



E-ISSN 3032-601X & P-ISSN 3032-7105

Vol. 3, No. 1, 2026

MISTER

**Journal of Multidisciplinary Inquiry in Science,
Technology and Educational Research**

**Jurnal Penelitian Multidisiplin dalam Ilmu
Pengetahuan, Teknologi dan Pendidikan**

**UNIVERSITAS SERAMBI MEKKAH
KOTA BANDA ACEH**

mister@serambimekkah.ac.id

Journal of Multidisciplinary Inquiry in Science Technology
and Educational Research

Journal of MISTER

Vol. 3, No. 1, 2026

Pages: 1108–1118

Clustering Pelanggan dengan K-Means Menggunakan Parameter Frekuensi Belanja dan Total Pengeluaran

Rahman Budianto, Hasbi Firmansyah, Rizki Prasetyo Tulodo,
Wahyu Asriyani

Universitas Pancasakti Tegal

Article in Journal of MISTER

Available at : <https://jurnal.serambimekkah.ac.id/index.php/mister/index>

DOI : <https://doi.org/10.32672/mister.v3i1.4038>

How to Cite this Article

APA : Budianto, R., Hasbi Firmansyah, Rizki Prasetyo Tulodo, & Wahyu Asriyani. (2025). Clustering Pelanggan dengan K-Means Menggunakan Parameter Frekuensi Belanja dan Total Pengeluaran. *Journal of Multidisciplinary Inquiry in Science, Technology and Educational Research*, 3(1), 1108 – 1118. <https://doi.org/10.32672/mister.v3i1.4038>

Others Visit : <https://jurnal.serambimekkah.ac.id/index.php/mister/index>

MISTER: *Journal of Multidisciplinary Inquiry in Science, Technology and Educational Research* is a scholarly journal dedicated to the exploration and dissemination of innovative ideas, trends and research on the various topics include, but not limited to functional areas of Science, Technology, Education, Humanities, Economy, Art, Health and Medicine, Environment and Sustainability or Law and Ethics.

MISTER: *Journal of Multidisciplinary Inquiry in Science, Technology and Educational Research* is an open-access journal, and users are permitted to read, download, copy, search, or link to the full text of articles or use them for other lawful purposes. Articles on Journal of MISTER have been previewed and authenticated by the Authors before sending for publication. The Journal, Chief Editor, and the editorial board are not entitled or liable to either justify or responsible for inaccurate and misleading data if any. It is the sole responsibility of the Author concerned.



ISSN 3032-7105

ISSN 3032-601X



Clustering Pelanggan dengan K-Means Menggunakan Parameter Frekuensi Belanja dan Total Pengeluaran

Rahman Budianto¹, Hasbi Firmansyah², Rizki Prasetyo Tulodo³, Wahyu Asriyani⁴
Program Studi Informatika, Fakultas Teknik dan Ilmu Komputer, Universitas Pancasakti Tegal, Tegal, Indonesia^{1,2,4}
Program Studi Pendidikan Bahasa dan Sastra Indonesia, Fakultas Keguruan dan Ilmu Pendidikan, Universitas Pancasakti Tegal, Tegal, Indonesia³

*Email: rahmanbudianto181@mail.com, hasbifirmansyah@upstegal.ac.id,
rizki.prasetyo.tulodo@gmail.com, asriyani1409@gmail.com

Diterima: 13-12-2025

| Disetujui: 23-12-2025

| Diterbitkan: 25-12-2025

ABSTRACT

Customer segmentation is an important strategy for understanding consumer behavior patterns. This study aims to group customers using the K-Means algorithm based on purchase frequency and total expenditure. The research process was conducted using RapidMiner as a data mining tool for data processing and analysis. The research stages include dataset collection, data preprocessing, implementation of the K-Means algorithm, and evaluation of clustering quality using the Performance (Clustering) method. The number of clusters used in this study is two clusters ($k = 2$). The results show that the K-Means algorithm is able to form customer groups with different characteristics. Model evaluation produces a positive Silhouette Coefficient and a relatively low Davies-Bouldin Index, indicating good internal cluster cohesion and clear separation between clusters. Therefore, the application of the K-Means method using RapidMiner can be utilized as a solution to support customer segmentation and assist decision-making in marketing strategies.

Keywords: K-Means; Clustering; RapidMiner; Customer Segmentation; Data Mining.

