



Clustering Pemesanan Spare Part K-Means di PT Prima Multi Peralatan Cabang Belawan

M Tsaqif Al Mutawakkil
Simbolon*

Universitas Islam Negeri
Sumatera Utara,
Indonesia

Sriani

Universitas Islam Negeri
Sumatera Utara,
Indonesia

***Corresponding author:**

M Tsaqif Al Mutawakkil Simbolon, Universitas
Islam Negeri Sumatera Utara, Indonesia.
✉ tsaqifmuhammad19@gmail.com

Article Info:

Article history:

Received: July 18, 2026

Revised: August 20, 2026

Accepted: August 21, 2026

Available Online: August 26, 2026

Keywords:

Clustering; Data Mining; Inventory;
K-Means Clustering; Spare Parts.



This is an open access article under the
[CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

Copyright © 2026 by Author. Published by
Politeknik Siber Cerdika International

Kata Kunci:

Clustering; Data Mining;
Persediaan; K-Means Clustering;
Sparepart.

Abstract

Background: Effective spare parts inventory management is essential to support the operational continuity of port equipment. PT Prima Multi Peralatan Belawan Branch manages a large volume of spare parts ordering data, requiring systematic classification based on demand levels to support inventory decision-making.

Objective: This study aims to implement the K-Means Clustering algorithm to classify spare parts ordering data at the Container Crane (CC) Site of PT Prima Multi Peralatan Belawan Branch based on order quantity and ordering frequency.

Methods: The study used purchase requisition data from the Container Crane (CC) Site collected from January to June 2024. The research stages included spare parts filtering, data cleaning and aggregation, normalization using Min-Max Scaling, determination of the optimal number of clusters using the Elbow Method, K-Means clustering, and cluster evaluation using Sum of Squared Error (SSE), Davies-Bouldin Index (DBI), and Silhouette Coefficient Index (SCI).

Results: A total of 188 spare parts were classified into three clusters. Cluster 0 (Rarely Ordered) contained 169 items (89.9%), Cluster 1 (Moderately Ordered) contained 16 items (8.5%), and Cluster 2 (Frequently Ordered) contained 3 items (1.6%). The clustering evaluation produced an SSE of 1.670368, DBI of 0.369410, and SCI of 0.875059, indicating good clustering quality.

Conclusion: The K-Means Clustering algorithm effectively classified spare parts according to their ordering characteristics and can support inventory decision-making by helping the company identify different levels of spare parts demand.

Abstrak

Latar belakang: Pengelolaan persediaan sparepart yang efektif sangat penting untuk mendukung kelancaran operasional peralatan pelabuhan. PT Prima Multi Peralatan Cabang Belawan mengelola data pemesanan sparepart dalam jumlah besar sehingga diperlukan pengelompokan secara sistematis berdasarkan tingkat kebutuhan untuk mendukung pengambilan keputusan persediaan.

Tujuan: Penelitian ini bertujuan menerapkan algoritma K-Means Clustering untuk mengelompokkan data pemesanan sparepart pada Site Container Crane (CC) PT Prima Multi Peralatan Cabang Belawan berdasarkan jumlah dan frekuensi pemesanan.

Metode: Penelitian menggunakan data purchase requisition pada Site Container Crane (CC) periode Januari hingga Juni 2024. Tahapan penelitian meliputi filtering data sparepart, cleaning dan agregasi data, normalisasi menggunakan Min-Max Scaling, penentuan jumlah cluster

optimal menggunakan Elbow Method, proses clustering menggunakan K-Means, serta evaluasi cluster menggunakan Sum of Squared Error (SSE), Davies-Bouldin Index (DBI), dan Silhouette Coefficient Index (SCI).

Hasil: Sebanyak 188 sparepart berhasil dikelompokkan menjadi tiga cluster. Cluster 0 (Jarang Dipesan) terdiri atas 169 item (89,9%), Cluster 1 (Cukup Sering Dipesan) sebanyak 16 item (8,5%), dan Cluster 2 (Sering Dipesan) sebanyak 3 item (1,6%). Evaluasi clustering menghasilkan nilai SSE sebesar 1,670368, DBI sebesar 0,369410, dan SCI sebesar 0,875059 yang menunjukkan kualitas clustering yang baik.

Kesimpulan: Algoritma K-Means Clustering mampu mengelompokkan sparepart berdasarkan karakteristik pemesanannya dan dapat mendukung pengambilan keputusan persediaan dengan membantu perusahaan mengidentifikasi perbedaan tingkat kebutuhan sparepart.

To cite this article: Simbolon, M. T. A. M. & Sriani. (2026). Clustering Pemesanan Spare Part K-Means di PT Prima Multi Peralatan Cabang Belawan. *Journal of Business, Social and Technology*, 7(4), 1695-1710. <https://doi.org/10.59261/jbt.v7i4.801>

