

Analisis Preferensi Pilihan Rasa Kopi Menggunakan Teknik Data Mining Pada Survey Konsumen Menggunakan Algoritma K-means Clustering

Muhammad Rafli Siregar¹, Ishak², Fery Setiawan³

^{1,2,3} Sistem Informasi, STMIK Triguna Dharma

Email: ¹rafli260401@gmail.com, ²ishakmkom@gmail.com, ³ferysetiawan13@gmail.com

Email Penulis Korespondensi: rafli260401@gmail.com

Abstrak

Kopi merupakan salah satu minuman yang banyak digemari oleh masyarakat dari berbagai kalangan. Setiap individu memiliki preferensi rasa kopi yang berbeda, seperti rasa manis, pahit, atau asam. Untuk memahami pola preferensi konsumen terhadap rasa kopi, dibutuhkan metode analisis yang mampu mengelompokkan data secara efektif. Penelitian ini menggunakan teknik data mining dengan algoritma *K-Means Clustering* untuk menganalisis hasil survei konsumen mengenai rasa kopi. Data dikumpulkan melalui kuesioner yang menilai atribut rasa berdasarkan skala tertentu, kemudian diproses menggunakan algoritma *K-Means* untuk membentuk beberapa kelompok preferensi. Hasil dari proses klastering menunjukkan bahwa konsumen terbagi ke dalam beberapa klaster dengan kecenderungan rasa tertentu. Temuan ini dapat digunakan oleh pelaku usaha kopi untuk merancang produk yang lebih sesuai dengan selera konsumen di setiap klaster. Penelitian ini membuktikan bahwa teknik *K-Means Clustering* efektif dalam mengidentifikasi pola preferensi rasa kopi dari data survei konsumen.

Kata Kunci: Kopi, Preferensi Rasa, Data Mining, K-means Clustering, Survei Konsumen

Abstract

Coffee is a popular beverage enjoyed by people from various backgrounds. Each individual has different coffee taste preferences, such as sweet, bitter, or sour. To understand consumer preference patterns for coffee flavors, an analytical method capable of effectively grouping data is needed. This study uses data mining techniques with the K-Means Clustering algorithm to analyze the results of a consumer survey on coffee flavors. Data were collected through a questionnaire that assessed taste attributes based on a specific scale, then processed using the K-Means algorithm to form several preference groups. The results of the clustering process indicate that consumers are divided into several clusters with specific taste tendencies. These findings can be used by coffee businesses to design products that better suit consumer tastes in each cluster. This study proves that the K-Means Clustering technique is effective in identifying coffee flavor preference patterns from consumer survey data.

Keywords: Coffee, Taste Preference, Data Mining, K-means Clustering, Consumer Survey

1. PENDAHULUAN

Kopi merupakan minuman yang banyak digemari karena aroma dan rasanya yang khas serta manfaatnya bagi tubuh (Najiyati & Danarti, 2008). Berdasarkan data International Coffee Organization (ICO), produksi kopi Indonesia tumbuh 2,75%, dan konsumsi meningkat 4,04% selama 2017–2021. Pertumbuhan ini mendorong variasi produk kopi dan persaingan antar produsen, sehingga pemahaman terhadap selera konsumen menjadi sangat penting [1].

Perkembangan teknologi dan komunikasi saat ini memberikan peluang besar dalam menganalisis data yang berkaitan dengan perilaku konsumen. Data mining, khususnya teknik *clustering*, memudahkan peneliti untuk mengelompokkan data berdasarkan kesamaan karakteristik. Salah satu metode yang banyak digunakan dalam data mining adalah algoritma *K-means Clustering*. Algoritma ini dapat membantu dalam mengidentifikasi pola-pola dalam data yang besar dan kompleks, sehingga memberikan wawasan yang lebih mendalam mengenai preferensi konsumen.

