

Original Research

K-Means Clustering for Drug Inventory Analysis at Anugrah Pharmacy Bekasi

Alya Priscilla Putri ^a, Adi Muhajirin ^{b*}, Fata Nidaul Khasanah ^c

^{a,b,c} Informatics, Faculty of Computer Science, Universitas Bhayangkara Jakarta Raya, Bekasi, Indonesia

*Corresponding Author: adi.muhammad@dsn.ubharajaya.ac.id

ABSTRACT

Improper drug inventory management can cause stockouts, overstocking, and inefficient procurement decisions, especially in small-scale retail pharmacies that still rely on manual estimation. This study analyzes drug sales patterns and groups drug inventory at Anugrah Pharmacy Bekasi using the K-Means clustering algorithm within the Cross Industry Standard Process for Data Mining (CRISP-DM) framework. The dataset consisted of 5,312 sales transactions from January to June 2025. The transaction records were aggregated into 731 drug items using three variables: transaction frequency, sales volume, and transaction value. Data preparation included aggregation, missing-value checking, duplicate checking, transformation, and Min-Max normalization. The optimal number of clusters was determined using the Elbow Method, which indicated three clusters ($k = 3$). The K-Means results grouped the 731 drug items into 28 Fast Moving items (3.83%), 129 Medium Moving items (17.65%), and 574 Slow Moving items (78.52%). The centroid analysis shows that each cluster has distinct sales-movement characteristics. The results can support inventory decision-making by helping the pharmacy prioritize replenishment for fast-moving drugs, maintain controlled stock for medium-moving drugs, and limit excessive procurement for slow-moving drugs. This study demonstrates that CRISP-DM and K-Means clustering can provide practical information for data-driven drug inventory management in a retail pharmacy context.

Copyright © 2026 by Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0).

How to Cite:

A. P. Putri, A. Muhajirin, and F. N. Khasanah, "K-Means Clustering for Drug Inventory Analysis at Anugrah Pharmacy Bekasi," *J. Kecerdasan Buatan dan Teknol. Inf.*, vol. 5, no. 3, pp. 623–631, 2026.

Doi: <https://doi.org/10.69916/jkbt1.v5i3.577>

KEYWORDS

Data Mining
CRISP-DM
K-Means Clustering
Drug Inventory
Sales Pattern

ARTICLE INFO

Received: July 7, 2026
Revised: August 20, 2026
Available Online: September 1, 2026

1. INTRODUCTION

Effective drug inventory management is a critical operational imperative in the pharmaceutical supply chain, directly dictating service continuity, patient safety, and financial viability. In retail pharmacies, inventory control is inherently complex due to the perishability of products, strict regulatory requirements, and fluctuating consumer demand. Suboptimal inventory practices frequently lead to a dichotomous operational failure: stockouts of high-demand essential medicines, which compromise patient care, and the overstocking of low-demand items, which results in capital tie-up and expiry-related financial losses [1–3]. While large hospital networks and public health institutions have increasingly adopted Enterprise Resource Planning (ERP) systems and automated forecasting models, small and medium-sized enterprise (SME) retail pharmacies often remain resource-constrained. Consequently, their procurement and stock-control decisions are predominantly driven by manual estimation, intuitive experience, and rudimentary spreadsheet tracking, leaving them highly vulnerable to market inefficiencies [4, 5].

Anugrah Pharmacy Bekasi exemplifies this prevalent challenge among SME retail pharmacies. Despite possess-

ing a wealth of historical operational data, recording 5,312 distinct drug sales transactions between January and June 2025, the pharmacy suffers from a "data-rich but information-poor" paradox. The raw transaction logs contain latent patterns regarding drug movement and financial contribution; however, this data has not been systematically transformed into actionable knowledge. Without a structured mechanism to classify inventory into distinct movement categories (e.g., fast, medium, and slow-moving), the pharmacy struggles to optimize safety stock levels, prioritize procurement budgets, and mitigate the risk of drug expiration [6, 7].

