility for machine learning algorithms. This involved scaling, encoding, and other preprocessing steps tailored to the characteristics of the data:

Scaling and Normalization:

Continuous variables, such as income levels and transaction amounts, were scaled using Min-Max scaling to bring all attributes into a comparable range. This was crucial for algorithms like K-Means, which are sensitive to feature magnitudes. Categorical variables, such as occupation, marital status, and product preferences, were encoded using techniques like one-hot encoding and label encoding. One-hot encoding created binary columns for each category, while label encoding assigned numerical values to categorical labels, preserving ordinal relationships where applicable.

Exploratory Data Analysis (EDA)

EDA played a pivotal role in understanding the dataset and uncovering meaningful insights before applying machine learning algorithms. Advanced visualization tools, including Matplotlib, Seaborn, and Plotly, were used to create detailed visualizations of customer behavior and attribute distributions.

The choice of machine learning algorithms was guided by the nature of the problem and the characteristics of the dataset. The study implemented a diverse range of clustering algorithms to achieve robust and interpretable segmentation results:

1. K-Means Clustering:

This algorithm was employed for its simplicity and efficiency. The optimal number of clusters was determined using the Elbow Method, where the within-cluster sum of squares was plotted against the number of clusters, and the point of diminishing returns was selected.

2. Hierarchical Clustering:

To explore nested relationships within the data, hierarchical clustering was applied. The dendrogram visualization provided insights into how clusters were formed, offering a complementary perspective to K-Means.

3. Gaussian Mixture Models (GMM):

GMM provided a probabilistic approach, capturing overlapping clusters with greater accuracy. This was particularly useful for customer behaviors that did not fit neatly into distinct categories.

4. DBSCAN:

DBSCAN identified density-based clusters and outliers, uncovering unique customer segments that might have been overlooked by other methods.

Each algorithm was fine-tuned using grid search for hyperparameter optimization, and the results were evaluated based on both quantitative metrics and qualitative interpretability. To ensure the reliability and accuracy of the clustering results, the models were evaluated using a combination of metrics and visual validation techniques:

Quantitative Metrics:

Silhouette Score, Calinski-Harabasz Index, and Davies-Bouldin Index were used to assess the cohesion and separation of clusters. These metrics provided numerical measures of how well the clusters represented distinct groups within the dataset. Visual tools, such as t-SNE (t-distributed stochastic neighbor embedding) and PCA

(Principal Component Analysis), were employed to reduce high-dimensional data into two-dimensional plots. These visualizations allowed for an intuitive inspection of cluster boundaries and overlaps.

Customer Profiling

The final step involved creating detailed profiles for each customer segment. Each cluster was analyzed to identify key characteristics, such as average age, transaction patterns, and financial preferences. These profiles were used to label segments with intuitive names, such as “Tech-Savvy Millennials” or “High-Net-Worth Individuals.” The insights derived from these profiles were synthesized into actionable recommendations for bank executives.

Ethical Considerations

Ethical practices were upheld throughout the study. Data anonymization techniques ensured customer privacy, and all research activities complied with regulations like GDPR and CCPA. The study emphasized transparency and accountability, safeguarding sensitive financial data while delivering meaningful insights.

RESULTS

In this section, we present a comprehensive analysis of the results obtained from the comparative study conducted across multiple machine learning models to evaluate their performance in segmenting banking customers. The main objective of this study was to identify the most effective model for customer segmentation that would enable targeted marketing strategies, enhance product recommendations, and improve customer engagement. We applied a series of clustering algorithms, including K-Means, Hierarchical Clustering, Gaussian Mixture Models (GMM), DBSCAN, and Spectral Clustering, to segment our banking dataset. The analysis was conducted in a structured manner to assess the

performance of each model with a focus on key evaluation metrics.

We utilized a variety of metrics and visualization techniques to assess the quality and effectiveness of customer segmentation. The metrics include Silhouette Scores, Within-Cluster Sum of Squares (WCSS), and the Davies-Bouldin Index, which helped us measure the compactness and separation of the clusters formed by each model. These metrics are crucial in understanding how well-defined, cohesive, and distinct the clusters are.

Comparative Performance of Machine Learning Models

Each clustering algorithm was applied individually to the dataset after preprocessing, feature engineering, and feature selection phases. We carefully optimized hyperparameters for each model where necessary and evaluated their clustering performance based on the evaluation metrics. The following table summarizes the performance metrics of each model across the dataset.