PENDAHULUAN

Sparepart merupakan komponen pengganti pada mesin atau peralatan yang mengalami kerusakan atau keausan sehingga keberadaannya penting dalam menjaga kinerja dan umur operasional mesin (Andrian, 2026; Widari et al., 2025; Wohon et al., 2023). Dalam kegiatan industri, sparepart berperan dalam mendukung kelancaran operasional dan proses produksi. Pada industri otomotif, sparepart dapat berupa filter, pelumas, rem, baterai, dan ban, sedangkan pada industri manufaktur dapat berupa bearing, gear, dan komponen elektronik. Ketersediaan sparepart yang sesuai dengan kebutuhan menjadi penting karena berkaitan dengan upaya mengurangi downtime dan menjaga produktivitas operasional (Aprilia et al., 2026; Arizqi & Putri, 2025; Ramadani et al., 2025). Oleh karena itu, pengelolaan persediaan sparepart menjadi bagian penting dalam aktivitas pemeliharaan agar peralatan tetap dapat beroperasi sesuai dengan fungsinya. Ketersediaan sparepart yang memadai juga memungkinkan proses pemeliharaan dilakukan secara lebih efektif sehingga gangguan operasional dan waktu henti peralatan dapat diminimalkan (Abdelmeguid & Corsini, 2026; Fahri et al., 2025; Zhang et al., 2021). Permasalahan pengelolaan sparepart juga ditemukan pada PT Prima Multi Peralatan Cabang Belawan. Perusahaan memiliki data pemesanan sparepart dalam jumlah besar sehingga pengelolaan dan pemantauan kebutuhan barang menjadi lebih sulit. Kondisi tersebut menyebabkan perusahaan mengalami kesulitan dalam mengelompokkan sparepart berdasarkan tingkat kebutuhannya, mulai dari sparepart yang sering dipesan hingga yang jarang dipesan. Apabila data pemesanan tidak dikelompokkan secara sistematis, perusahaan akan lebih sulit menentukan prioritas kebutuhan persediaan. Kondisi tersebut dapat berdampak pada efisiensi pengelolaan persediaan, keterlambatan pengadaan, penumpukan stok yang tidak diperlukan, serta berpotensi mengganggu produktivitas operasional.

Salah satu pendekatan yang dapat digunakan untuk mengelompokkan data pemesanan tersebut adalah K-Means Clustering. K-Means merupakan metode clustering yang mengelompokkan data berdasarkan kemiripan karakteristik sehingga data dengan karakteristik yang relatif sama berada dalam satu kelompok, sedangkan data dengan karakteristik berbeda ditempatkan pada kelompok lainnya. K-Means termasuk metode *partitional clustering* karena proses pengelompokan dilakukan dengan menentukan sejumlah cluster dan centroid sebagai pusat masing-masing kelompok (Shepyantoni et al., 2024; Wananda & Nugraha, 2023). Dalam konteks penelitian ini, metode tersebut digunakan untuk mengelompokkan sparepart berdasarkan jumlah pemesanan dan frekuensi pemesanan sehingga pola kebutuhan sparepart dapat diidentifikasi secara lebih sistematis.

Sejumlah penelitian sebelumnya telah menunjukkan penggunaan metode clustering dalam pengelompokan data. Sitinjak et al. (2024) menerapkan algoritma K-Means Clustering untuk menentukan strategi penjualan sparepart pada CV. Exa Jaya Motor. Penelitian tersebut mengelompokkan 60 jenis barang ke dalam tiga kategori berdasarkan tingkat penjualannya, yaitu barang paling laris, laris, dan kurang laris. Hasil evaluasi menggunakan rasio BCV/WCV sebesar 0,2487 serta Davies-Bouldin Index (DBI) sebesar 0,00395 menunjukkan bahwa cluster yang terbentuk memiliki tingkat kekompakan dan pemisahan yang baik. Temuan tersebut menunjukkan bahwa K-Means dapat digunakan untuk menyederhanakan pengelompokan data

sehingga hasilnya dapat membantu proses pengambilan keputusan ([Alvianatinova et al., 2024](#); [Sitinjak et al., 2024](#); [Wibowo & Elisa, 2026](#)).

Penelitian Ramdhan et al. (2022) juga menerapkan K-Means untuk clustering data persediaan barang. Penelitian tersebut memperlihatkan bahwa K-Means dapat dimanfaatkan dalam pengelompokan data persediaan. Meskipun demikian, penelitian tersebut belum secara spesifik membahas pengelompokan pemesanan sparepart berdasarkan karakteristik jumlah dan frekuensi pemesanan pada Site Container Crane (CC) PT Prima Multi Peralatan Cabang Belawan. Penelitian sebelumnya juga belum menjadi dasar untuk menjelaskan pola kebutuhan sparepart pada objek penelitian ini. Dengan demikian, masih terdapat ruang penelitian untuk menerapkan proses clustering yang secara khusus diarahkan pada data pemesanan sparepart berdasarkan karakteristik kebutuhan pada Site CC.

Clustering adalah proses analisa informasi yang kerap digunakan sebagai salah satu proses untuk Data Mining yang bertujuan untuk mengumpulkan informasi yang memiliki karakter yang sama pada satu kawasan yang sama dan informasi yang memiliki karakter yang berbeda ke kawasan lain ([Rahayu et al., 2024](#); [Suryadi & Supriatna, 2019](#)).

K-Means Clustering adalah suatu metode penganalisaan data atau metode Data Mining yang melakukan proses pemodelan Unsupervised learning dan menggunakan metode yang mengelompokkan data berbagai partisi. K-Means Clustering memiliki objective yaitu meminimalisasi object function yang telah di atur pada proses clasterisasi. Dengan cara minimalisasi variasi antar 1 cluster dengan maksimalisasi variasi dengan data di cluster lainnya ([Ramdhan et al., 2022](#)). Berikut ini merupakan rumus dari K-Means.

$$\bar{v}_{ij} = \frac{1}{N_i} \sum_{k=1}^{N_i} x_{kj}$$

Clustering Algoritma (K-Means) memiliki tujuan untuk meminimalisasikan fungsi objective yang telah di set dalam proses K-Means clustering. Tujuan tersebut dilakukan dengan cara meminimalikan variasi data yang ada didalam cluster dan memaksimalkan variasi data yang ada di cluster lainnya ([Ayu et al., 2026](#); [Darmi & Setiawan, 2016](#)).

Python adalah bahasa pemrograman tingkat tinggi yang dikembangkan oleh Guido van Rossum pada tahun 1991. Python didesain dengan filosofi yang menekankan keterbacaan kode dan sintaksis yang sederhana, sehingga membuatnya menjadi pilihan yang populer untuk pemula maupun pengembang berpengalaman ([Astrina et al., 2025](#)).