The advent of data science and data-driven decision-making offers a robust pathway to bridge this operational gap. By applying unsupervised machine learning techniques, organizations can extract hidden patterns from transactional data to support objective, evidence-based inventory stratification [8, 9]. Among clustering algorithms, K-Means remains the most widely adopted for inventory segmentation due to its computational efficiency, scalability, and effectiveness in handling multi-dimensional numerical datasets [10]. To ensure the data mining process aligns with business objectives, the Cross-Industry Standard Process for Data Mining (CRISP-DM) is utilized. Unlike older, rigid frameworks such as Knowledge Discovery in Databases (KDD), CRISP-DM provides an iterative, business-centric workflow that ensures analytical outputs directly address operational pain points [11].

Despite the proven efficacy of K-Means in healthcare inventory management, a critical review of contemporary literature reveals four significant research gaps. First, contextual bias: the majority of existing studies focus on large-scale hospital pharmacies or public health centers [12–14], largely neglecting the unique, resource-constrained environment of SME retail pharmacies. Second, dimensional limitations: many studies rely on single-dimensional metrics, such as static stock levels or basic sales volume [15, 16], failing to capture the holistic relationship between demand frequency, physical stock movement, and financial revenue contribution. Third, methodological rigidity: several prominent studies continue to utilize the outdated KDD framework rather than the agile, business-aligned CRISP-DM methodology [17, 18]. Fourth, and most crucially, the "last-mile" deployment gap: the vast majority of academic data mining studies terminate at the theoretical evaluation phase. They produce mathematical clusters but fail to operationalize these findings into accessible, user-friendly decision support systems (DSS) for non-technical pharmacy staff [19, 20].

To address these lacunae, this study proposes an end-to-end, data-driven inventory analytics framework tailored specifically for SME retail pharmacies. By applying the CRISP-DM methodology and K-Means clustering to real-world historical sales data at Anugrah Pharmacy Bekasi, this research transitions from theoretical modeling to practical operationalization. Furthermore, to bridge the deployment gap, the resulting clustering model is integrated into a web-based dashboard using Streamlit and GitHub, enabling pharmacy managers to interactively upload new sales data and generate instant inventory stratification reports without requiring programming expertise [21]. The specific novelties and contributions of this study are threefold:

1. **Multi-Dimensional Inventory Stratification:** Unlike previous studies that rely on limited variables, this study clusters 731 drug items using a tri-dimensional feature set (Transaction Frequency, Sales Volume, and Transaction Value), providing a comprehensive view of both physical movement and financial impact.
2. **Contextual Application of CRISP-DM in SMEs:** This research explicitly demonstrates how the enterprise-grade CRISP-DM framework can be adapted to structure data mining workflows in resource-limited retail pharmacy settings.
3. **Actionable Deployment via Web-Based DSS:** Moving beyond theoretical clustering, this study delivers a deployed, browser-based application that translates complex K-Means outputs into immediate, actionable procurement strategies (Fast, Medium, and Slow-Moving categories) for everyday pharmacy operations.

Table 1. Summary of Related Studies

Study	Object/Data	Method	Main Finding and Relevance
Nugroho et al. [22]	Hospital drug data	K-Means and KDD	Grouped drugs into high- and low-use clusters; not focused on retail pharmacy SMEs.
Hidayati et al. [23]	Public health-center drug data	K-Means	Formed high, medium, and low drug-use clusters; did not use CRISP-DM.
Kunaifi Kurnia et al. [24]	Pharmacy drug sales data	K-Means and KDD	Produced two sales clusters; different framework from this study.
Prasetyo et al. [25]	Clinic drug inventory	K-Means and Elbow Method	Produced fast, medium, and slow clusters using different variables and object.
Hidayah et al. [26]	Hospital drug data	K-Means in Rapid-Miner	Grouped hospital drugs by characteristics; not a small retail pharmacy context.
Supriyanto et al. [15]	Pharmaceutical inventory data	K-Means, Elbow, Gap Statistic	Generated three optimal clusters, but used stock-oriented data.
Fitriyani et al. [12]	Pharmacy inventory data	K-Means and KDD	Clustered pharmacy inventory; did not adopt CRISP-DM.
Wahyuni et al. [27]	Pharmacy sales data	K-Means	Determined fast, medium, and slow moving products; variables were more limited than this study.