Memahami preferensi rasa kopi konsumen bukan hanya membantu dalam pengembangan produk baru, tetapi juga dalam strategi pemasaran yang lebih efektif. Dengan menggunakan data mining, produsen kopi dapat mengidentifikasi segmen-segmen konsumen berdasarkan preferensi rasa, sehingga dapat menyesuaikan penawaran produk mereka.

Tujuan penelitian ini berdasarkan penjelasan dari batasan masalah yang telah dipaparkan yaitu: Untuk menganalisis preferensi konsumen terhadap berbagai rasa kopi berdasarkan data survei yang dikumpulkan, menggunakan algoritma *k-means clustering* untuk mengelompokkan konsumen berdasarkan kesamaan preferensi rasa kopi konsumen, memberikan saran kepada produsen kopi tentang rasa yang populer.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini dapat dilakukan dengan beberapa metode penelitian yaitu sebagai berikut:

1. Data Collecting (Teknik Pengumpulan Data)
 - a. Wawancara
 - b. Kuesioner

2.2 Kopi

Hasil produksi kopi di Indonesia didominasi oleh kopi yang diusahakan di lahan perkebunan rakyat yang mencapai 95,5% per tahun. Sedangkan produksi kopi yang berasal dari perkebunan milik negara berkontribusi sekitar 2,59% dan perkebunan swasta sekitar 1,91% [2].

Kopi adalah salah satu minuman yang paling populer di dunia, dengan berbagai jenis dan rasa. Preferensi rasa kopi dapat dipengaruhi berbagai faktor, termasuk budaya, pengalaman pribadi, dan karakteristik demografi. Penelitian sebelumnya telah menunjukkan bahwa preferensi rasa kopi berbeda-beda antara individu dan kelompok.

2.3 Data Mining

Data mining merupakan proses untuk menemukan informasi baru yang sebelumnya tidak teridentifikasi secara manual dari suatu basis data. Proses ini dilakukan dengan mengekstraksi serta mengenali pola-pola penting dalam data. Data mining umumnya diterapkan pada basis data berukuran besar dan dikenal juga sebagai Knowledge Discovery in Databases (KDD) [3]. Data mining, atau penambangan data, adalah proses analisis data besar-besaran untuk menemukan pola, hubungan, atau pengetahuan baru yang berguna bagi pengambilan keputusan [4]. Dalam konteks data mining, kualitas data input sangat mempengaruhi keakuratan dan kegunaan informasi yang dihasilkan [5]. Tujuan utama dari data mining adalah untuk menemukan pola-pola penting yang dapat mendukung proses pengambilan keputusan dalam organisasi maupun bisnis [6]. Pembagian metode data mining dibagi menjadi beberapa kelompok, sebagai berikut [7]:

- a. *Association*
- b. *Clustering*
- c. *Classification*
- d. *Estimation*

2.4 K-means Clustering

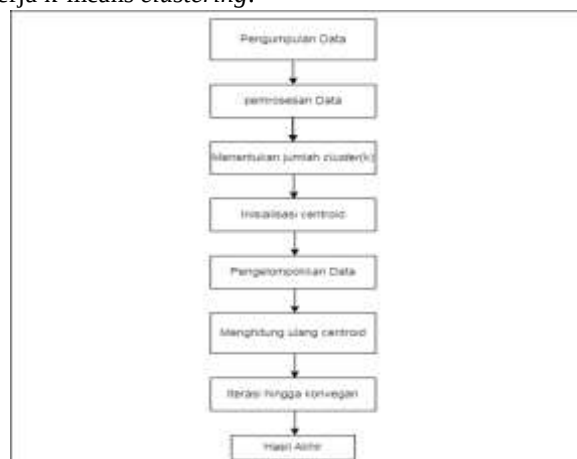
K-Means adalah metode clustering berbasis jarak yang membagi data ke dalam sejumlah cluster dan algoritma ini hanya bekerja pada atribut numeric. Algoritma K-Means termasuk partitioning clustering yang memisahkan data ke k daerah bagian yang terpisah [8]. Algoritma K-Means untuk Clustering telah digunakan dalam sejumlah penelitian sebelumnya untuk menerapkan topik analisis Clustering [9]. Clustering merupakan teknik pengelompokan data berdasarkan kesamaan karakteristik yang bertujuan untuk menemukan pola tersembunyi dalam dataset besar. Dalam dunia bisnis, hasil clustering dapat dimanfaatkan untuk mengevaluasi performa produk, mengenali perilaku pelanggan, dan menyusun strategi pemasaran yang lebih personal dan terarah [10]. Untuk mencapai tujuan tersebut, fungsi jarak digunakan sebagai ukuran kemiripan dalam klaster. Dengan melakukan perhitungan jarak terpendek antara data dan titik centroid, kita dapat memperoleh pemaksimalan kemiripan data dalam klaster tersebut [11]. Langkah-langkah melakukan clustering dengan metode K-Means adalah sebagai berikut [12]:

