



An Analytical Study of Clustering Algorithms for Large-Scale Customer Segmentation

Samruddhi Sujit Patil¹, Prashant Devidas Shimpi²

¹PG Student, Department of Computer Science & Engg.,  : 0009-0002-4303-3550
Godavari Foundations, Godavari College of Engineering, Jalgaon, India, 425001

²Assistant Professor, Department of Computer Science & Engg.,  : 0009-0004-9588-4298
Godavari Foundations, Godavari College of Engineering, Jalgaon, India, 425001

Abstract – Customer segmentation is a cornerstone of data-driven marketing, personalization, and decision-making. With the rapid growth of digital platforms, organizations now handle large-scale, high-dimensional customer data, making traditional segmentation techniques insufficient. This analytical study examines major clustering algorithms used for customer segmentation, evaluates their suitability for large datasets, and compares their performance in terms of scalability, accuracy, interpretability, and computational complexity. The proposed framework leverages Mini-Batch K-Means, a batch-based learning algorithm that processes data incrementally in small subsets, significantly reducing memory usage and improving execution speed. Experimental results demonstrate that the proposed scalable clustering framework achieves faster convergence and better resource utilization compared to conventional techniques, proving suitable for real-world big data applications.

Keywords – Customer Segmentation, Clustering Algorithms, Mini-Batch K-Means, K-Means, DBSCAN, Scalable Clustering, Big Data, Unsupervised Learning

I. INTRODUCTION

Customer segmentation involves grouping customers with similar characteristics—such as purchasing behavior, demographics, or engagement patterns—to enable targeted strategies. Clustering, an unsupervised machine learning technique, is widely used for this task because it does not require labeled data.

In the modern digital era, the rapid advancement of information technology has resulted in the generation of massive volumes of data from business transactions, social media platforms, healthcare systems, IoT devices, and scientific research. This unprecedented growth, commonly referred to as Big Data, has made traditional data processing and analysis techniques inadequate.

Large-scale customer data introduces significant challenges: massive volume involving millions of records, high dimensionality spanning hundreds of features, presence of noise and outliers, and the need for real-time or near-real-time processing. Clustering is a fundamental technique in unsupervised machine learning that aims to group similar data objects into clusters such that objects within the same cluster exhibit high similarity, while objects in different clusters show significant dissimilarity.

Conventional clustering algorithms such as K-Means, Hierarchical Clustering, and DBSCAN were originally designed for small to medium-sized datasets. When applied to large-scale datasets, these algorithms face excessive memory consumption, increased execution time, and limited scalability. To address these challenges, scalable techniques such as Mini-Batch K-Means have been developed that process data incrementally in batches, enabling efficient clustering without loading entire datasets into memory.

Python has emerged as a powerful language for data science due to its simplicity and extensive library support. Libraries such as NumPy, Pandas, and Scikit-learn provide robust tools for preprocessing, clustering, and performance evaluation. This study focuses on an analytical review and implementation of a Scalable Data Clustering Framework for Large-Scale Customer Segmentation using batch-based techniques.

II. LITERATURE REVIEW

MacQueen [1] introduced K-Means clustering as a method for classification of multivariate observations, establishing the foundational centroid-based partitioning approach. Arthur and Vassilvitskii [2] proposed K-Means++ to improve initialization and reduce sensitivity to initial centroid selection, significantly improving convergence behavior.

Ester et al. [4] introduced DBSCAN, a density-based algorithm capable of identifying clusters of arbitrary shape and handling noise effectively. Zhang et al. [5] proposed BIRCH for very large databases, introducing a CF-tree structure that enables incremental and scalable clustering without loading the entire dataset into memory.

Sculley [6] presented web-scale K-Means clustering using mini-batches, demonstrating that processing small random subsets of data can approximate full-batch K-Means quality with substantially reduced computation time. This work laid the foundation for scalable clustering in big data environments.

Zaharia et al. [14] introduced Apache Spark as a unified engine for big data processing, enabling distributed clustering at scale. Research by Bifet et al. [13] on MOA (Massive Online Analysis) demonstrated the viability of streaming and incremental learning approaches. Pedregosa et al. [17] formalized the Scikit-learn framework, providing efficient Python implementations of Mini-Batch K-Means and related algorithms.

The collective literature indicates that batch-based and distributed clustering techniques provide an effective balance between performance and scalability, forming the basis for the proposed framework presented in this study.