Penelitian sebelumnya yang dilakukan oleh Ramdhan et al. (2022) menunjukkan penerapan K-Means untuk K-Means clustering data persediaan barang secara umum. Namun, penelitian tersebut belum secara spesifik menganalisis pengaruh inisialisasi centroid pada perolehan hasil K-Means clustering yang optimal. Kontribusi penelitian ini terletak pada penerapan K-Means dengan pemilihan centroid yang lebih objektif menggunakan Elbow Method untuk optimalisasi penentuan jumlah cluster, serta analisis mendalam terhadap distribusi pemesanan spare parts pada Site Container Crane (CC) di PT Prima Multi Peralatan Cabang Belawan. Penelitian ini dirancang untuk memberikan dasar pengambilan keputusan inventory yang lebih efektif melalui kategorisasi demand berdasarkan hasil K-Means clustering.

Berdasarkan kondisi tersebut, kebutuhan penelitian tidak hanya terletak pada penerapan algoritma K-Means, tetapi juga pada bagaimana data pemesanan sparepart dapat diolah menjadi kelompok kebutuhan yang dapat diinterpretasikan. Data pemesanan yang sebelumnya berupa kumpulan transaksi perlu diolah berdasarkan karakteristik jumlah pemesanan dan frekuensi pemesanan agar perusahaan dapat membedakan sparepart yang jarang dipesan, cukup sering dipesan, dan sering dipesan. Pengelompokan tersebut diperlukan untuk memberikan gambaran yang lebih terstruktur mengenai pola pemesanan sparepart dan mendukung pengelolaan persediaan sesuai dengan karakteristik masing-masing kelompok.

Kontribusi penelitian ini terletak pada penerapan K-Means Clustering terhadap data pemesanan sparepart pada Site Container Crane (CC) PT Prima Multi Peralatan Cabang Belawan dengan menggunakan jumlah pemesanan dan frekuensi pemesanan sebagai variabel pengelompokan. Penelitian juga menggunakan Min-Max Scaling untuk menyamakan rentang nilai antarvariabel dan Elbow Method untuk membantu menentukan jumlah cluster yang digunakan dalam proses clustering. Selanjutnya, kualitas cluster dievaluasi menggunakan Sum of Squared

Error (SSE), Davies-Bouldin Index (DBI), dan Silhouette Coefficient Index (SCI). Dengan tahapan tersebut, pengelompokan tidak hanya menghasilkan kategori sparepart, tetapi juga disertai evaluasi terhadap kualitas cluster yang terbentuk.

Berdasarkan uraian tersebut, penelitian ini bertujuan menerapkan algoritma K-Means Clustering untuk mengelompokkan data pemesanan sparepart pada Site Container Crane (CC) PT Prima Multi Peralatan Cabang Belawan berdasarkan jumlah dan frekuensi pemesanan. Hasil pengelompokan diharapkan dapat memberikan informasi mengenai karakteristik kebutuhan sparepart berdasarkan kategori jarang dipesan, cukup sering dipesan, dan sering dipesan. Secara praktis, hasil penelitian dapat digunakan sebagai informasi pendukung dalam pengambilan keputusan pengelolaan persediaan sparepart. Secara metodologis, penelitian ini memberikan penerapan proses pengelompokan data pemesanan melalui tahapan preprocessing, normalisasi, penentuan jumlah cluster, K-Means Clustering, evaluasi cluster, dan visualisasi hasil menggunakan Python.

METODE

Penelitian ini menggunakan pendekatan kuantitatif dengan penerapan algoritma K-Means Clustering untuk mengelompokkan data pemesanan sparepart berdasarkan karakteristik jumlah dan frekuensi pemesanan. Data penelitian bersumber dari dokumen *purchase requisition* PT Prima Multi Peralatan Cabang Belawan pada Site Container Crane (CC) periode Januari hingga Juni 2024. Data tersebut digunakan untuk mengidentifikasi pola kebutuhan sparepart melalui proses pengelompokan sehingga setiap item dapat dikategorikan berdasarkan tingkat pemesanannya.

Data awal yang digunakan berjumlah 268 item. Pada tahap *preprocessing*, dilakukan penyaringan data untuk memastikan bahwa data yang dianalisis hanya berasal dari kategori sparepart. Data dari kategori lain, seperti *consumable*, jasa, *utility*, dan *tool*, tidak digunakan karena penelitian difokuskan pada pemesanan sparepart. Selanjutnya, dilakukan *data cleaning* terhadap nilai kosong (*missing value*) dan data yang tidak valid serta pemeriksaan format data numerik agar dapat diproses pada tahap berikutnya. Setelah proses penyaringan dan pembersihan, data transaksi diintegrasikan berdasarkan item sparepart sehingga diperoleh 188 item sparepart yang digunakan dalam proses clustering.

Variabel yang digunakan dalam proses clustering terdiri atas jumlah pemesanan (X1) dan frekuensi pemesanan (X2). Jumlah pemesanan menunjukkan jumlah unit sparepart yang dipesan, sedangkan frekuensi pemesanan menunjukkan intensitas pemesanan setiap item berdasarkan nomor *purchase requisition* yang berbeda. Kedua variabel tersebut digunakan karena dapat merepresentasikan karakteristik pemesanan setiap sparepart dan menjadi dasar untuk membedakan tingkat kebutuhan antaritem.

Sebelum proses clustering dilakukan, kedua variabel dinormalisasi menggunakan metode Min-Max Scaling sehingga nilai berada pada rentang 0 sampai 1. Normalisasi dilakukan karena jumlah pemesanan dan frekuensi pemesanan memiliki rentang nilai yang berbeda. Proses ini bertujuan mencegah variabel dengan nilai yang lebih besar mendominasi perhitungan jarak pada algoritma K-Means. Normalisasi dilakukan menggunakan persamaan Min-Max Scaling sebagai berikut:

$$X' = (X - X_{\min}) / (X_{\max} - X_{\min})$$

Keterangan:

X' = nilai hasil normalisasi;

X = nilai data awal;

Xmin = nilai minimum pada variabel;

Xmax = nilai maksimum pada variabel.