2. RESEARCH METHODS

2.1. Research Framework

This study uses the CRISP-DM framework to guide the data mining process. The business-understanding phase identified the main problem: Anugrah Pharmacy Bekasi had not grouped drug items based on sales movement, making procurement and stock-control decisions less data-driven. The data-understanding phase explored transaction records from January to June 2025. The data-preparation phase cleaned, transformed, aggregated, and normalized the data. The modeling phase applied K-Means clustering, the evaluation phase interpreted the resulting centroids and cluster distribution, and the deployment phase converted the findings into inventory-management recommendations.

2.2. Dataset and Variables

The dataset consisted of 5,312 sales transactions obtained from Anugrah Pharmacy Bekasi. The transaction attributes included transaction date, transaction identifier, drug code, drug name, quantity sold, unit, unit price, total price, and drug category. For clustering, the transaction-level data were aggregated by drug name into 731 unique drug items. Three numerical variables were used as clustering inputs: Transaction Frequency, Sales Volume, and Transaction Value. These variables were selected because they represent demand intensity, quantity movement, and financial contribution [28,29].

Table 2. Clustering Variables

Variable	Description	Role in Inventory Analysis
Transaction Frequency	Number of sales transactions involving a drug item	Identifies how often a drug is purchased.
Sales Volume	Total number of drug units sold	Identifies the magnitude of physical stock movement.
Transaction Value	Total sales value generated by a drug item	Identifies the financial contribution of each drug item.

2.3. Data Preparation

Data preparation began with aggregation by drug item. The transaction frequency was calculated by counting transaction occurrences, sales volume was calculated by summing quantities sold, and transaction value was calculated by summing total sales value for each drug. After aggregation, missing values and duplicate records were checked. The final aggregated dataset contained 731 drug items with no missing or duplicate records.

The numerical variables were normalized using Min-Max Scaling so that all features were placed within the 0 to 1 range:

$$X' = \frac{X - X_{\min}}{X_{\max} - X_{\min}} \quad (1)$$

where X' is the normalized value, X is the original value, $X_{\min}$ is the minimum value, and $X_{\max}$ is the maximum value.

2.4. K-Means Modelling and Cluster Selection

K-Means clustering groups data into k clusters by minimizing the distance between each data point and the centroid of its assigned cluster. The distance was calculated using Euclidean Distance:

$$d(x, c) = \sqrt{\sum_{i=1}^n (x_i - c_i)^2} \quad (2)$$

The optimal number of clusters was determined using the Elbow Method, which compares Within-Cluster Sum of Squares (WCSS) values across different values of k :

$$\text{WCSS} = \sum_{k=1}^K \sum_{x \in C_k} \|x - \mu_k\|^2 \quad (3)$$

The experiment tested values of k from 1 to 10. The curve showed a major decrease in WCSS from $k = 1$ to $k = 3$, followed by a more gradual decrease after $k = 3$. Therefore, three clusters were selected for the final model.

2.5. Evaluation and Deployment

Evaluation was conducted by analyzing the cluster distribution and centroid values of each cluster. The final deployment output was integrated into a web application using Streamlit and GitHub to support procurement and stock-control decisions.

3. RESULTS AND DISCUSSION

3.1. Data Aggregation and Cleaning Results

The initial dataset contained 5,312 sales transaction records for the January to June 2025 period. After aggregation by drug item, the dataset contained 731 unique drug items. Each item had values for transaction frequency, sales volume, and transaction value. The cleaning process showed that all three main variables had no missing values or duplicate records. Examples from the aggregated data show differences in drug movement: Amoxicillin 500 mg appeared in 64 transactions with 93 units sold (Rp 513,000), Paratusin tablets appeared in 58 transactions with 74 units sold (Rp 1,238,000), whereas Dulcolax 10S had only 3 transactions and 4 units sold.