III. METHODOLOGY

The proposed Scalable Clustering Framework for Large-Scale Customer Segmentation follows a modular pipeline-based architecture comprising five key components: Data Collection, Data Preprocessing, Batch Processing, Clustering Engine, and Evaluation Module. The framework is implemented in Python using Scikit-learn and Pandas.

Table 1. Comparison of Clustering Algorithms for Large-Scale Data

Algorithm	Scalability	Memory	Noise Handling	Speed	Best For
K-Means	Medium	High	Low	Moderate	Small/Medium
Mini-Batch K-Means	High	Low	Low	Fast	Large-Scale
DBSCAN	Low	High	High	Slow	Noise-heavy
Hierarchical	Low	Very High	Medium	Slow	Small Data
BIRCH	High	Low	Medium	Fast	Incremental

A. Data Collection

Large datasets are collected from publicly available repositories or generated synthetically in CSV format. The datasets contain numerical attributes relevant to customer segmentation such as transaction amount, purchase frequency, recency, and demographic attributes.

B. Data Preprocessing

Data preprocessing ensures effective clustering through: (i) handling missing values by mean imputation, (ii) removing duplicate records, and (iii) normalizing features using StandardScaler to ensure all dimensions contribute equally. Normalization is learned incrementally from the first data pass and applied consistently across all batches.

C. Batch-Wise Processing

Instead of loading the entire dataset into memory, data is processed in configurable small batches. This approach reduces memory usage from $O(n)$ to $O(B)$, where B is the batch size, and enables processing of datasets that exceed available RAM capacity.

D. Mini-Batch K-Means Clustering

Let $D = \{x_1, x_2, \dots, x_n\}$ be a dataset with n data points. The objective of clustering is to minimize the Within-Cluster Sum of Squares (WCSS):

$$J = \sum_{i=1}^K \sum_{x_j \in C_i} ||x_j - \mu_i||^2$$

where C_i represents the i -th cluster and μ_i is its centroid. Mini-Batch K-Means minimizes J using incremental centroid updates. Algorithm: (1) Initialize K centroids randomly; (2) Select random batch B ; (3) Assign each point to nearest centroid via Euclidean distance; (4) Update centroids using `partial_fit`; (5) Repeat steps 2–4 until convergence; (6) Assign final cluster labels.

E. Performance Evaluation

Clustering performance is evaluated using: execution time (seconds), peak memory utilization (MB), final inertia (WCSS value), and cluster stability across iterations. Results are compared against standard K-Means on equivalent data samples.

IV. RESULTS AND DISCUSSION

Experimental evaluation was conducted on synthetically generated customer datasets ranging from 100,000 to 5,000,000 records with 10–50 features. Mini-Batch K-Means was compared against standard K-Means and BIRCH across scalability, memory, and quality metrics.

Table 2. Performance Results on Large Customer Dataset (1M records)

Metric	K-Means	Mini-Batch K-Means	BIRCH
Exec. Time (s)	482	38	65
Memory (MB)	3840	210	380
Final Inertia	1.24×10^7	1.31×10^7	1.28×10^7
Scalability	Low	High	High

Results demonstrate that Mini-Batch K-Means achieves approximately 12.7x speedup over standard K-Means while consuming only 5.5% of the memory footprint on a one-million record dataset. The inertia values remain comparable, indicating that clustering quality is not significantly compromised by the batch-based approach. BIRCH shows intermediate performance, with slightly higher memory requirements than Mini-Batch K-Means.

For the customer segmentation use-case, Mini-Batch K-Means consistently produced stable cluster formations corresponding to identifiable customer groups such as high-frequency buyers, seasonal shoppers, and dormant customers. The batch processing approach proved particularly effective for datasets exceeding available RAM, enabling incremental processing without data truncation.

V. CONCLUSION

This analytical study highlights that no single clustering algorithm is universally optimal for large-scale customer segmentation. K-Means and its scalable variants remain dominant due to simplicity and efficiency. Density-based and model-based methods offer advantages in handling noise and complex cluster structures. The choice of algorithm should be guided by data characteristics, scale, and business objectives.

The proposed Mini-Batch K-Means based Scalable Clustering Framework achieves significant improvements in execution time and memory consumption compared to traditional approaches, with minimal loss in clustering quality. The framework is suitable for real-world big data applications including customer segmentation, IoT sensor data analysis, and large-scale pattern discovery. Future work will explore integration with distributed processing frameworks such as Apache Spark to further enhance scalability.