1. Menentukan jumlah *cluster(k)* yang diinginkan
2. Menginisialisasi centroid secara acak
3. Mengelompokkan data berdasarkan jarak terdekat ke centroid
4. Menghitung ulang posisi centroid berdasarkan rata-rata disetiap cluster
5. Mengulangi langkah 3 dan 4 sampai centroid tidak berubah atau iterasi maksimum tercapai

3. HASIL DAN PEMBAHASAN

3.1 Penerapan Metode K-means Clustering

Berikut ini adalah kerangka kerja k-means clustering:



Gambar 1 Kerangka Kerja K-means

Berikut adalah data yang digunakan dalam penelitian ini

Tabel 1. Data Preferensi Kopi

Nama	Frekuensi Konsumsi kopi(perminggu)	Tingkat Kemanisan dari 1 -5	Tingkat Keasaman dari 1 -5	Tingkat Kepahitan dari 1 -5	Tingkat aroma kopi dari 1 -5
Dafa	2	1	3	4	4
Ridho	3	3	1	2	3
Razaq	3	1	4	5	5
Resia	2	1	1	4	4
Amel	1	3	3	3	4
Tribela	2	3	1	3	4
Vozy	1	3	2	3	5
Kanaya	1	2	2	3	4
Alecia	1	4	1	1	5
Aisyah	4	2	1	3	3
Riski	2	2	2	4	5
Syifa	2	3	1	1	3
Eca	2	3	2	4	3
Cantika	1	3	2	1	3
Alya	1	2	3	3	4
Febrin	4	3	1	2	5
Aprila	1	3	1	3	4
Haliza	4	3	3	3	5
Tiara	1	5	1	1	4
Lidya	3	3	2	1	3
Putri	1	2	2	1	3
Nurul	1	2	1	1	4
Rifki	2	4	2	2	3
Faisal	1	2	2	4	4
Naufal	1	4	1	2	4
Rian	2	2	3	3	4
Ipan	2	2	3	4	3
Soraya	3	2	3	3	5
Sofyan	2	1	3	2	3
Jaki	3	2	2	4	4

Penetapan jumlah *cluster* yaitu 3. Setelah menetapkan jumlah *cluster* , tentukan titik pusat awal *cluster*(Centroid).

Berikut ini titik centroid yang dipilih :

Tabel 2. Centroid Awal

	V	W	X	Y	Z
m1	1	4	1	1	5
m2	2	1	3	2	3
m3	2	1	3	4	4

Menghitung jarak setiap variabel pada data terhadap pusat *cluster* yang sudah ditentukan.