3.2. Elbow Method and K-Means Cluster Formation

The Elbow Method identified $k = 3$ as the optimal number of clusters. The final K-Means model grouped the 731 drug items into three clusters labeled Fast Moving, Medium Moving, and Slow Moving.

Table 3. K-Means Cluster Distribution

Cluster Category	Number of Drug Items	Percentage
Fast Moving	28	3.83%
Medium Moving	129	17.65%
Slow Moving	574	78.52%
Total	731	100.00%

Table 3 illustrates the distribution of the 731 drug items across the three generated clusters. The K-Means algorithm revealed a highly skewed, long-tail distribution characteristic of retail inventory. The Fast Moving cluster comprises a mere 28 items (3.83% of the total inventory). Despite their small proportion, these items represent the core drivers of daily pharmacy transactions. The Medium Moving cluster contains 129 items (17.65%), representing stable products. Most notably, the Slow Moving cluster overwhelmingly dominates the inventory landscape, accounting for 574 items (78.52%).

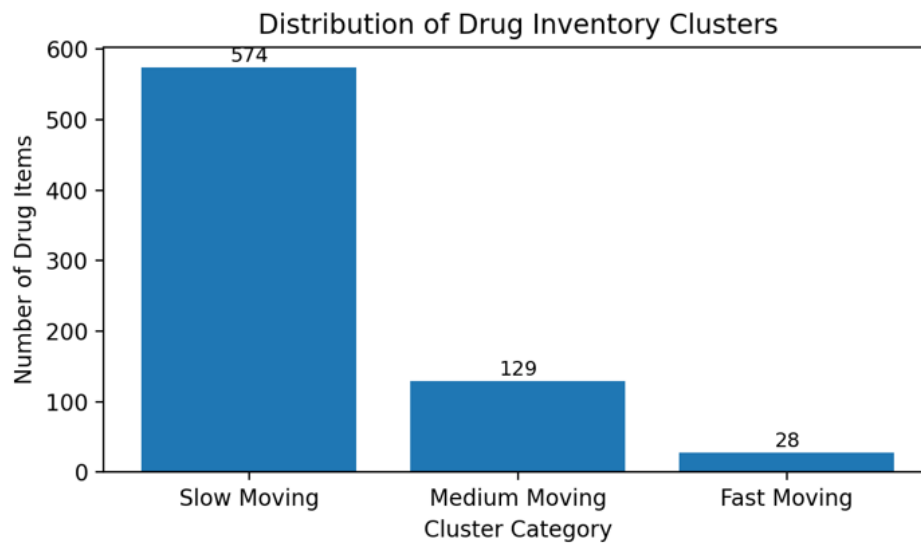


Figure 1. Distribution of drug inventory clusters based on K-Means results

Figure 1 presents the distribution of the 731 drug items across the three clusters produced by the K-Means model. The bar chart reveals a highly skewed, long-tail structure in the pharmacy's inventory: the Slow Moving cluster dominates with 574 items (78.52%), the Medium Moving cluster contains 129 items (17.65%), and the Fast Moving cluster comprises only 28 items (3.83%). This steep disparity indicates that only a small fraction of the catalog drives the majority of sales activity, while most items exhibit limited movement over the six-month observation period. Overall, Figure 1 confirms that the clustering output yields clearly separable movement categories that can serve as an empirical basis for differentiated inventory control policies at Anugrah Pharmacy Bekasi.

3.3. Cluster Characteristics

The centroid analysis confirms that the clusters have distinct characteristics, as shown in Table 4.

Table 4. Normalized Centroid Characteristics by Cluster

Cluster Category	Transaction Frequency	Sales Volume	Transaction Value
Fast Moving	0.5186	0.5955	0.4007
Medium Moving	0.1534	0.1422	0.2048
Slow Moving	0.0234	0.0221	0.0355

3.4. Inventory Management Implications

The clustering results can be translated into practical inventory-management actions summarized in Table 5.