ACKNOWLEDGMENT

The authors acknowledge the support of their respective institutes and the open-source community for providing tools and datasets that facilitated this research.

REFERENCES

- [1] J. MacQueen, "Some methods for classification and analysis of multivariate observations," *Proc. 5th Berkeley Symp. on Mathematical Statistics and Probability*, vol. 1, pp. 281–297, 1967.
- [2] D. Arthur and S. Vassilvitskii, "k-means++: The advantages of careful seeding," *Proc. 18th Annual ACM-SIAM Symp. on Discrete Algorithms*, pp. 1027–1035, 2007.
- [3] S. Lloyd, "Least squares quantization in PCM," *IEEE Transactions on Information Theory*, vol. 28, no. 2, pp. 129–137, Mar. 1982.
- [4] M. Ester, H.-P. Kriegel, J. Sander, and X. Xu, "A density-based algorithm for discovering clusters in large spatial databases with noise," *Proc. 2nd Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, pp. 226–231, 1996.
- [5] T. Zhang, R. Ramakrishnan, and M. Livny, "BIRCH: An efficient data clustering method for very large databases," *ACM SIGMOD Record*, vol. 25, no. 2, pp. 103–114, 1996.
- [6] D. Sculley, "Web-scale k-means clustering," *Proc. 19th Int. World Wide Web Conf.*, pp. 1177–1178, 2010.
- [7] P. Berkhin, "A survey of clustering data mining techniques," *Grouping Multidimensional Data*, Springer, pp. 25–71, 2006.
- [8] J. Han, M. Kamber, and J. Pei, *Data Mining: Concepts and Techniques*, 3rd ed., Morgan Kaufmann, 2011.
- [9] A. Jain, M. Murty, and P. Flynn, "Data clustering: A review," *ACM Computing Surveys*, vol. 31, no. 3, pp. 264–323, 1999.
- [10] I. Dhillon, Y. Guan, and B. Kulis, "Weighted graph cuts without eigenvectors: A multilevel approach," *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol. 29, no. 11, pp. 1944–1957, 2007.
- [11] M. Kantardzic, *Data Mining: Concepts, Models, Methods, and Algorithms*, 2nd ed., Wiley-IEEE Press, 2011.
- [12] L. Kaufman and P. Rousseeuw, *Finding Groups in Data: An Introduction to Cluster Analysis*, Wiley, 2009.
- [13] A. Bifet, G. Holmes, R. Kirkby, and B. Pfahringer, "MOA: Massive online analysis," *Journal of Machine Learning Research*, vol. 11, pp. 1601–1604, 2010.
- [14] M. Zaharia et al., "Apache Spark: A unified engine for big data processing," *Communications of the ACM*, vol. 59, no. 11, pp. 56–65, 2016.
- [15] T. White, *Hadoop: The Definitive Guide*, 4th ed., O'Reilly Media, 2015.
- [16] J. Dean and S. Ghemawat, "MapReduce: Simplified data processing on large clusters," *Communications of the ACM*, vol. 51, no. 1, pp. 107–113, 2008.
- [17] F. Pedregosa et al., "Scikit-learn: Machine learning in Python," *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.
- [18] G. Hamerly and J. Drake, "Accelerating Lloyd's algorithm for k-means clustering," *Partitional Clustering Algorithms*, Springer, pp. 41–78, 2015.
- [19] A. Singh, A. Yadav, and A. Rana, "K-means with three different distance metrics," *Int. Journal of Computer Applications*, vol. 67, no. 10, pp. 13–17, 2013.
- [20] C. Aggarwal, *Data Clustering: Algorithms and Applications*, Chapman & Hall/CRC, 2013.
- [21] S. Theodoridis and K. Koutroumbas, *Pattern Recognition*, 4th ed., Academic Press, 2009.
- [22] R. Xu and D. Wunsch, "Survey of clustering algorithms," *IEEE Transactions on Neural Networks*, vol. 16, no. 3, pp. 645–678, May 2005.
- [23] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, pp. 436–444, 2015.
- [24] A. Likas, N. Vlassis, and J. Verbeek, "The global k-means clustering algorithm," *Pattern Recognition*, vol. 36, no. 2, pp. 451–461, 2003.
- [25] S. Raschka and V. Mirjalili, *Python Machine Learning*, 3rd ed., Packt Publishing, 2019.