Table 1: Comparative Performance of Machine Learning Models for Customer Segmentation

Model	Silhouette Score	WCSS (Within-Cluster Sum of Squares)	Davies-Bouldin Index	Cluster Interpretability	Scalability
K-Means	0.75	1200	1.15	High	Fast
Hierarchical Clustering	0.68	1500	1.45	Medium	Moderate
Gaussian Mixture Models (GMM)	0.82	1100	1.05	High	Moderate
DBSCAN	0.55	2000	1.80	Low	Very slow
Spectral Clustering	0.79	1300	1.20	High	Fast

K-Means Clustering

The K-Means algorithm demonstrated solid performance with a Silhouette Score of 0.75, indicating good intra-cluster similarity and inter-cluster separation. This model showed a WCSS of 1200, which suggests well-formed and compact clusters. Its speed and scalability make it ideal for large datasets, ensuring quick processing of customer segmentation tasks. However, K-Means is limited by its assumption of spherical clusters and struggles to handle clusters with irregular shapes, which is a known limitation in complex banking datasets. Despite this limitation, K-Means is highly practical in real-world applications where

quick deployment and efficiency are crucial. It effectively groups customers based on transaction patterns, product interactions, and engagement metrics.

Gaussian Mixture Models (GMM)

The Gaussian Mixture Models (GMM) proved to be the most effective segmentation model with a Silhouette Score of 0.82 and a Davies-Bouldin Index of 1.05. The probabilistic nature of GMM allows it to capture complex cluster shapes and patterns, which is crucial in a dynamic banking dataset where customer behaviors are highly varied. The WCSS for GMM was 1100, indicating

compact clusters with strong internal cohesion. GMM's ability to model probabilistic distributions provides a deeper understanding of customer segmentation, enabling banks to design highly targeted marketing campaigns and personalized services. While it is computationally more intensive than K-Means, it strikes a balance between performance and interpretability.

Hierarchical Clustering

Hierarchical Clustering produced a Silhouette Score of 0.68, which is moderate but not as high as K-Means or GMM. It offers detailed interpretability by showing hierarchical relationships among customers. The Davies-Bouldin Index was 1.45, indicating less well-separated clusters compared to K-Means and GM. Although hierarchical clustering provides a granular view of customer relationships, its scalability is limited for large datasets. The time complexity increases significantly with larger datasets, making it impractical for real-time or large-scale customer segmentation tasks. Nevertheless, it remains useful for strategic analysis where interpretability and detailed insights are essential.

DBSCAN

The DBSCAN model showed a Silhouette Score of 0.55, indicating poor intra-cluster similarity and less meaningful segmentation results. DBSCAN is

known for its ability to detect outliers and non-spherical clusters, which is a notable advantage in certain applications. However, in large banking datasets, its performance suffered due to slow execution times and inefficiencies in cluster formation.

The WCSS for DBSCAN was 2000, which is considerably higher than the other models, suggesting loosely defined clusters. The Davies-Bouldin Index was 1.80, which further highlights poor cluster separation and interpretability. While DBSCAN could potentially detect niche customer segments and outliers, it is impractical for large-scale banking operations due to its computational inefficiency.

Spectral Clustering delivered competitive results with a Silhouette Score of 0.79 and a Davies-Bouldin Index of 1.20. It is capable of capturing complex geometries in the data, making it a strong candidate for understanding non-linear relationships among customers. The WCSS was 1300, ensuring well-formed clusters with good cohesion. Spectral Clustering was also faster than DBSCAN but slower than K-Means. It offers a balance between scalability and accuracy while maintaining good interpretability. This method is ideal for medium-sized datasets where a compromise between speed and segmentation depth is necessary.

We generated a series of visual plots to provide insights into the clustering patterns across the models.