Table 5. Inventory Recommendations Based on Cluster Category

Cluster Category	Observed Characteristic	Recommended Action
Fast Moving	High transaction frequency, high volume, and high transaction value.	Increase monitoring frequency, prioritize replenishment, and maintain safety stock.
Medium Moving	Moderate transaction frequency and relatively stable sales movement.	Use periodic procurement and review demand trends regularly.
Slow Moving	Low transaction frequency, low volume, and low transaction value.	Limit procurement quantity, evaluate product display, and prevent stock accumulation.

3.5. Deployment: Web-Based Application Implementation

Deployment is the final stage in the CRISP-DM methodology. In this study, the data mining results were integrated into a web-based Streamlit application deployed via GitHub.

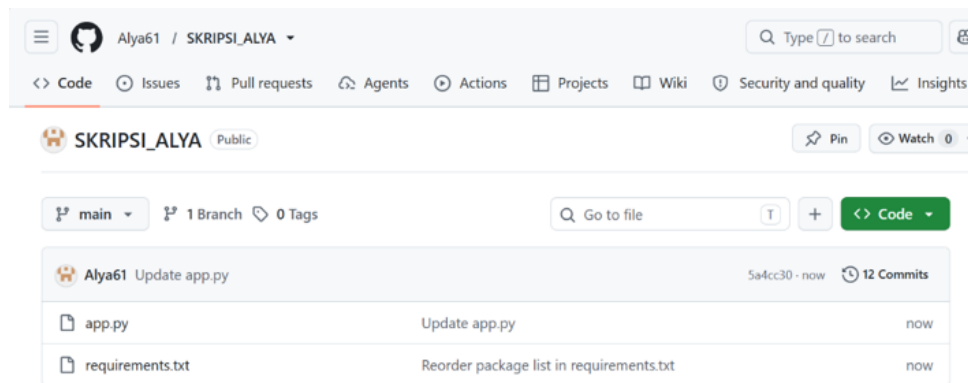


Figure 2. Deployment Process

Figure 2 shows the deployment process using GitHub and Streamlit. GitHub was used as a repository to store and manage the program code, while Streamlit was used to present the implementation results as a web application accessible via browser.

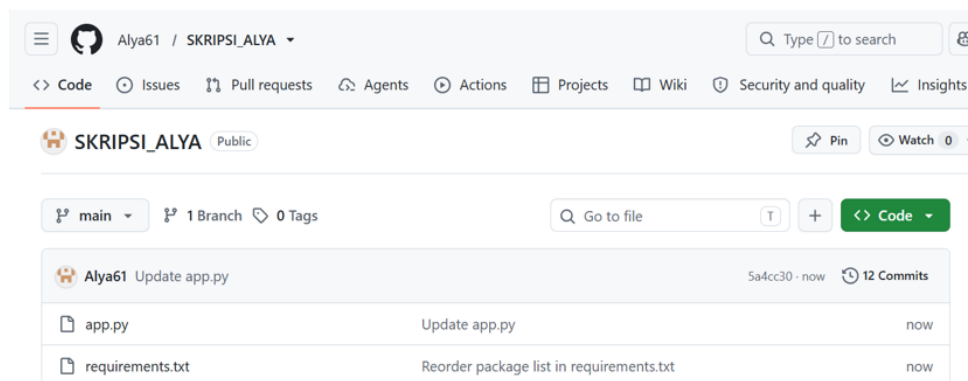


Figure 3. Deployment Result Page

Figure 3 shows the deployment result page of the Drug Inventory Clustering application. The page provides an Upload Sales Data feature that allows users to upload sales datasets in Excel (.xlsx) or CSV format. After uploading, the system processes the data using the trained K-Means model, displaying the cluster assignments and generating downloadable Excel reports.