$$d(x,y) = \sqrt{\{(x_1 - y_1)^2 + (x_2 - y_2)^2 + \dots + (x_n - y_n)^2\}}$$

a. Jarak antara Data 1 ke Pusat 1

$$\begin{aligned}
 & d(Da,ale) \\
 &= \sqrt{(2 - 1)^2 + (1 - 4)^2 + (3 - 1)^2 + (4 - 1)^2 + (4 - 5)^2} \\
 &= \sqrt{24} = 4,90
 \end{aligned}$$

b. Jarak antara Data 1 ke Pusat 2

$$\begin{aligned}
 d(Da, sof) &= \sqrt{(2-2)^2 + (1-1)^2 + (3-3)^2 + (4-2)^2 + (4-3)^2} \\
 &= \sqrt{5} = 2,24
 \end{aligned}$$

c. Jarak antara Data 1 ke Pusat 3

$$\begin{aligned}
 d(Da, Da) &= \sqrt{(2-2)^2 + (1-1)^2 + (3-3)^2 + (4-4)^2 + (4-4)^2} \\
 &= \sqrt{0} = 0,0
 \end{aligned}$$

Untuk mengetahui data 1 masuk ke *cluster* yang mana, maka diambil dari nilai minimum, maka data 1 masuk ke cluster 3. Lanjutkan perhitungan ke setiap data.

Berdasarkan proses perhitungan diatas maka diperoleh hasil jarak setiap variabel dengan *centroid* awal m1, m2, m3 sebagai berikut.

Tabel 3. Hasil Perhitungan Iterasi 1

Nama	m1	m2	m3	Min
Da	4,90	2,24	0,00	0,00
Ri	3,16	3,00	3,74	3,00
Ra	6,16	3,87	2,00	2,00
Re	4,47	3,00	2,00	2,00
Am	3,16	2,65	2,45	2,45
Tri	2,65	3,16	3,00	2,65
Vo	2,45	3,32	2,83	2,45
Kan	3,16	2,24	2,00	2,00
Ale	0,00	4,36	4,90	0,00
Ai	4,58	3,16	3,32	3,16
Ris	3,87	3,16	1,73	1,73
Syi	2,45	3,00	4,24	2,45
Eca	4,00	3,00	2,45	2,45
Can	2,45	2,65	4,00	2,45
Alya	3,61	2,00	1,73	1,73
Feb	3,32	4,00	4,12	3,32
....				
Jak	4,36	2,83	1,73	1,73

Data yang masuk ke cluster1 adalah data dari (tri, vo, ale, syi, can, feb, apr, tia, nur, rif, nau), data yang masuk ke cluster2 adalah data dari (ri, ai, lid, put, sof), data yang masuk ke cluster3 adalah data dari (da, ra, re, am, kan, ris, eca, alya, hali, fais, rian, ip, sora, jak).

Lanjutkan perhitungan ke iterasi 2 dengan menghitung ulang *centroid*

Tabel 4. Centroid baru Iterasi 2

	V	W	X	Y	Z
m1	1,5	3,4	1,3	1,8	4,0
m2	2,6	2,2	1,8	1,8	3,0
m3	2,1	2,0	2,6	3,6	4,1

Hitung lagi jarak antar data dengan *centroid* yang baru.

Tabel 5. Hasil Perhitungan Iterasi 2

Nama	m1	m2	m3	Min
Da	3,68	3,01	1,16	1,16
Ri	1,83	1,22	2,89	1,22
Ra	5,12	4,55	2,55	2,55
Re	3,26	2,88	1,90	1,90
Am	2,19	2,66	1,66	1,66
Tri	1,35	2,02	1,98	1,35
Vo	1,83	2,95	1,90	1,83
Kan	2,02	2,25	1,38	1,38
Ale	1,56	3,33	3,92	1,56
Ai	3,22	2,02	2,81	2,02
Ris	2,89	3,05	1,09	1,09