ABSTRAK

Segmentasi pelanggan merupakan salah satu strategi penting dalam memahami pola perilaku konsumen. Penelitian ini bertujuan untuk melakukan pengelompokan pelanggan menggunakan algoritma K-Means berdasarkan frekuensi belanja dan total pengeluaran. Proses penelitian dilakukan dengan memanfaatkan aplikasi RapidMiner sebagai alat bantu pengolahan dan analisis data. Tahapan penelitian meliputi pengumpulan dataset pelanggan, preprocessing data, penerapan algoritma K-Means, serta evaluasi kualitas clustering menggunakan metode Performance (Clustering). Jumlah cluster yang digunakan pada penelitian ini adalah dua cluster ($k = 2$). Hasil penelitian menunjukkan bahwa algoritma K-Means mampu membentuk kelompok pelanggan dengan karakteristik yang berbeda. Evaluasi model menghasilkan nilai Silhouette Coefficient bernilai positif dan Davies-Bouldin Index yang relatif rendah, yang mengindikasikan bahwa cluster yang terbentuk memiliki tingkat kemiripan internal yang baik serta pemisahan antar cluster yang cukup jelas.

Katakunci: K-Means; Clustering; RapidMiner; Segmentasi Pelanggan; Data Mining.

PENDAHULUAN

Perkembangan dunia bisnis yang semakin kompetitif menuntut setiap perusahaan atau pelaku usaha untuk memahami perilaku pelanggan secara lebih mendalam (Atmoko et al., 2025). Informasi mengenai kebiasaan belanja pelanggan, seperti frekuensi kunjungan dan total pengeluaran, menjadi sangat penting dalam menentukan strategi pemasaran yang tepat (Hidayat et al., 2025). Namun, data tersebut sering kali berjumlah besar dan sulit dianalisis secara manual tanpa adanya metode pengolahan data yang sistematis (Elektro, 2024).

Salah satu teknik data mining yang banyak digunakan untuk menganalisis pola dan mengelompokkan data adalah algoritma K-Means Clustering (Arifianto et al., 2024). K-Means merupakan metode unsupervised learning yang bertujuan untuk mengelompokkan data ke dalam beberapa cluster berdasarkan tingkat kemiripan antar objek (Safitri et al., 2025). Dengan membagi pelanggan ke dalam kelompok tertentu, perusahaan dapat mengetahui karakteristik setiap segmen dan merancang strategi pelayanan maupun promosi yang lebih efektif (Malang, 2019).

Pada penelitian ini, algoritma K-Means diterapkan untuk mengelompokkan pelanggan berdasarkan frekuensi belanja dan total pengeluaran (Pengeluaran et al., 2023). Kedua variabel tersebut dipilih karena mampu mencerminkan tingkat loyalitas pelanggan serta nilai transaksi yang dihasilkan (Kristanti et al., 2024). Hasil clustering diharapkan dapat membantu pengelola usaha dalam mengidentifikasi pelanggan prioritas, pelanggan dengan potensi peningkatan transaksi, serta kelompok pelanggan yang membutuhkan pendekatan pemasaran yang berbeda (Jihan et al., 2025).

Dengan demikian, penelitian ini tidak hanya memberikan gambaran mengenai pola perilaku belanja pelanggan, tetapi juga berkontribusi dalam memberikan rekomendasi strategi bisnis yang lebih tepat sasaran. Melalui analisis clustering, pelaku usaha dapat mengoptimalkan pengambilan keputusan berbasis data dan meningkatkan efisiensi dalam pengelolaan hubungan pelanggan (Rahma et al., 2025).

METODE PENELITIAN

Penelitian ini merupakan penelitian kuantitatif deskriptif dengan pendekatan data mining, khususnya menggunakan metode K-Means Clustering untuk mengelompokkan pelanggan berdasarkan dua variabel numerik, yaitu frekuensi belanja dan total pengeluaran (Sari et al., 2020). Tujuan penelitian ini adalah untuk mengidentifikasi pola kelompok pelanggan yang memiliki karakteristik serupa sehingga dapat memberikan dasar bagi pengambilan keputusan bisnis (Dona & Rifqi, 2022).