3.6. Comparison with Previous Studies

The findings align with earlier K-Means research demonstrating that drug inventory can be segmented into distinct movement categories [10–17]. However, this study uniquely focuses on SME retail pharmacies, adopts the CRISP-DM framework, utilizes a tri-dimensional feature space, and bridges the deployment gap via an operational Streamlit web application.

3.7. Discussion

The results provide empirical evidence that K-Means clustering within CRISP-DM effectively transforms raw retail transactions into actionable inventory intelligence. The long-tail distribution (78.52% Slow Moving, 3.83% Fast Moving) highlights significant capital tie-up and expiration risk in SME operations. The tri-dimensional feature set prevented misclassification of high-value, moderate-turnover drugs. Practical deployment via Streamlit empowers non-technical pharmacy personnel to execute data-driven procurement strategies seamlessly.

4. CONCLUSION

This study applied the CRISP-DM framework and K-Means clustering algorithm to analyze drug sales patterns at Anugrah Pharmacy Bekasi. The dataset consisted of 5,312 transactions from January to June 2025, which were aggregated into 731 drug items using transaction frequency, sales volume, and transaction value. After data cleaning, transformation, and Min-Max normalization, the Elbow Method indicated that three clusters were optimal. The final K-Means model grouped the drug items into 28 Fast Moving items, 129 Medium Moving items, and 574 Slow Moving items. The centroid analysis showed that each cluster had distinct characteristics, with Fast Moving drugs having the highest sales movement and Slow Moving drugs having the lowest movement. The study contributes a practical data-driven approach for supporting pharmacy inventory management. The clustering results can help the pharmacy prioritize replenishment for fast-moving drugs, maintain controlled stock for medium-moving drugs, and reduce excessive procurement of slow-moving drugs. Future studies can extend this work by using a longer sales period, adding variables such as expiry date, purchase price, supplier lead time, and stock level, and evaluating clustering quality with additional metrics such as Silhouette Score or Davies-Bouldin Index.

DECLARATIONS

Author Contributions:

Alya Priscilla Putri: Conceptualization, methodology, data curation, software, formal analysis, visualization, and writing – original draft.

Adi Muhajirin: Supervision, methodology, and writing – review and editing.

Fata Nidaul Khasanah: Supervision, validation, and writing – review and editing.

Funding:

The authors received no financial support for the research, authorship, and/or publication of this article.

Conflicts of Interest:

The authors declare that they have no competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data Availability:

The data that support the findings of this study are available from the corresponding author upon reasonable request.

AI Use Statement:

During the preparation of this work, the authors used Qwen AI in order to improve the language and readability of the manuscript. After using this tool, the authors reviewed and edited the content as needed and take full responsibility for the final content of the published article.

Additional Information:

No additional information is available for this paper.

REFERENCES

- [1] D. A. Ramadhanty, R. Syafitri, E. Rasywir, and D. Meisak, "Implementasi Data Mining Untuk Menentukan Persediaan Stok Obat Di Apotek K-24 Menggunakan Metode K-Means Clustering," *J. Inform. Dan Rekayasa Komput.*, vol. 2, no. 1, pp. 155–160, 2022, doi: 10.33998/jakakom.2022.2.1.31.
- [2] F. D. Ragestu and A. J. P. Sibarani, "Penerapan Metode Fuzzy Tsukamoto Dalam Pemilihan Siswa Teladan di Sekolah," *Teknika*, vol. 9, no. 1, pp. 9–15, 2020, doi: 10.34148/teknika.v9i1.251.