Jumlah cluster ditentukan menggunakan *Elbow Method* dengan mengamati perubahan nilai Sum of Squared Error (SSE) pada beberapa jumlah cluster. Jumlah cluster dipilih pada titik yang menunjukkan pola siku (*elbow*) pada grafik. Berdasarkan proses tersebut, jumlah cluster yang digunakan dalam penelitian adalah tiga cluster. Selanjutnya, algoritma K-Means diterapkan terhadap data yang telah dinormalisasi untuk mengelompokkan sparepart berdasarkan

kedekatan karakteristik jumlah dan frekuensi pemesanannya.

Proses K-Means dilakukan dengan menghitung jarak setiap data terhadap centroid menggunakan *Euclidean Distance*. Data kemudian ditempatkan pada cluster dengan jarak centroid terdekat. Setelah seluruh data memperoleh cluster, centroid diperbarui berdasarkan nilai rata-rata anggota masing-masing cluster dan proses perhitungan diulang sampai tidak terjadi perubahan cluster atau telah mencapai kondisi konvergen. Jarak Euclidean dihitung menggunakan persamaan berikut:

$$d(x,c) = \sqrt{[(x_1 - c_1)^2 + (x_2 - c_2)^2]}$$

Keterangan:

$d(x,c)$ = jarak data terhadap centroid;

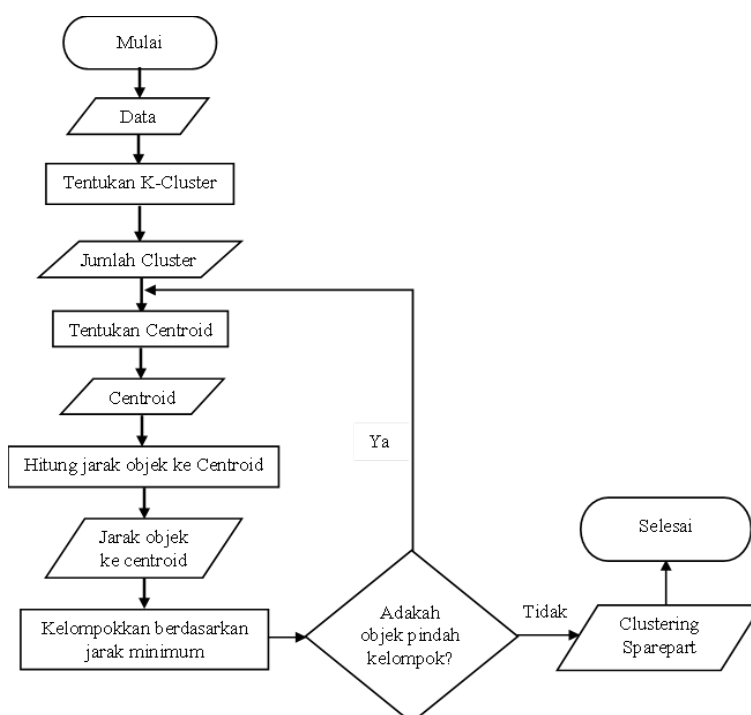
x_1 dan x_2 = nilai variabel pada data;

c_1 dan c_2 = nilai variabel pada centroid.

Hasil clustering kemudian diinterpretasikan berdasarkan karakteristik centroid masing-masing kelompok. Cluster dengan karakteristik jumlah dan frekuensi pemesanan rendah diberi label Jarang Dipesan, cluster dengan karakteristik sedang diberi label Cukup Sering Dipesan, dan cluster dengan karakteristik tinggi diberi label Sering Dipesan. Proses pengolahan data, normalisasi, penentuan jumlah cluster, clustering, evaluasi, dan visualisasi dilakukan menggunakan Python melalui Jupyter Notebook.

Kualitas hasil clustering dievaluasi menggunakan tiga ukuran, yaitu Sum of Squared Error (SSE), Davies-Bouldin Index (DBI), dan Silhouette Coefficient Index (SCI). SSE digunakan untuk menggambarkan tingkat kedekatan data terhadap pusat cluster, sedangkan DBI dan SCI digunakan untuk mengevaluasi kualitas pemisahan dan kekompakan cluster yang terbentuk. Hasil akhir clustering selanjutnya divisualisasikan menggunakan *scatter plot* untuk memperlihatkan persebaran data dan posisi centroid pada masing-masing cluster.

Alur penelitian secara keseluruhan meliputi pengumpulan data *purchase requisition*, filtering data sparepart, *data cleaning*, agregasi data, normalisasi menggunakan Min-Max Scaling, penentuan jumlah cluster menggunakan *Elbow Method*, proses K-Means Clustering, evaluasi hasil clustering menggunakan SSE, DBI, dan SCI, serta visualisasi hasil. Tahapan tersebut dirangkum dalam flowchart penelitian pada Gambar 1.



Gambar 1. Flowchart Algoritma K-Means

HASIL DAN PEMBAHASAN

Hasil Penelitian

Penerapan metode K-Means Clustering untuk pengelompokan spare parts di PT Prima Multi Peralatan Cabang Belawan bertujuan untuk mengidentifikasi pola pemesanan berdasarkan karakteristik jumlah unit yang dipesan dan frekuensi pemesanan setiap item.