Gambar 1. Tahap Penelitian

Penelitian ini dilakukan melalui beberapa tahapan, yang diilustrasikan pada diagram alur berikut:

1. Pengumpulan Dataset

Dataset yang digunakan dalam penelitian ini merupakan data pelanggan yang diperoleh dari catatan transaksi toko/usaha dalam periode tertentu. Data berisi dua atribut utama, yaitu Frekuensi Belanja (jumlah kunjungan pelanggan dalam periode waktu tertentu) dan Total Pengeluaran (jumlah uang yang dibelanjakan pelanggan) (Ramadhan & Saprudin, 2022). Dataset disusun dalam format tabel untuk memudahkan proses pengolahan pada Microsoft Excel dan RapidMiner Studio (Faidah & Fatah, 2025).

2. Preprocessing

Sebelum dilakukan proses clustering, data harus melalui tahap preprocessing untuk memastikan kualitas data yang optimal. Preprocessing mencakup:

1. Pemeriksaan Missing Value

Memastikan seluruh data terisi lengkap. Jika terdapat data kosong, maka dilakukan:

- penghapusan baris (remove), atau
- imputasi nilai, jika memungkinkan.

2. Pemeriksaan Outlier

Outlier yang terlalu ekstrem dapat memengaruhi hasil clustering sehingga perlu dicek dan ditangani sesuai kebutuhan.

3. Pemilihan Atribut (Attribute Selection)

Atribut yang digunakan hanya Frekuensi dan Pengeluaran, sesuai dengan tujuan clustering.

4. Normalisasi Data (Opsional)

Jika rentang nilai antar atribut sangat berbeda, normalisasi dilakukan dengan metode:

$$Z = \frac{x - \mu}{\sigma}$$

atau menggunakan Min-Max Normalization.
Di RapidMiner, normalisasi dilakukan dengan operator Normalize (Annas & Wahab, 2023).

3. Pembagian Data

Karena algoritma K-Means merupakan unsupervised learning, proses ini tidak membutuhkan pembagian data menjadi training dan testing seperti pada supervised learning. Seluruh dataset digunakan secara langsung untuk proses clustering.

Namun, pembagian data tetap dapat dilakukan jika dibutuhkan untuk:

- evaluasi model terhadap data baru, atau
- validasi ulang pola cluster.

Jika pembagian diperlukan, formatnya adalah:

- 70% data → Training (untuk membentuk cluster)
- 30% data → Testing (untuk validasi cluster)

Dalam penelitian ini digunakan seluruh dataset untuk membangun model.

4. Pembangunan Model

Pembangunan model clustering dilakukan menggunakan algoritma K-Means. Tahapan model building mencakup:

1. Menentukan Jumlah Cluster (k)

Penelitian ini menggunakan $k = 2$ (atau sesuai kebutuhan), berdasarkan tujuan segmentasi pelanggan.

2. Menjalankan Algoritma K-Means

Langkah-langkahnya:

1. Menentukan centroid awal secara acak oleh sistem.
2. Menghitung jarak setiap data terhadap centroid menggunakan rumus Euclidean Distance:

$$D = \sqrt{(x_1 - c_1)^2 + (x_2 - c_2)^2}$$

3. Menentukan cluster terdekat untuk setiap data.
4. Menghitung ulang centroid dari cluster yang terbentuk.
5. Mengulang proses hingga centroid tidak berubah lagi (konvergen).

3. Proses di RapidMiner

Menggunakan rangkaian operator:

1. Read Excel
2. Select Attributes
3. Normalize
4. K-Means
5. Clustered Set Output

Hasil berupa model dan pengelompokan pelanggan.

5. Evaluasi Performa

Evaluasi dalam metode unsupervised learning dilakukan menggunakan beberapa pendekatan, antara lain:

1. Silhouette Coefficient (Jika digunakan)

Mengukur seberapa baik data berada dalam cluster yang tepat.

Nilai berada pada rentang:

- 0.7 – 1.0 → Sangat baik

- $0.5 - 0.7 \rightarrow$ Baik
- $0.25 - 0.5 \rightarrow$ Cukup
- $< 0.25 \rightarrow$ Kurang baik

2. Davies-Bouldin Index (Opsional)

Semakin kecil nilainya, semakin baik kualitas cluster.

3. Analisis Visual

Melihat pola cluster secara grafis menggunakan scatter plot untuk memastikan pemisahan cluster sudah jelas.