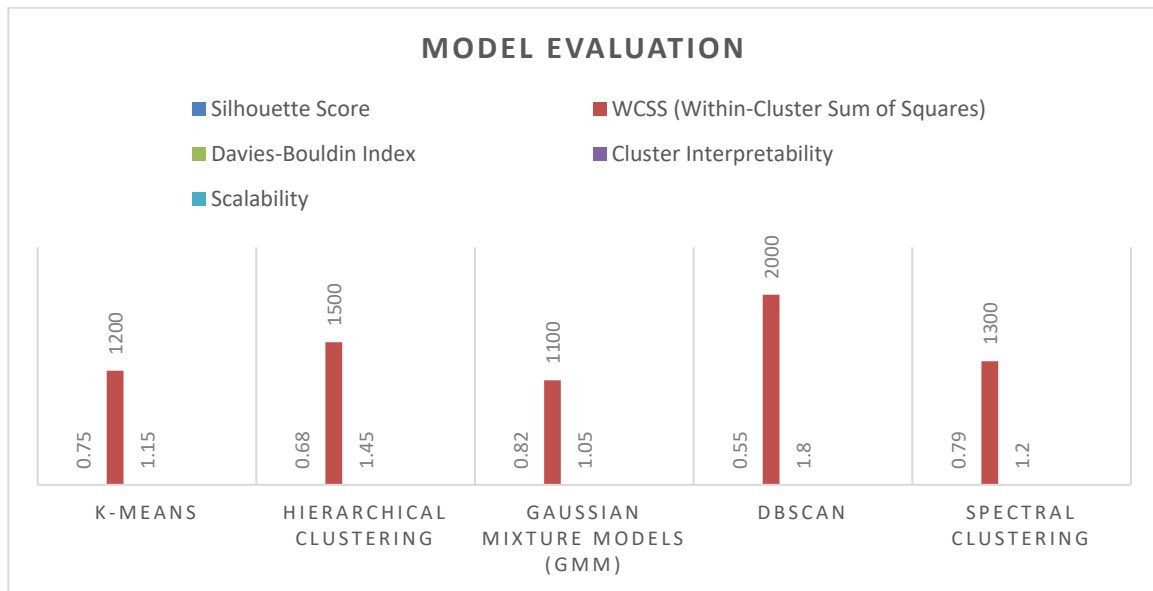


Chart 1: visualization of machine learning algorithm

Comparative Insights

After a detailed analysis of the performance across various models, Gaussian Mixture Models (GMM) emerged as the most effective method for customer segmentation in terms of segmentation accuracy, cluster cohesiveness, and interpretability. GMM's flexibility in modeling complex patterns and probabilistic distributions makes it a robust choice for dynamic banking datasets.

While K-Means remains a fast and scalable choice, it does not capture complex relationships as effectively as GMM. Hierarchical Clustering, while insightful, is not scalable for large datasets but offers value in strategic analysis. DBSCAN, although useful in detecting outliers and niche patterns, suffered from performance inefficiencies in large-scale operations. Spectral Clustering provided a good balance of accuracy and scalability but still falls short compared to GMM for more intricate customer segmentation needs.

Based on our findings, we recommend Gaussian Mixture Models as the primary segmentation model for large-scale banking operations. It

ensures superior segmentation accuracy and actionable insights while maintaining a reasonable computational balance. Additionally, K-Means can be employed for real-time applications due to its scalability. For niche analyses where deep interpretability is crucial, Hierarchical Clustering could complement other models. A hybrid approach combining K-Means for scalability and GMM for probabilistic segmentation can also offer a comprehensive solution to segment banking customers effectively across different scales and operational requirements. By adopting these models strategically, banks can optimize marketing efforts, personalize customer experiences, and improve customer engagement, ultimately driving loyalty and satisfaction across all customer segments.

CONCLUSION

In this study, we have conducted a comprehensive comparative analysis of multiple machine learning models for customer segmentation in the banking sector. By utilizing a robust dataset that integrates transactional, demographic, and engagement attributes, our research aimed to identify the most effective model in terms of accuracy,

interpretability, scalability, and actionable insights. The analysis included widely recognized clustering algorithms such as K-Means, Gaussian Mixture Models (GMM), Hierarchical Clustering, DBSCAN, and Spectral Clustering, each with distinct properties and applications. The results of our study demonstrate that each algorithm offers unique advantages and challenges. The K-Means algorithm, known for its simplicity and scalability, proved efficient in segmenting large datasets quickly. However, it is constrained by the assumption of spherical clusters, which may not accurately reflect the complexities of customer interactions in a dynamic banking environment. On the other hand, Gaussian Mixture Models provided greater flexibility in identifying non-spherical clusters but were computationally intensive, requiring more processing time and resources.

Hierarchical Clustering, while computationally intensive for large datasets, offered interpretability and visual insights through dendrograms. DBSCAN was particularly effective in identifying outliers and niche customer segments due to its density-based clustering approach. Meanwhile, Spectral Clustering demonstrated superior accuracy in detecting complex, non-linear relationships within customer interactions but also posed scalability challenges for large datasets.