Syi	1,44	1,51	3,43	1,44
Eca	2,57	2,42	1,66	1,66
Can	1,62	1,97	3,28	1,62
Alya	2,56	2,55	1,33	1,33
Feb	2,70	2,70	3,26	2,70
Apr	1,38	2,51	2,25	1,38
Hali	3,40	3,08	2,46	2,46
Ti	1,93	3,56	4,43	1,93
Lid	2,11	1,22	3,24	1,22
Put	2,09	1,81	3,12	1,81
Nur	1,70	2,21	3,26	1,70
Rif	1,47	1,92	2,89	1,47
Fais	2,73	2,91	1,27	1,27
Nau	0,90	2,73	3,22	0,90
Rian	2,54	2,07	0,79	0,79
Ip	3,29	2,58	1,27	1,27
Sora	3,06	2,66	1,48	1,48
Sof	3,13	1,81	2,28	1,81
Jak	3,04	2,47	1,16	1,16

Dari tabel diatas didapat data baru dari *cluster1* {tri, vo, ale, syi, can, feb, apr, tia, nur, rif, nau}, *cluster2* {ri, ai, lid, put, sof}, *cluster3* {da, ra, re, am, kan, ris, eca, alya, hali, fais, rian, ip, soray, jaki}.

Lanjutkan ke iterasi 3 dikarenakan nilai *centroid* iterasi 1 dan 2 berbeda

Tabel 6. Centroid Baru Iterasi 3

	V	W	X	Y	Z
m1	1,5	3,4	1,3	1,8	4,0
m2	2,6	2,2	1,8	1,8	3,0
m3	2,1	2,0	2,6	3,6	4,1

Lalu hitung kembali data dengan rumus jarak antar data dengan centroid baru, hasil dari perhitungan pada tabel dibawah ini :

Tabel 7. Hasil Perhitungan Iterasi 3

Nama	m1	m2	m3	Min
Da	3,68	3,01	1,16	1,16
Ri	1,83	1,22	2,89	1,22
Ra	5,12	4,55	2,55	2,55
Re	3,26	2,88	1,90	1,90
Am	2,19	2,66	1,66	1,66
Tri	1,35	2,02	1,98	1,35
Vo	1,83	2,95	1,90	1,83
Kan	2,02	2,25	1,38	1,38
Ale	1,56	3,33	3,92	1,56
Ai	3,22	2,02	2,81	2,02
Ris	2,89	3,05	1,09	1,09
Syi	1,44	1,51	3,43	1,44
Eca	2,57	2,42	1,66	1,66
Can	1,62	1,97	3,28	1,62
Alya	2,56	2,55	1,33	1,33
Feb	2,70	2,70	3,26	2,70
Apr	1,38	2,51	2,25	1,38
Hali	3,40	3,08	2,46	2,46
Ti	1,93	3,56	4,43	1,93
Lid	2,11	1,22	3,24	1,22
Put	2,09	1,81	3,12	1,81
Nur	1,70	2,21	3,26	1,70

Rif	1,47	1,92	2,89	1,47
fais	2,73	2,91	1,27	1,27
Nau	0,90	2,73	3,22	0,90
Rian	2,54	2,07	0,79	0,79
Ip	3,29	2,58	1,27	1,27
Sora	3,06	2,66	1,48	1,48
Sof	3,13	1,81	2,28	1,81
Jak	3,04	2,47	1,16	1,16

Dari tabel diatas didapat data baru dari *cluster1* {tri, vo, ale, syi, can, feb, apr, tia, nur, rif, nau}, *cluster2* {ri, ai, lid, put, sof}, *cluster3* {da, ra, re, am, kan, ris, eca, alya, hali, fais, rian, ip, soray, jaki}.

untuk memastikan apakah iterasi berhenti di iterasi 3 maka, dilakukan menghitung ulang centroid untuk memastikan bahwa nilai centroid tidak berubah. Maka ini hasil centroid yang didapat :

Tabel 8. Hasil Hitung Ulang Centroid

	V	W	X	Y	Z
m1	1,5	3,4	1,3	1,8	4,0
m2	2,6	2,2	1,8	1,8	3,0
m3	2,1	2,0	2,6	3,6	4,1

Setelah dilakukan penghitungan ulang centroid, ternyata nilai centroid tidak berubah atau tetap, maka iterasi diberhentikan sampai iterasi 3.