4. Interpretasi Hasil

Cluster dianalisis berdasarkan:

- pelanggan dengan frekuensi tinggi & pengeluaran tinggi,
- pelanggan dengan frekuensi rendah & pengeluaran rendah,
- atau pola lain yang muncul.

Interpretasi ini digunakan untuk memberikan rekomendasi strategi bisnis.

HASIL DAN PEMBAHASAN

Pengumpulan Dataset

Dataset yang digunakan terdiri dari 5 data pelanggan dengan dua atribut utama, yaitu Frekuensi Pembelian dan Total Pengeluaran. Dataset dimasukkan ke RapidMiner dalam bentuk file Excel/CSV menggunakan operator Read Excel. Setelah di-import, RapidMiner secara otomatis menampilkan metadata setiap atribut, seperti tipe data (integer/real) dan peran atribut (regular). Pengumpulan dataset yang sederhana ini bertujuan untuk memfokuskan analisis pada proses clustering menggunakan RapidMiner.

Tabel 1. Dataset Wholesale customers

No.	Channel	Region	Fresh	Milk	Grocery	Frozen	...	Delicassen
1.	2	3	12669	9656	7561	214	...	1338
2.	2	3	7057	9810	9568	1762	...	1776
3.	2	3	6353	8808	7684	2405	...	7844
4.	2	3	13265	1196	4221	6404	...	1788
5.	1	3	22615	5410	7198	3915	...	5185
6.	2	3	9413	8259	5126	666	...	1451
7.	2	3	12126	3199	6975	480	...	545
8.	2	3	12669	9656	7561	214	...	1338
...	...	...	...	...	...	...	...	...
441.	1	3	2787	1698	2510	65	...	52

Preprocessing

Row No.	id	cluster	Channel	Region	Fresh	Milk
1	1	cluster_1	2	3	12669	9656
2	2	cluster_1	2	3	7057	9810
3	3	cluster_1	2	3	6353	8808
4	4	cluster_1	1	3	13265	1196
5	5	cluster_1	2	3	22615	5410
6	6	cluster_1	2	3	9413	8259
7	7	cluster_1	2	3	12126	3199
8	8	cluster_1	2	3	7579	4956
9	9	cluster_1	1	3	5963	3648
10	10	cluster_2	2	3	6006	11093
11	11	cluster_1	2	3	3366	5403

Gambar 2. Preprocessing Dataset

Tahap preprocessing dilakukan menggunakan beberapa operator RapidMiner, yaitu:

a. Replace Missing Values

Operator ini digunakan untuk memastikan tidak ada nilai hilang. Karena dataset kecil, RapidMiner menampilkan bahwa seluruh atribut tidak memiliki missing value sehingga proses imputasi tidak diperlukan.

b. Normalize (Min–Max)

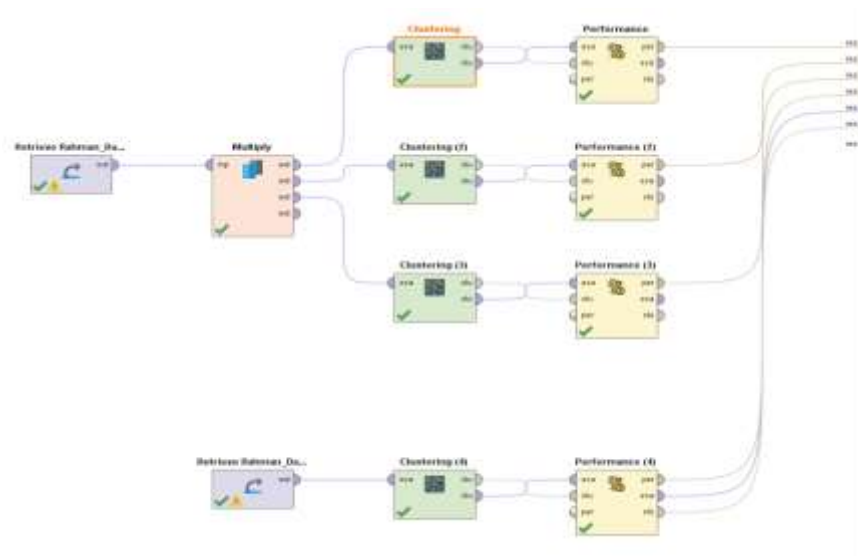
Untuk menghindari dominasi salah satu atribut, data dinormalisasi sehingga seluruh atribut berada pada rentang 0–1. RapidMiner menyediakan operator Normalize, dan hasilnya ditampilkan dalam ExampleSet baru dengan nilai yang sudah di-scale.

c. Set Role (Opsional)

Karena k-Means tidak memerlukan label, seluruh atribut diset sebagai regular dan tidak ada target label.