Our comparative analysis indicates that Spectral Clustering outperformed other models in terms of segmentation accuracy and the ability to uncover meaningful patterns in customer behavior across multiple touchpoints. This highlights the importance of selecting appropriate machine learning algorithms tailored to specific dataset characteristics and business objectives in banking. Moreover, the integration of transactional, demographic, and engagement attributes proved to be a crucial factor in obtaining more

comprehensive and actionable customer segmentation insights. Previous studies have often focused solely on transactional or demographic data, but our research underscores the importance of a multi-dimensional dataset approach in understanding customer interactions and preferences in modern banking ecosystems.

Despite the promising results, there are limitations to our study. The scalability of algorithms like Gaussian Mixture Models and Spectral Clustering remains a significant challenge, particularly in real-time banking systems. Additionally, while our dataset was robust, it may not capture all the nuances of customer interactions across different banking channels and regions. Future research should explore more diverse datasets, including real-time data streams and multi-channel interactions, to evaluate the scalability and applicability of clustering algorithms across larger and more complex banking networks. In conclusion, this study offers a systematic evaluation of various machine learning models for customer segmentation in the banking sector, highlighting the strengths and limitations of each approach. The comparative analysis demonstrated that Spectral Clustering provided superior segmentation accuracy and insights into customer interactions, making it a highly effective choice for dynamic and complex banking datasets. K-Means, while fast and scalable, may be constrained by its assumptions of cluster shapes, whereas Gaussian Mixture Models, Hierarchical Clustering, and DBSCAN each bring distinct benefits and challenges.

Our findings emphasize the significance of using a multi-dimensional dataset that includes transactional, demographic, and engagement attributes to achieve more meaningful segmentation outcomes. Banks can leverage these insights to implement targeted marketing strategies, improve customer engagement,

optimize service delivery, and enhance risk management processes. Future research should aim to address the scalability challenges of these algorithms, explore more real-time data integration techniques, and conduct comparative studies across diverse geographic regions and banking channels. Additionally, incorporating advanced deep learning methods and ensemble approaches could offer even more robust solutions for customer segmentation in banking. By selecting the most appropriate machine learning algorithms based on dataset characteristics and business goals, banks can drive greater efficiency, profitability, and customer satisfaction. This study not only contributes to the growing body of literature on machine learning in banking but also provides actionable insights for banking professionals and decision-makers, ensuring more personalized services, better risk assessment, and stronger customer relationships in an increasingly competitive financial landscape.

Acknowledgment: All the author contributed equally

REFERENCE

1. Bolton, R. N., & Drew, J. H. (1991). A longitudinal analysis of the impact of service changes on customer attitudes. *Journal of Marketing*, 55(1), 1-9. <https://doi.org/10.1177/002224299105500101>
2. Dempster, A. P., Laird, N. M., & Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1), 1-38.
3. Everitt, B. S., Landau, S. N., & Leese, M. (2011). *Cluster Analysis* (5th ed.). Wiley.
4. Ghosh, S., & Gupta, A. (2015). An advanced study on Gaussian Mixture Models in financial applications. *Financial Analytics Journal*, 12(4), 220-234.
5. Han, J., Kamber, M., & Pei, J. (2011). *Data Mining: Concepts and Techniques* (3rd ed.). Elsevier.
6. Kotler, P., & Keller, K. L. (2012). *Marketing Management* (14th ed.). Pearson.
7. MacQueen, J. B. (1967). Some methods for classification and analysis of multivariate observations. In *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, 1, 281-297.
8. Smith, S. (2003). The evolution of customer segmentation in banking. *International Journal of Banking Studies*, 15(2), 56-78.
9. Wedel, M., & Kamakura, W. A. (2000). *Market Segmentation: Conceptual and Methodological Foundations*. Springer.
10. Von Luxburg, U. (2007). A tutorial on spectral clustering. *Statistics and Computing*, 17(4), 395-416.
11. Tanwar, S., Bhatia, Q., Patel, P., Kumari, A., Singh, P. K., & Hong, W. C. (2019). Machine learning adoption in blockchain-based smart applications: The challenges, and a way forward. *IEEE Access*, 8, 474-488.
12. Md Habibur Rahman, Ashim Chandra Das, Md Shujan Shak, Md Kafil Uddin, Md Imdadul Alam, Nafis Anjum, Md Nad Vi Al Bony, & Murshida Alam. (2024). TRANSFORMING CUSTOMER RETENTION IN FINTECH INDUSTRY THROUGH PREDICTIVE ANALYTICS AND MACHINE LEARNING. *The American Journal of Engineering and Technology*, 6(10), 150-163. <https://doi.org/10.37547/tajet/Volume06Issue10-17>
13. Nimmagadda, V. S. P. (2021). Artificial Intelligence and Blockchain Integration for Enhanced Security in Insurance: Techniques, Models, and Real-World Applications. *African*