Berikut ini adalah hasil akhir yang didapat:

Tabel 9. Hasil Akhir

Nama	m1	m2	m3	Min
Da	3,68	3,01	1,16	1,16
Ri	1,83	1,22	2,89	1,22
Ra	5,12	4,55	2,55	2,55
Re	3,26	2,88	1,90	1,90
Am	2,19	2,66	1,66	1,66
Tri	1,35	2,02	1,98	1,35
Vo	1,83	2,95	1,90	1,83
Kan	2,02	2,25	1,38	1,38
Ale	1,56	3,33	3,92	1,56
Ai	3,22	2,02	2,81	2,02
Ris	2,89	3,05	1,09	1,09
Syi	1,44	1,51	3,43	1,44
Eca	2,57	2,42	1,66	1,66
Can	1,62	1,97	3,28	1,62
Alya	2,56	2,55	1,33	1,33
Feb	2,70	2,70	3,26	2,70
Apr	1,38	2,51	2,25	1,38
Hali	3,40	3,08	2,46	2,46
Ti	1,93	3,56	4,43	1,93
Lid	2,11	1,22	3,24	1,22
Put	2,09	1,81	3,12	1,81
Nur	1,70	2,21	3,26	1,70
Rif	1,47	1,92	2,89	1,47
fais	2,73	2,91	1,27	1,27
Nau	0,90	2,73	3,22	0,90
Rian	2,54	2,07	0,79	0,79
Ip	3,29	2,58	1,27	1,27
Sora	3,06	2,66	1,48	1,48
Sof	3,13	1,81	2,28	1,81
Jak	3,04	2,47	1,16	1,16

Maka didapat hasil bahwa *cluster1* menyukai kopi rasa manis berjumlah 11, *cluster2* menyukai kopi rasa asam berjumlah 5, *cluster3* menyukai kopi rasa pahit berjumlah 14. Didapat bahwasanya rasa kopi yang paling diminati para penyuka kopi adalah rasa manis dan pahit dengan aroma kopi yang kuat.

3.2 Implementasi Sistem

Berikut ini adalah tahapan dari pengaplikasian Data Mining dengan menggunakan metode *K-means* dalam mengelompokkan data preferensi rasa kopi :

1. Form Survey

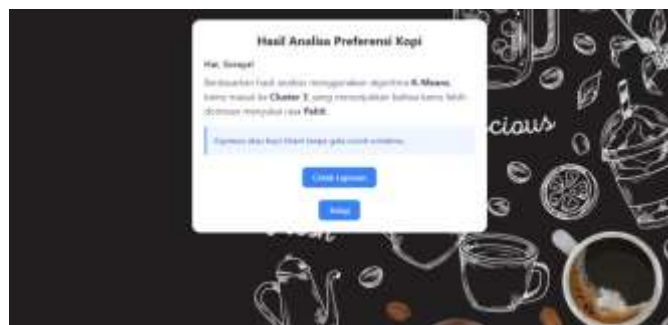
Form survey merupakan tampilan yang pertama muncul ketika sistem diakses



Gambar 2 Form Survey

2. Form Laporan

Form ini muncul ketika responden telah mensubmit data di form survey



Gambar 3 Form Laporan

4. KESIMPULAN

Kesimpulan merupakan jawaban dari rumusan masalah yang menggambarkan hasil penelitian yang dilakukan. Maka dibutuhkan sebuah sistem yang mampu mengelompokkan data preferensi rasa kopi, maka dari itu di rancanglah sebuah sistem yang mampu menerapkan pemodelan algoritma K-Means Clustering sehingga dapat membantu menganalisa rasa kopi konsumen. Aplikasi yang di bangun berbasis Website telah mampu mengimplementasikan Data Mining dalam melakukan pengelompokkan data berdasarkan pilihan rasa kopi konsumen. Metode *K-means* terbukti efisiensinya dalam menganalisis preferensi rasa kopi konsumen, mengelompokkan data berdasarkan kesamaan karakteristik dalam memilih rasa kopi.