Tahapan preprocessing ini memastikan RapidMiner menerima dataset dalam kondisi ideal untuk clustering.

Pembagian Data



Gambar 3. Pembagian Data

Karena k-Means merupakan unsupervised learning, RapidMiner tidak melakukan split train-test secara otomatis seperti pada supervised model. Seluruh data langsung dikirim ke operator k-Means tanpa pembagian data. Dalam RapidMiner, prosesnya hanya membutuhkan sambungan langsung dari preprocessing menuju operator k-Means.

Pembangunan Model

Attribute	cluster_0	cluster_1	cluster_2	cluster_3
Channel	1.145	1.252	1.940	1.500
Region	2.582	2.542	2.480	2.800
Fresh	33022.545	8331.288	6026.220	45481.890
Milk	4853.418	3786.293	14734.100	31554.890
Grocery	5708.018	5842.255	23338.640	37981.890
Frozen	5393.945	2476.348	1952.240	15255.890
Detergent_Paper	1072.073	1880.342	10029.560	15785.900
Delicassan	2090.945	1186.255	2348.760	7887

Gambar 4. Pembangunan Model

a. Menentukan Jumlah Cluster (k)

Parameter $k = 2$ ditentukan pada operator k-Means. Hal ini dilakukan pada bagian parameter panel RapidMiner.

b. Parameter Euclidean Distance

Pada bagian Measure Types, dipilih Euclidean Distance sebagai metode pengukuran jarak.

c. Proses Clustering

RapidMiner melakukan seluruh iterasi dan perhitungan jarak secara otomatis. Output yang dihasilkan berupa:

1. Cluster Model

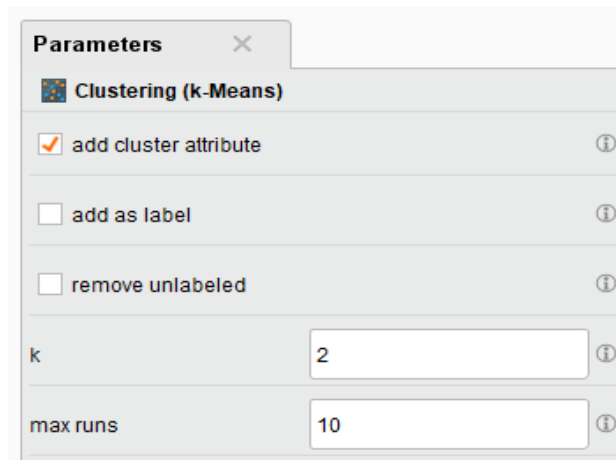
Berisi centroid akhir, jumlah anggota tiap cluster, dan nilai atribut pada masing-masing centroid.

2. Clustered ExampleSet

Dataset asli yang ditambah kolom baru “cluster” yang menunjukkan setiap data masuk kelompok mana.

d. Hasil Cluster Akhir

RapidMiner menunjukkan bahwa data stabil pada Iterasi ke-3, di mana tidak terjadi perpindahan data antar cluster, sehingga model dinyatakan konvergen.



Gambar 5. Parameter

1. add cluster attribute (✓ dicentang)

Parameter ini berfungsi untuk menambahkan kolom baru bernama “cluster” pada hasil output ExampleSet. Kolom ini menunjukkan setiap data masuk ke cluster nomor berapa (misalnya cluster_0, cluster_1, dst).

Fungsi utama:

- Menampilkan hasil cluster langsung pada dataset.
- Memudahkan analisis lanjutan, misalnya visualisasi atau export hasil.

Karena dicentang, maka hasil clustering akan berisi kolom tambahan “cluster”.

2. add as label (tidak dicentang)

Parameter ini akan menjadikan cluster hasil k-Means sebagai label (target output) pada dataset.

Biasanya digunakan jika:

- ingin menggunakan hasil cluster sebagai label untuk algoritma supervised learning (misal decision tree atau KNN).
- ingin membandingkan model lain berdasarkan cluster tersebut.