Representasi Data

Data penelitian berasal dari dokumen purchase requisition PT Prima Multi Peralatan Cabang Belawan Site Container Crane (CC) periode Januari hingga Juni 2024. Setelah proses seleksi dan preprocessing, diperoleh 268 item sparepart. Data kemudian diagregasi berdasarkan item sparepart sehingga menghasilkan 188 item yang digunakan dalam proses clustering. Data hasil agregasi memuat nama sparepart, jumlah pemesanan (X1), dan frekuensi pemesanan (X2). Sepuluh data yang ditampilkan pada Tabel 1 merupakan contoh representasi dari dataset yang digunakan untuk memperlihatkan karakteristik data sebelum proses clustering.

Tabel 1

Representasi Data

No	Nama Spare part	Jumlah Pemesanan (X1)	Frekuensi Pemesanan (X2)
1	Seal Dia 75 x Dia 110mm x 25mm	1	1
2	Battery 12 V 200AH - PN190HS2-N200 YUASA	6	2
3	Varistor	6	2
4	Racor Separator Filter Merk Sakura SF-79370	8	4
5	Fuel Separator	20	3
6	Water Filter	20	4
7	Lube Filter	30	3
8	Racor Separator Filter Merk Parker PM 2020	40	5
9	Oli Filter Merk CAT Pn. 1R - 0726	66	5
10	Fuel Separator Filter Merk CAT Pn. 1R - 0756	125	5

Implementasi K-Means

Sebelum proses K-Means clustering dilakukan, data dinormalisasi menggunakan metode Min-Max Scaling agar kedua variabel berada pada rentang yang sama (0 sampai 1). Tujuannya agar variabel yang memiliki nilai besar (X1) tidak mendominasi perhitungan jarak Euclidean. Rumus Min-Max Scaling adalah sebagai berikut.

$$X' = \frac{X - X_{min}}{X_{max} - X_{min}}$$

Data 1 : X1 = 1, X2 = 1

$$X1' = \frac{1 - 1}{125 - 1} = \frac{0}{124} = 0.000000$$

$$X2' = \frac{1 - 1}{5 - 1} = \frac{0}{4} = 0.000000$$

Data 2 : X1 = 6, X2 = 2

$$X1' = \frac{6 - 1}{125 - 1} = \frac{5}{124} = 0.040323$$

$$X2' = \frac{2 - 1}{5 - 1} = \frac{1}{4} = 0.250000$$

Setelah selesai melakukan perhitungan normalisasi data pada semua data, dapatlah hasil seperti pada tabel 2.

Tabel 2*Normalisasi Data*

No	Nama <i>Spare part</i>	X1	X2	X1'	X2'
1	Seal Dia 75 x Dia 110mm x 25mm	1	1	0.000000	0.000000
2	Battery 12 V 200AH - PN190HS2-N200 YUASA	6	2	0.040323	0.250000
3	Varistor	6	2	0.040323	0.250000
4	Racor Separator Filter Merk Sakura SF-79370	8	4	0.056452	0.750000
5	Fuel Separator	20	3	0.153226	0.500000
6	Water Filter	20	4	0.153226	0.750000
7	Lube Filter	30	3	0.233871	0.500000
8	Racor Separator Filter Merk Parker PM 2020	40	5	0.314516	1.000000
9	Oli Filter Merk CAT Pn. 1R - 0726	66	5	0.524194	1.000000
10	Fuel Separator Filter Merk CAT Pn. 1R - 0756	125	5	1.000000	1.000000

Setelah proses normalisasi, jumlah cluster ditentukan menggunakan Elbow Method berdasarkan perubahan nilai SSE. Hasil pengujian menunjukkan jumlah cluster yang digunakan dalam proses K-Means adalah tiga cluster. Berdasarkan karakteristik centroid yang terbentuk, ketiga cluster tersebut kemudian diinterpretasikan sebagai Cluster 0 (Jarang Dipesan), Cluster 1 (Cukup Sering Dipesan), dan Cluster 2 (Sering Dipesan). Untuk memperlihatkan mekanisme perhitungan K-Means secara manual, sepuluh data pada Tabel 2 digunakan sebagai contoh perhitungan. Pada ilustrasi tersebut, centroid awal C0, C1, dan C2 menggunakan nilai dari data ke-1, ke-4, dan ke-10 sebagaimana ditampilkan pada Tabel 3. Perhitungan manual ini digunakan untuk menunjukkan proses penentuan jarak, pengelompokan data, dan pembaruan centroid, sedangkan hasil clustering keseluruhan mengacu pada pengolahan dataset menggunakan Python.

Tabel 3*Centroid Awal*

Centroid Awal	Jumlah Pemesanan	Frekuensi Pemesanan
Centroid 0	0.000000	0.000000
Centroid 1	0.056452	0.750000
Centroid 2	1.000000	1.000000

Setelah centroid awal ditentukan, tahap selanjutnya adalah menghitung jarak antara setiap data dengan masing-masing centroid yang telah ditetapkan menggunakan metode Euclidean Distance, dengan rumus yang seperti berikut ini:

$$d(D_i, C_j) = \sqrt{(X1_i - X1_j)^2 + (X2_i - X2_j)^2}$$

perhitungan centroid terdekat pada data ke-1, ke-4, dan data ke-10

$$d(1,0) = \sqrt{(0.000000 - 0.000000)^2 + (0.000000 - 0.000000)^2} = 0.000000$$

$$d(1,1) = \sqrt{(0.000000 - 0.056452)^2 + (0.000000 - 0.750000)^2} = 0.752122$$

$$d(1,2) = \sqrt{(0.000000 - 1.000000)^2 + (0.000000 - 1.000000)^2} = 1.414214$$

$$d(4,0) = \sqrt{(0.056452 - 0.000000)^2 + (0.750000 - 0.000000)^2} = 0.752122$$

$$d(4,1) = \sqrt{(0.056452 - 0.056452)^2 + (0.750000 - 0.750000)^2} = 0.000000$$

$$d(4,2) = \sqrt{(0.056452 - 1.000000)^2 + (0.750000 - 1.000000)^2} = 0.976106$$

$$d(10,0) = \sqrt{(1.000000 - 0.000000)^2 + (1.000000 - 0.000000)^2} = 1.414214$$

$$d(10,1) = \sqrt{(1.000000 - 0.056452)^2 + (1.000000 - 0.750000)^2} = 0.976106$$

$$d(10,2) = \sqrt{(1.000000 - 1.000000)^2 + (1.000000 - 1.000000)^2} = 0.000000$$

Setelah melakukan perhitungan, didapat hasil seperti yang terlihat pada tabel 4.