- [3] A. J. P. Sibarani, “Persediaan Obat Di Apotek Bumi Husada,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 4, pp. 4430–4435, 2024.
- [4] S. V. G. Subrahmanya et al., “The Role of Data Science in Healthcare Advancements: Applications, Benefits, and Future Prospects,” *Ir. J. Med. Sci.*, vol. 191, no. 4, pp. 1473–1483, 2022, doi: 10.1007/s11845-021-02730-z.
- [5] N. Elgendy, A. Elragal, and T. Päivärinta, “DECAS: A Modern Data-Driven Decision Theory for Big Data and Analytics,” *J. Decis. Syst.*, vol. 31, no. 4, pp. 337–373, 2022, doi: 10.1080/12460125.2021.1894674.
- [6] Y. Yesica, T. Sitorus, and E. Purwanto, “Pengaruh Tata Kelola Perusahaan yang Baik dan Tanggung Jawab Sosial Perusahaan terhadap Kinerja Keuangan,” *J. Bus. Appl. Manag.*, vol. 13, no. 2, p. 191, 2020, doi: 10.30813/jbam.v13i2.2356.
- [7] T. Ardiansah and D. Hidayatullah, “Penerapan Metode Waterfall Pada Aplikasi Reservasi Lapangan Futsal Berbasis Web,” *J. Inf. Technol. Softw. Eng. Comput. Sci.*, vol. 1, no. 1, pp. 6–13, 2022, doi: 10.58602/it-secs.v1i1.8.
- [8] B. Chong, “K-Means Clustering Algorithm: A Brief Review,” *Acad. J. Comput. Inf. Sci.*, vol. 4, no. 5, pp. 37–40, 2021, doi: 10.25236/ajcis.2021.040506.
- [9] I. P. Mulyadi, “Klasterisasi Menggunakan Metode Algoritma K-Means dalam Meningkatkan Penjualan Tupperware,” *J. Inform. Ekon. Bisnis*, vol. 4, pp. 172–179, 2022, doi: 10.37034/infeb.v4i4.164.
- [10] M. Elkalawy, A. Al-Sakkaf, E. M. Abdelkader, and G. Alfalah, “CRISP-DM-Based Data-Driven Approach for Building Energy Prediction Utilizing Indoor and Environmental Factors,” *Sustainability*, vol. 16, no. 17, 2024, doi: 10.3390/su16177249.
- [11] N. Lebkiri et al., “Using Machine Learning for Prediction Students Failure in Morocco: An Application of the CRISP-DM Methodology,” *Int. J. Educ. Inf. Technol.*, vol. 15, pp. 344–352, 2021, doi: 10.46300/9109.2021.15.36.
- [12] D. Fitriyani, M. Jajuli, and G. Garo, “Implementasi Algoritma K-Means Untuk Klasterisasi Dalam Pengelolaan Persediaan Obat (Studi Kasus: Apotek Naza),” *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, pp. 2841–2848, 2024, doi: 10.23960/jitet.v12i3.4921.
- [13] B. Shen, “E-commerce Customer Segmentation via Unsupervised Machine Learning,” in *CONFCDs 2021*, 2021.
- [14] D. Juniati and A. E. Suwanda, “Klasifikasi Penyakit Mata Berdasarkan Citra Fundus Retina Menggunakan Dimensi Fraktal Box Counting Dan Fuzzy K-Means,” *Prox. J. Penelit. Mat. dan Pendidik. Mat.*, vol. 5, no. 1, pp. 10–18, 2022, doi: 10.30605/proximal.v5i1.1623.
- [15] H. Supriyanto, M. Al Hafidz, A. C. Puspitaningrum, R. A. P. Firmansyah, and R. Zuhdi, “Klasterisasi Data Obat Farmasi Berdasarkan Jumlah Persediaan Dengan Menggunakan Metode K-Means,” *Teknika*, vol. 13, no. 3, pp. 361–369, 2024, doi: 10.34148/teknika.v13i3.987.
- [16] L. A. Prasetyo, I. Himawan, and R. A. Sihombing, “Pengelompokan Persediaan Obat Dengan Metode K-Means,” *J. Rekayasa Sist. Inf. dan Teknol.*, vol. 2, no. 2, pp. 759–768, 2024.
- [17] H. Supriyanto, M. Al Hafidz, A. C. Puspitaningrum, R. A. P. Firmansyah, and R. Zuhdi, “Klasterisasi Data Obat Farmasi Berdasarkan Jumlah Persediaan Dengan Menggunakan Metode K-Means,” *Teknika*, vol. 13, no. 3, pp. 361–369, 2024, doi: 10.34148/teknika.v13i3.987.
- [18] E. Wahyuni, F. Yunita, B. Rianto, U. I. Indragiri, I. Hilir, and K. Clustering, “Implementasi Metode K-Means Clustering Dalam Menentukan Produk Fast, Medium dan Slow Moving Pada Apotek Zerina,” *Unknown*, vol. 8, pp. 1–10, 2026.
- [19] W. E. Sari, M. B, and S. Rani, “Perbandingan Metode SAW dan Topsis pada Sistem Pendukung Keputusan Seleksi Penerima Beasiswa,” *J. Sisfokom (Sistem Inf. dan Komputer)*, vol. 10, no. 1, pp. 52–58, 2021, doi: 10.32736/sisfokom.v10i1.1027.