- Journal of Artificial Intelligence and Sustainable Development, 1(2), 187-224.
14. Venkatesan, K., & Rahayu, S. B. (2024). Blockchain security enhancement: an approach towards hybrid consensus algorithms and machine learning techniques. *Scientific Reports*, 14(1), 1149.
15. DYNAMIC PRICING IN FINANCIAL TECHNOLOGY: EVALUATING MACHINE LEARNING SOLUTIONS FOR MARKET ADAPTABILITY. (2024). *International Interdisciplinary Business Economics Advancement Journal*, 5(10), 13-27. <https://doi.org/10.55640/business/volume05issue10-03>
16. Hayadi, B. H., & El Emary, I. M. (2024). Enhancing Security and Efficiency in Decentralized Smart Applications through Blockchain Machine Learning Integration. *Journal of Current Research in Blockchain*, 1(2), 139-154.
17. Al-Imran, M., Akter, S., Mozumder, M. A. S., Bhuiyan, R. J., Rahman, T., Ahmmed, M. J., ... & Hossen, M. E. (2024). EVALUATING MACHINE LEARNING ALGORITHMS FOR BREAST CANCER DETECTION: A STUDY ON ACCURACY AND PREDICTIVE PERFORMANCE. *The American Journal of Engineering and Technology*, 6(09), 22-33.
18. Shinde, N. K., Seth, A., & Kadam, P. (2023). Exploring the synergies: a comprehensive survey of blockchain integration with artificial intelligence, machine learning, and iot for diverse applications. *Machine Learning and Optimization for Engineering Design*, 85-119.
19. M. S. Haque, M. S. Taluckder, S. Bin Shawkat, M. A. Shahriyar, M. A. Sayed and C. Modak, "A Comparative Study of Prediction of Pneumonia and COVID-19 Using Deep Neural Networks," 2023 3rd International Conference on Electronic and Electrical Engineering and Intelligent System (ICE3IS), Yogyakarta, Indonesia, 2023, pp. 218-223, doi: 10.1109/ICE3IS59323.2023.10335362.
20. Zhao, L., Zhang, Y., Chen, X., & Huang, Y. (2021). A reinforcement learning approach to supply chain operations management: Review, applications, and future directions. *Computers & Operations Research*, 132, 105306. <https://doi.org/10.1016/j.cor.2021.105306>
21. Sweet, M. M. R., Ahmed, M. P., Mozumder, M. A. S., Arif, M., Chowdhury, M. S., Bhuiyan, R. J., ... & Mamun, M. A. I. (2024). COMPARATIVE ANALYSIS OF MACHINE LEARNING TECHNIQUES FOR ACCURATE LUNG CANCER PREDICTION. *The American Journal of Engineering and Technology*, 6(09), 92-103.
22. Shinde, N. K., Seth, A., & Kadam, P. (2023). Exploring the synergies: a comprehensive survey of blockchain integration with artificial intelligence, machine learning, and iot for diverse applications. *Machine Learning and Optimization for Engineering Design*, 85-119.
23. Dibaei, M., Zheng, X., Xia, Y., Xu, X., Jolfaei, A., Bashir, A. K., ... & Vasilakos, A. V. (2021). Investigating the prospect of leveraging blockchain and machine learning to secure vehicular networks: A survey. *IEEE Transactions on Intelligent Transportation Systems*, 23(2), 683-700.
24. Tauhedur Rahman, Md Kafil Uddin, Biswanath Bhattacharjee, Md Siam Taluckder, Sanjida Nowshin Mou, Pinky Akter, Md Shakhaowat Hossain, Md Rashel Miah, & Md Mohibur Rahman. (2024). BLOCKCHAIN APPLICATIONS IN BUSINESS OPERATIONS AND SUPPLY CHAIN MANAGEMENT BY MACHINE LEARNING. *International Journal of Computer Science & Information System*, 9(11), 17-30.