Tabel 4

Iterasi 1

No	Data	Cluster 0	Cluster 1	Cluster 2	Cluster Terpilih
1	D1	0.000000	0.752122	1.414214	0
2	D2	0.253231	0.500260	1.217982	0
3	D3	0.253231	0.500260	1.217982	0
4	D4	0.752122	0.000000	0.976106	1
5	D5	0.522951	0.268077	0.983375	1
6	D6	0.765492	0.096774	0.882908	1
7	D7	0.551992	0.306558	0.914852	1
8	D8	1.048294	0.359301	0.685484	1
9	D9	1.129061	0.530361	0.475806	2
10	D10	1.414214	0.976106	0.000000	2

Setelah perhitungan menggunakan Euclidean Distance selesai, hitung centroid baru dari setiap cluster berdasarkan rata-rata nilai data anggotanya sebelum melanjutkan ke iterasi kedua. Menentukan centroid C0, C1, dan C2 yang baru dari nilai rata-rata anggota yang ada pada masing-masing Cluster.

$$C0 (X1') = (0.000000 + 0.040323 + 0.040323)/3 = 0.026882$$

$$C0 (X2') = (0.000000 + 0.250000 + 0.250000)/3 = 0.166667$$

$$C1 (X1') = (0.056452 + 0.153226 + 0.153226 + 0.233871 + 0.314516)/5 = 0.182258$$

$$C1 (X2') = (0.750000 + 0.500000 + 0.750000 + 0.500000 + 1.000000)/5 = 0.700000$$

$$C2 (X1') = (0.524194 + 1.000000)/2 = 0.762097$$

$$C2 (X2') = (1.000000 + 1.000000)/2 = 1.000000$$

Centroid baru yang diperoleh dapat dilihat pada tabel 5.

Tabel 5

Centroid Baru

Centroid awal	Jumlah pemesanan	Frekuensi pemesanan
Centroid 0	0.026882	0.166667
Centroid 1	0.182258	0.700000
Centroid 2	0.762097	1.000000

Setelah dapat Centroid baru, lakukan proses perhitungan iterasi seperti di awal:

$$d(1,0) = \sqrt{(0.000000 - 0.026882)^2 + (0.000000 - 0.166667)^2} = 0.168821$$

$$d(1,1) = \sqrt{(0.000000 - 0.182258)^2 + (0.000000 - 0.700000)^2} = 0.723338$$

$$d(1,2) = \sqrt{(0.000000 - 0.762097)^2 + (0.000000 - 1.000000)^2} = 1.257295$$

$$d(4,0) = \sqrt{(0.056452 - 0.026882)^2 + (0.750000 - 0.166667)^2} = 0.584082$$

$$d(4,1) = \sqrt{(0.056452 - 0.182258)^2 + (0.750000 - 0.700000)^2} = 0.135378$$

$$d(4,2) = \sqrt{(0.056452 - 0.762097)^2 + (0.750000 - 1.000000)^2} = 0.748622$$

$$d(10,0) = \sqrt{(1.000000 - 0.026882)^2 + (1.000000 - 0.166667)^2} = 1.281173$$

$$d(10,1) = \sqrt{(1.000000 - 0.182258)^2 + (1.000000 - 0.700000)^2} = 0.871035$$

$$d(10,2) = \sqrt{(1.000000 - 0.762097)^2 + (1.000000 - 1.000000)^2} = 0.237903$$

Setelah perhitungan selesai maka akan didapat hasil yang baru dapat dilihat pada Tabel 6.

Tabel 6

Iterasi 2

No	Data	Cluster 0	Cluster 1	Cluster 2	Cluster Terpilih
1	D1	0.168821	0.723338	1.257295	0
2	D2	0.084410	0.471853	1.040893	0
3	D3	0.084410	0.471853	1.040893	0
4	D4	0.584082	0.135378	0.748622	1
5	D5	0.356474	0.202096	0.787860	1
6	D6	0.596859	0.057818	0.658197	1
7	D7	0.392372	0.206552	0.727339	1
8	D8	0.881577	0.327860	0.447581	1
9	D9	0.970445	0.454884	0.237903	2
10	D10	1.281173	0.871035	0.237903	2

Dalam perhitungan K-Means, jika centroid tidak mengalami perubahan atau sudah konvergensi, maka proses perhitungan telah selesai. Karena cluster pada perhitungan pertama dan kedua tidak berubah, maka perhitungan telah dinyatakan selesai.

Pembahasan

Penerapan Python

Tahapan yang akan dilakukan sekarang adalah penerapan. Penerapan yang akan dilakukan penulis menggunakan python dengan aplikasi jupyter notebook.

Input library

Pada tahap awal, dilakukan proses import library Python yang digunakan dalam penelitian.

```
import pandas as pd
import numpy as np
import matplotlib.pyplot as plt

from sklearn.preprocessing import MinMaxScaler
from sklearn.cluster import KMeans
from sklearn.metrics import davies_bouldin_score, silhouette_score
```

Gambar 2. Input Library

Library yang digunakan meliputi pandas untuk pengolahan data, numpy untuk perhitungan numerik, matplotlib untuk visualisasi data, serta library dari scikit-learn yang digunakan dalam proses normalisasi data, K-Means clustering K-Means, dan evaluasi model K-Means clustering.