- [20] F. N. Handayani, I. Diasih, V. A. Salsabilla, and A. Pramudita, “Sistem Pendukung Keputusan Dalam Pemilihan Smartphone Terbaik Menggunakan Metode SAW,” *Bridge J. Publ. Sist. Inf. dan Telekomun.*, vol. 2, no. 3, pp. 130–141, 2024, doi: 10.62951/bridge.v2i3.127.
- [21] J. N. Oruh and B. C. Iwuji, “A Hybrid Prediction Model for Classifying StudentS Academic Performance Using Voting Ensemble Method,” *Fudma J. Sci.*, vol. 10, no. 3, pp. 364–376, 2026, doi: 10.33003/fjs-2026-1003-4751.
- [22] M. R. Nugroho, I. E. Hendrawan, T. Informatika, U. Singa, P. Karawang, and D. Obat, “Penerapan Algoritma K-Means Untuk Klasterisasi Data Obat Pada Rumah Sakit,” vol. 16, pp. 125–133, 2022.
- [23] P. Aliya, S. Nurbaeti, A. Firdaus, F. Suparman, and M. Pd, “Title,” vol. 11, no. April, pp. 156–165, 2025.
- [24] K. Kurnia Abdullah, I. Maulana, A. Suharso, G. Garino, and A. Primajaya, “Penerapan Algoritme K-Means Dalam Klasterisasi Data Penjualan Obat Apotek Kidangrangga,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 2, pp. 1274–1279, 2023, doi: 10.36040/jati.v7i2.7061.
- [25] L. A. Prasetyo, I. Himawan, and R. A. Sihombing, “Pengelompokan Persediaan Obat Dengan Metode K-Means,” *J. Rekayasa Sist. Inf. dan Teknol.*, vol. 2, no. 2, pp. 759–768, 2024.
- [26] A. H. Siregar, D. D. Sihotang, B. A. Wijaya, and S. D. Siregar, “Implementasi Algoritma K - Means Menggunakan RapidMiner untuk Klasterisasi Data Obat,” *J. Teknol. dan Ilmu Komput. Prima*, vol. 7, no. 2, pp. 200–211, 2024.
- [27] E. Wahyuni, F. Yunita, B. Rianto, U. I. Indragiri, I. Hilir, and K. Clustering, “IMPLEMENTASI METODE K-MEANS CLUSTERING DALAM MENENTUKAN PRODUK FAST, MEDIUM DAN SLOW MOVING PADA APOTEK ZERINA,” vol. 8, pp. 1–10, 2026.
- [28] R. Maoulana, B. Irawan, and A. Bahtiar, “Data Mining Dalam Konteks Transaksi Penjualan Hijab Dengan Menggunakan Algoritma Clustering K-Means,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 515–521, 2024, doi: 10.36040/jati.v8i1.8504.
- [29] S. A. Arnomo, N. Adhiatma, and H. Nuryanto, “Data Mining Analysis for Inventory Management,” *J. Data Sci. Technol.*, vol. 4, no. 1, pp. 54–59, 2025.