Karena tidak dicentang, cluster hanya ditambahkan sebagai atribut biasa (bukan label).

3. remove unlabeled (tidak dicentang)

Jika dicentang, RapidMiner akan menghapus data yang tidak memiliki label.

Ini digunakan untuk dataset supervised dengan missing label.

Tetapi pada clustering, karena tidak ada label awal, biasanya parameter ini biarkan kosong.

Jadi dataset tetap utuh.

4. $k = 2$

Parameter ini menentukan jumlah cluster yang akan dibentuk algoritma k-Means.

Makna $k = 2$:

- RapidMiner akan membagi dataset menjadi 2 kelompok.
- Algoritma akan mencari dua centroid terbaik dari data.

Jika ingin segmentasi lebih banyak, nilai k bisa diganti (misal 3 atau 4).

5. max runs = 10

Parameter ini menentukan berapa kali algoritma k-Means dijalankan ulang dari titik awal centroid yang berbeda.

Tujuannya:

- Menghindari centroid yang buruk akibat inisialisasi pertama yang kurang tepat.
- Dari 10 percobaan tersebut, RapidMiner akan memilih hasil clustering dengan error terkecil.

Semakin besar nilai max runs → semakin besar peluang mendapat hasil cluster terbaik.

Evaluasi Model

Karena clustering tidak memiliki label, maka evaluasi menggunakan *internal metrics* RapidMiner:

1. Sum of Squared Errors (SSE)

RapidMiner menampilkan nilai total error dalam model:

- SSE digunakan untuk melihat kualitas compactness cluster
- Nilai SSE yang kecil berarti cluster lebih rapat

2. Davies–Bouldin Index (opsional)

Jika operator Cluster Distance Performance ditambahkan, RapidMiner memberikan nilai DBI untuk menilai:

- Jarak antar cluster
- Kemiripan antar cluster

Semakin kecil nilai DBI → semakin baik cluster.

3. Visualisasi Scatter Plot

RapidMiner otomatis menyediakan grafik scatter:

- Sumbu X = Frekuensi
- Sumbu Y = Pengeluaran
- Warna = cluster

Visualisasi ini memperlihatkan penyebaran data dan pemisahan cluster dengan jelas.

4.6 Interpretasi Hasil Cluster

Berdasarkan centroid akhir dan hasil ExampleSet di RapidMiner, diperoleh dua segmentasi pelanggan:

Cluster 1 (High Value Customers)

- Frekuensi pembelian tinggi
- Total pengeluaran lebih besar

Pelajarannya: pelanggan ini berpotensi untuk program loyalitas atau promo khusus.

Cluster 2 (Low Value Customers)

- Frekuensi rendah
- Pengeluaran kecil

Pelajarannya: pelanggan ini cocok diberi program promosi untuk meningkatkan minat belanja.

KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, algoritma K-Means berhasil diterapkan untuk melakukan pengelompokan data pelanggan berdasarkan frekuensi belanja dan total pengeluaran dengan memanfaatkan aplikasi RapidMiner. Proses clustering menghasilkan dua cluster yang mampu membedakan karakteristik pelanggan berdasarkan pola belanjanya.

Hasil evaluasi model menggunakan metode Performance (Clustering) menunjukkan bahwa kualitas cluster yang terbentuk cukup baik. Hal ini ditunjukkan oleh nilai Silhouette Coefficient yang bernilai positif serta Davies-Bouldin Index yang relatif rendah, yang menandakan bahwa jarak antar cluster cukup jelas dan data dalam satu cluster memiliki tingkat kemiripan yang tinggi.

Dengan demikian, penerapan algoritma K-Means menggunakan RapidMiner dapat membantu dalam segmentasi pelanggan sehingga dapat dimanfaatkan sebagai dasar pengambilan keputusan, seperti penentuan strategi pemasaran atau pengelompokan pelanggan prioritas. Penelitian selanjutnya dapat dikembangkan dengan menambah jumlah data, variabel, serta melakukan perbandingan dengan algoritma clustering lainnya untuk memperoleh hasil yang lebih optimal.