Membaca data excel

Tahap berikutnya adalah membaca data penelitian yang berasal dari file Excel menggunakan library pandas. Data yang digunakan merupakan data *spare parts* yang telah melalui proses sortir manual sehingga lebih terstruktur dan mudah diproses pada tahap K-Means clustering.

```
file_path = "data CC clean manual.xlsx"

df = pd.read_excel(file_path)

df.head()
```

Gambar 3 Membaca Data Excel

Pengecekan Struktur Data

Pada tahap ini dilakukan pengecekan struktur data untuk mengetahui jumlah baris, nama kolom, serta memastikan data berhasil dibaca dengan baik oleh sistem. Selain itu, dilakukan pengecekan terhadap kategori PR Status untuk memastikan data spare parts dapat dipisahkan dari kategori lainnya seperti consumable, jasa, utility, dan tool.

```
print("Jumlah baris dan kolom:", df.shape)
print("\nNama kolom:")
print(df.columns)

print("\nJumlah data berdasarkan PR Status:")
print(df['Pr Status'].value_counts())
```

Gambar 4 Pengecekan Struktur Data

Filtering Data Sparepart

Tahap filtering dilakukan untuk mengambil data dengan kategori Sparepart saja. Hal ini dilakukan karena penelitian berfokus pada pengelompokan data spare parts, sehingga data selain sparepart tidak digunakan dalam proses K-Means clustering.

```
df_sparepart = df[df['Pr Status'].str.lower().str.strip() == 'sparepart'].copy()

print("Jumlah data sparepart:", df_sparepart.shape[0])
df_sparepart.head()
```

Gambar 5 Filtering Data Sparepart

Cleaning Data

Tahap cleaning data dilakukan untuk membersihkan data dari nilai kosong (missing value), data yang tidak valid, serta memastikan format data numerik dapat diproses dengan baik. Proses ini bertujuan agar hasil K-Means clustering yang diperoleh menjadi lebih akurat dan optimal.

```
df_sparepart = df_sparepart[['Pr No', 'Material', 'Qty Request']].copy()

df_sparepart = df_sparepart.dropna(subset=['Pr No', 'Material', 'Qty Request'])

df_sparepart['Qty Request'] = pd.to_numeric(
    df_sparepart['Qty Request'],
    errors='coerce'
)

df_sparepart = df_sparepart.dropna(subset=['Qty Request'])

df_sparepart['Material'] = df_sparepart['Material'].astype(str).str.strip()
df_sparepart['Pr No'] = df_sparepart['Pr No'].astype(str).str.strip()

print("Jumlah data setelah cleaning:", df_sparepart.shape[0])
df_sparepart.head()
```

Gambar 6. Cleaning Data

Agregasi Data

Tahap agregasi dilakukan untuk mengubah data transaksi menjadi data item spare parts unik. Pada tahap ini dihitung total jumlah pemesanan dan frekuensi pemesanan setiap spare parts berdasarkan nomor PR yang berbeda. Hasil agregasi inilah yang kemudian digunakan sebagai variabel penelitian pada proses K-Means clustering.

```

df_pr = df_sparepart.groupby(['Material', 'Pr No'], as_index=False).agg(
    qty_per_pr=('Qty Request', 'sum')
)

data_cluster = df_pr.groupby('Material', as_index=False).agg(
    jumlah_pemesanan=('qty_per_pr', 'sum'),
    frekuensi_pemesanan=('Pr No', 'nunique')
)

print("Jumlah item sparepart unik:", data_cluster.shape[0])
data_cluster.head()

```

Gambar 7. Agregasi Data

Normalisasi Data

Sebelum proses K-Means clustering dilakukan, data dinormalisasi menggunakan metode Min-Max Scaling. Proses normalisasi bertujuan untuk menyamakan rentang nilai antarvariabel sehingga variabel dengan nilai besar tidak mendominasi perhitungan jarak Euclidean pada algoritma K-Means.

```

X = data_cluster[['jumlah_pemesanan', 'frekuensi_pemesanan']]

scaler = MinMaxScaler()
X_scaled = scaler.fit_transform(X)

data_cluster['jumlah_norm'] = X_scaled[:, 0]
data_cluster['frekuensi_norm'] = X_scaled[:, 1]

data_cluster.head()

```

Gambar 8. Normalisasi Data

Penentuan Jumlah Cluster Menggunakan Metode Elbow

Metode Elbow digunakan untuk menentukan jumlah cluster optimal berdasarkan nilai SSE (Sum of Squared Error). Pada tahap ini dilakukan pengujian beberapa jumlah cluster, kemudian dipilih jumlah cluster yang membentuk pola siku (elbow) pada grafik.

```

sse = []

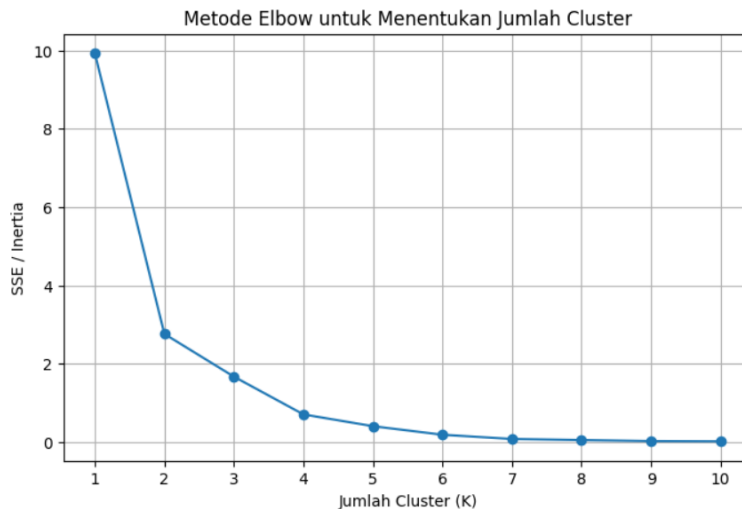
K_range = range(1, 11)

for k in K_range:
    kmeans = KMeans(n_clusters=k, random_state=42, n_init=10)
    kmeans.fit(X_scaled)
    sse.append(kmeans.inertia_)

plt.figure(figsize=(8, 5))
plt.plot(K_range, sse, marker='o')
plt.xlabel('Jumlah Cluster (K)')
plt.ylabel('SSE / Inertia')
plt.title('Metode Elbow untuk Menentukan Jumlah Cluster')
plt.xticks(K_range)
plt.grid(True)
plt.show()