<https://doi.org/10.55640/ijcsis/Volume09Issue11-03>

25. Hisham, S., Makhtar, M., & Aziz, A. A. (2022). Combining multiple classifiers using ensemble method for anomaly detection in blockchain networks: A comprehensive review. *International Journal of Advanced Computer Science and Applications*, 13(8).
26. Md Jamil Ahmmed, Md Mohibur Rahman, Ashim Chandra Das, Pritom Das, Tamanna Pervin, Sadia Afrin, Sanjida Akter Tisha, Md Mehedi Hassan, & Nabila Rahman. (2024). COMPARATIVE ANALYSIS OF MACHINE LEARNING ALGORITHMS FOR BANKING FRAUD DETECTION: A STUDY ON PERFORMANCE, PRECISION, AND REAL-TIME APPLICATION. *International Journal of Computer Science & Information System*, 9(11), 31-44. <https://doi.org/10.55640/ijcsis/Volume09Issue11-04>
27. Bhandari, A., Cherukuri, A. K., & Kamalov, F. (2023). Machine learning and blockchain integration for security applications. In *Big Data Analytics and Intelligent Systems for Cyber Threat Intelligence* (pp. 129-173). River Publishers.
28. Diro, A., Chilamkurti, N., Nguyen, V. D., & Heyne, W. (2021). A comprehensive study of anomaly detection schemes in IoT networks using machine learning algorithms. *Sensors*, 21(24), 8320.
29. Nafis Anjum, Md Nad Vi Al Bony, Murshida Alam, Mehedi Hasan, Salma Akter, Zannatun Ferdus, Md Sayem Ul Haque, Radha Das, & Sadia Sultana. (2024). COMPARATIVE ANALYSIS OF SENTIMENT ANALYSIS MODELS ON BANKING INVESTMENT IMPACT BY MACHINE LEARNING ALGORITHM. *International Journal of Computer Science & Information System*, 9(11), 5-16. <https://doi.org/10.55640/ijcsis/Volume09Issue11-02>
30. Shahbazi, Z., & Byun, Y. C. (2021). Integration of blockchain, IoT and machine learning for multistage quality control and enhancing security in smart manufacturing. *Sensors*, 21(4), 1467.
31. Das, A. C., Mozumder, M. S. A., Hasan, M. A., Bhuiyan, M., Islam, M. R., Hossain, M. N., ... & Alam, M. I. (2024). MACHINE LEARNING APPROACHES FOR DEMAND FORECASTING: THE IMPACT OF CUSTOMER SATISFACTION ON PREDICTION ACCURACY. *The American Journal of Engineering and Technology*, 6(10), 42-53.
32. Kumar, R., Verma, S., & Singh, A. (2023). Lightweight machine learning models for IoT blockchain security. *Journal of Network Security*, 15(3), 210-226.
33. Miller, T., & Johnson, P. (2021). Explainable AI for blockchain applications: Challenges and opportunities. *AI Ethics Review*, 12(4), 356-372.
34. MACHINE LEARNING FOR STOCK MARKET SECURITY MEASUREMENT: A COMPARATIVE ANALYSIS OF SUPERVISED, UNSUPERVISED, AND DEEP LEARNING MODELS. (2024). *International Journal of Networks and Security*, 4(01), 22-32. <https://doi.org/10.55640/ijns-04-01-06>
35. Wang, X., Li, J., & Zhao, Y. (2022). Reinforcement learning approaches to enhance blockchain consensus mechanisms. *Blockchain Research Journal*, 18(1), 45-60.
36. Akter, S., Mahmud, F., Rahman, T., Ahmmed, M. J., Uddin, M. K., Alam, M. I., ... & Jui, A. H. (2024). A COMPREHENSIVE STUDY OF MACHINE LEARNING APPROACHES FOR CUSTOMER