UCAPAN TERIMAKASIH

Terimakasih di ucapkan kepada Allah Subhanahuwata'ala yang telah memberikan Rahmat dan karunia sehingga mampu menyelesaikan jurnal ini. Kemudian kepada Bapak Ishak, dan Bapak Feri Setiawan atas arahan dan bimbingannya selama proses pengerjaan Skripsi hingga sampai ke penyusunan jurnal ini dan seluruh jajaran Manajemen, Dosen serta pegawai kampus STMIK Triguna Dharma yang telah banyak membantu baik dari segi informasi ataupun dukungan lainnya.

DAFTAR PUSTAKA

- [1] A. Nurul and R. Nurmalina, "SIKAP DAN PREFERENSI KONSUMEN TERHADAP," vol. 12, no. 1, pp. 120–130, 2024.
- [2] A. Info, "Peramalan Produksi , Volume Ekspor dan Nilai Ekspor Kopi Indonesia Abstrak," vol. 5, no. 2, pp. 178–

- 185, 2025.
- [3] M. Syahril, K. Erwansyah, and M. Yetri, “Penerapan Data Mining Untuk Menentukan Pola Penjualan Peralatan Sekolah Pada Brand Wigglo Dengan Menggunakan Algoritma Apriori,” vol. 3, no. 1, pp. 118–136, 2020.
 - [4] J. Penerapan, T. Informasi, and E. Saputri, “IT-EXPLORE Teknik dan aplikasi data mining di Indonesia : tinjauan literatur satu dekade (2015-2024),” vol. 04, pp. 138–149, 2025, doi: 10.24246/itexplore.v4i2.2025.pp138-149.
 - [5] A. Zaelani, M. Fikriansyah, M. Syahdan, and I. A. Hakim, “Peran Keamanan Basis Data Relasional dalam Menjamin Kualitas Data untuk Proses Data Mining : Studi Kasus Klasifikasi Aktivitas Akses Berisiko,” vol. 4, no. 2, pp. 5286–5292, 2025.
 - [6] A. A. Baskara, N. M. Piranti, and M. F. Romdendine, “FRAMEWORK DATA MINING : SEBUAH SURVEI,” vol. 9, no. 3, pp. 4886–4895, 2025.
 - [7] M. T. Informasi, U. Pembangunan, and P. Budi, “Implementasi Data Mining Untuk Clustering Produktivitas Bawang Merah Menggunakan Metode K-Means,” vol. 07, no. 02, pp. 109–121, 2025.
 - [8] S. Gantina and A. H. Nasyuha, “Implementasi Data Mining Dalam Pengelompokan Data Transaksi Penjualan Kosmetik di WN Kosmetik Dengan Menggunakan Metode K-Means Clustering,” no. x, pp. 1–11, 2020.
 - [9] M. A. Aldi *et al.*, “IMPLEMENTASI K-MEANS CLUSTER ING DALAM PENGELOMPOKAN DATA KUNJUNGAN,” vol. 2, no. 1, pp. 13–19, 2025.
 - [10] D. K. Clustering, S. Informasi, F. T. Informasi, and U. T. Digital, “2 * 1,2,” vol. 5, no. 1, pp. 435–448, 2025.
 - [11] R. Syahri, T. Informatika, and S. Selatan, “ALGORITMA K-MEANS CLUSTERING : SEBUAH STUDI LITERATUR K-MEANS CLUSTERING ALGORITHM : A LITERATUR STUDY,” vol. x, no. x, pp. 1–7, 2023, doi: 10.12345/juri.
 - [12] A. Sulistiyawati and E. Supriyanto, “Implementasi Algoritma K-means Clustering dalam Penentuan Siswa Kelas Unggulan,” vol. 15, no. 2, pp. 25–36, 2020.