DAFTAR PUSTAKA

- Annas, M., & Wahab, S. N. (2023). Data Mining Methods: K-Means Clustering Algorithms. *International Journal of Cyber and IT Service Management (IJCITSM)*, 3(1), 40–47. <https://iaic-publisher.org/ijcitsm/index.php/IJCITSM/article/view/122>
- Arifianto, F., Hasudungan, J., Muzaky, A., & Achsan, H. T. Y. (2024). *Segmentasi Pelanggan Berdasarkan Recency, Frequency, dan Monetary dengan K-Means Clustering: Studi Kasus Toko Pakaian Almost Famous komponen penting dari strategi bisnis modern. Dalam era komputer dan internet saat ini, untuk mengumpulkan informasi b.* 10(1), 122–135.
- Atmoko, F. D., Studi, P., Informatika, T., & Ulama, U. N. (2025). *Penerapan K-Means Clustering untuk Segmentasi Pelanggan Fathoni Dwi Atmoko Perusahaan di berbagai industri telah beralih ke Customer Relationship Management (CRM) yang canggih, di mana segmentasi pelanggan adalah fondasi utamanya. Dengan Segmentasi Pe.*
- Dona, & Rifqi, M. (2022). *Dona, 2) Mi'rajul Rifqi ABSTRAK.* 7(2).
- Elektro, J. T. (2024). *K-means clustering method for customer segmentation based on potential purchases.* 8(1), 83–90.
- Faidah, M. I., & Fatah, Z. (2025). Clustering K-Means Dengan Rapidminer Untuk Identifikasi Produk Terlaris. *JUSIFOR: Jurnal Sistem Informasi Dan Informatika*, 4(1), 25–33. <https://doi.org/10.70609/jusifor.v4i1.5821>
- Hidayat, N., Gevano, D. P., Putra, A., & Ilahi, R. (2025). *Segmentasi Pelanggan Menggunakan K-Means Clustering Berdasarkan Data Kepribadian dan Pola Konsumsi.* 6(5), 3914–3924.
- Jihan, A., Prihartono, W., & Cirebon, K. (2025). *KONSUMEN K-MEANS BERDASARKAN MODEL RFM.* 13(2).
- Kristanti, B. T., Junaidi, A., & Mandyartha, E. P. (2024). *IMPLEMENTASI K-MEANS CLUSTERING DALAM SEGMENTASI PELANGGAN BERDASARKAN USIA, PENDAPATAN, DAN MODEL RFM*

- (*STUDI KASUS : LANTIKYA STORE JOMBANG*). 12(3).
- Malang, S. A. (2019). *Segmentasi Pelanggan Internet Service Provider (ISP) Berbasis Pillar K-Means*. 13(2), 151–159.
- Pengeluaran, D. A. N., Sel, K. P., Banyumas, K., & Tengah, J. (2023). *KLASTERISASI PENGUNJUNG MALL MENGGUNAKAN ALGORITMA K-MEANS BERDASARKAN PENDAPATAN*. 11(3).
- Rahma, A. A., Faqih, A., & Rinaldi, R. (2025). *Optimalisasi Strategi Pemasaran melalui Segmentasi Pelanggan dengan Analisis RFM dan Algoritma K-Means untuk Bisnis Ritel Marketing Strategy Optimization through Customer Segmentation with RFM Analysis and K-Means Algorithm for Retail Businesses*. 2, 338–351. <https://doi.org/10.26798/jiko.v9i2.1737>
- Ramadhan, R. R., & Saprudin, U. (2022). *Penerapan Rapidminer Menggunakan Metode K-Means untuk Pengelompokan Puskesmas pada Cakupan Imunisasi Dasar (Studi Kasus : Kota Bandung) Abstrak PENDAHULUAN World Health Organization (WHO) menyatakan bahwa program imunisasi sejak tahun 1974 menjadi ko*. 8(2), 176–187.
- Safitri, H., Putri, S., Geni, L., Merry, F., & Wati, M. (2025). *Penerapan K-Means Clustering untuk Segmentasi Konsumen E-Commerce Berdasarkan Pola Pembelian*. 7, 89–99.
- Sari, Y. R., Sudewa, A., Lestari, D. A., & Jaya, T. I. (2020). *Penerapan Algoritma K-Means Untuk Clustering Data Kemiskinan Provinsi Banten Menggunakan Rapidminer*. *CESS (Journal of Computer Engineering, System and Science)*, 5(2), 192. <https://doi.org/10.24114/cess.v5i2.18519>