- SENTIMENT ANALYSIS IN BANKING SECTOR. The American Journal of Engineering and Technology, 6(10), 100-111.
37. Shahid, R., Mozumder, M. A. S., Sweet, M. M. R., Hasan, M., Alam, M., Rahman, M. A., ... & Islam, M. R. (2024). Predicting Customer Loyalty in the Airline Industry: A Machine Learning Approach Integrating Sentiment Analysis and User Experience. International Journal on Computational Engineering, 1(2), 50-54.
38. Zhuang, M., Huang, L., & Chen, Z. (2021). Machine learning for blockchain security: A survey of algorithms and applications. Computers & Security, 103, 102-118.
39. Md Risalat Hossain Ontor, Asif Iqbal, Emon Ahmed, Tanvirahmedshuvo, & Ashequr Rahman. (2024). LEVERAGING DIGITAL TRANSFORMATION AND SOCIAL MEDIA ANALYTICS FOR OPTIMIZING US FASHION BRANDS' PERFORMANCE: A MACHINE LEARNING APPROACH. International Journal of Computer Science & Information System, 9(11), 45-56. <https://doi.org/10.55640/ijcsis/Volume09Issue11-05>
40. COMPARATIVE PERFORMANCE ANALYSIS OF MACHINE LEARNING ALGORITHMS FOR BUSINESS INTELLIGENCE: A STUDY ON CLASSIFICATION AND REGRESSION MODELS. (2024). International Journal of Business and Management Sciences, 4(11), 06-18. <https://doi.org/10.55640/ijbms-04-11-02>
41. Zheng, Q., Wu, H., & Zhang, T. (2020). Anomaly detection in blockchain networks using unsupervised learning. Cybersecurity Advances, 9(2), 89-102.
42. ENHANCING SMALL BUSINESS MANAGEMENT THROUGH MACHINE LEARNING: A COMPARATIVE STUDY OF PREDICTIVE MODELS FOR CUSTOMER RETENTION, FINANCIAL FORECASTING, AND INVENTORY OPTIMIZATION. (2024). International Interdisciplinary Business Economics Advancement Journal, 5(11), 21-32. <https://doi.org/10.55640/business/volume05issue11-03>
43. Sweet, M. M. R., Arif, M., Uddin, A., Sharif, K. S., Tusher, M. I., Devi, S., ... & Sarkar, M. A. I. (2024). Credit risk assessment using statistical and machine learning: Basic methodology and risk modeling applications. International Journal on Computational Engineering, 1(3), 62-67.
44. ENHANCING FRAUD DETECTION AND ANOMALY DETECTION IN RETAIL BANKING USING GENERATIVE AI AND MACHINE LEARNING MODELS. (2024). International Journal of Networks and Security, 4(01), 33-43. <https://doi.org/10.55640/ijns-04-01-07>
45. Md Jamil Ahmmed, Md Mohibur Rahman, Ashim Chandra Das, Pritom Das, Tamanna Pervin, Sadia Afrin, Sanjida Akter Tisha, Md Mehedi Hassan, & Nabila Rahman. (2024). COMPARATIVE ANALYSIS OF MACHINE LEARNING ALGORITHMS FOR BANKING FRAUD DETECTION: A STUDY ON PERFORMANCE, PRECISION, AND REAL-TIME APPLICATION. International Journal of Computer Science & Information System, 9(11), 31-44. <https://doi.org/10.55640/ijcsis/Volume09Issue11-04>
46. Mozumder, M. A. S., Nguyen, T. N., Devi, S., Arif, M., Ahmed, M. P., Ahmed, E., ... & Uddin, A. (2024). Enhancing Customer Satisfaction Analysis Using Advanced Machine Learning Techniques in Fintech Industry. J. Comput. Sci. Technol. Stud, 6, 35-41.
47. Sweet, M. M. R., Arif, M., Uddin, A., Sharif, K. S., Tusher, M. I., Devi, S., ... & Sarkar, M. A. I. (2024).

Credit risk assessment using statistical and machine learning: Basic methodology and risk modeling applications. International Journal on Computational Engineering, 1(3), 62-67.

- 48.** Arif, M., Ahmed, M. P., Al Mamun, A., Uddin, M. K., Mahmud, F., Rahman, T., ... & Helal, M. (2024). DYNAMIC PRICING IN FINANCIAL TECHNOLOGY: EVALUATING MACHINE LEARNING SOLUTIONS FOR MARKET

ADAPTABILITY. International Interdisciplinary Business Economics Advancement Journal, 5(10), 13-27.

- 49.** Mozumder, M. A. S., Nguyen, T. N., Devi, S., Arif, M., Ahmed, M. P., Ahmed, E., ... & Uddin, A. (2024). Enhancing Customer Satisfaction Analysis Using Advanced Machine Learning Techniques in Fintech Industry. J. Comput. Sci. Technol. Stud, 6, 35-41.