```

Gambar 9. Penentuan Jumlah Cluster Menggunakan Metode Elbow

**Gambar 10.** Metode Elbow

Proses K-Means Clustering

Setelah jumlah cluster ditentukan, dilakukan proses K-Means clustering menggunakan algoritma K-Means dengan jumlah cluster sebanyak 3. Proses ini bertujuan untuk mengelompokkan data *spare parts* berdasarkan kemiripan jumlah dan frekuensi pemesanan.

```
kmeans = KMeans(n_clusters=3, random_state=42, n_init=10)

data_cluster['cluster_awal'] = kmeans.fit_predict(X_scaled)

data_cluster.head()
```

Gambar 11. Proses K-Means Clustering

Pelabelan Cluster

Tahap pelabelan cluster dilakukan untuk memberikan interpretasi terhadap hasil K-Means clustering. Cluster dengan nilai centroid rendah diberi label “Jarang Dipesan”, cluster sedang diberi label “Cukup Sering Dipesan”, dan cluster tertinggi diberi label “Sering Dipesan”.

```
centroid = pd.DataFrame(
    kmeans.cluster_centers_,
    columns=['jumlah_norm_centroid', 'frekuensi_norm_centroid']
)

centroid['cluster_awal'] = centroid.index
centroid['skor'] = centroid['jumlah_norm_centroid'] + centroid['frekuensi_norm_centroid']

centroid_sorted = centroid.sort_values('skor').reset_index(drop=True)

label_map = {
    centroid_sorted.loc[0, 'cluster_awal']: 'C1 - Jarang Dipesan',
    centroid_sorted.loc[1, 'cluster_awal']: 'C2 - Cukup Sering Dipesan',
    centroid_sorted.loc[2, 'cluster_awal']: 'C3 - Sering Dipesan'
}

data_cluster['cluster'] = data_cluster['cluster_awal'].map(label_map)

data_cluster.head()
```

Gambar 12. Pelabelan Cluster

Menampilkan Jumlah Anggota Cluster

Pada tahap ini dilakukan perhitungan jumlah anggota pada masing-masing cluster untuk mengetahui distribusi data *spare parts* pada setiap kelompok hasil K-Means clustering.

```

hasil_cluster = data_cluster['cluster'].value_counts().sort_index()

print("Jumlah anggota setiap cluster:")
print(hasil_cluster)

```

Gambar 13. Menampilkan Jumlah Anggota Cluster

Evaluasi Model Clustering

Tahap evaluasi dilakukan menggunakan metode SSE (Sum of Squared Error), Davies-Bouldin Index (DBI), dan Silhouette Coefficient Index (SCI). Evaluasi ini bertujuan untuk mengukur kualitas cluster yang terbentuk berdasarkan tingkat kedekatan data dan pemisahan antar cluster.

```

SSE = kmeans.inertia_
DBI = davies_bouldin_score(X_scaled, data_cluster['cluster_awal'])
SCI = silhouette_score(X_scaled, data_cluster['cluster_awal'])

print("Hasil Evaluasi K-Means")
print("SSE :", SSE)
print("DBI :", DBI)
print("SCI :", SCI)

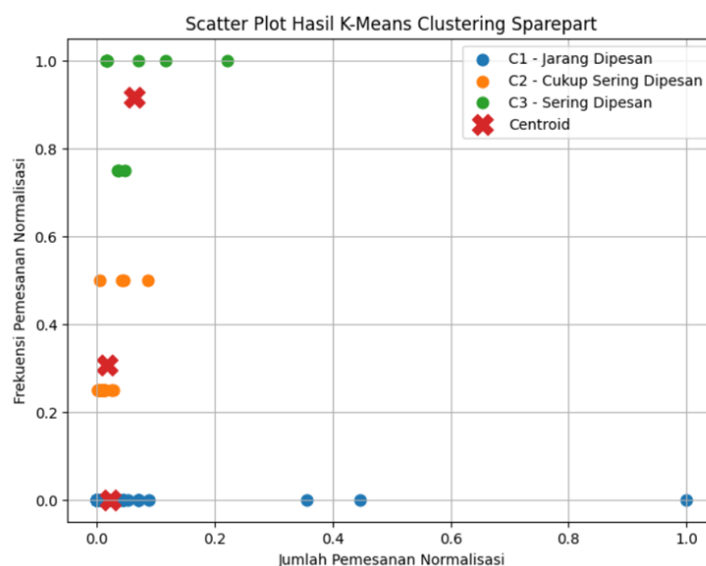
```

Hasil Evaluasi K-Means
 SSE : 1.670368180607603
 DBI : 0.3694095810638318
 SCI : 0.8750592883728601

Gambar 14. Evaluasi Model Clustering

Visualisasi Hasil Clustering Scatter Plot

Visualisasi scatter plot digunakan untuk memperlihatkan persebaran data hasil K-Means clustering berdasarkan variabel jumlah pemesanan dan frekuensi pemesanan yang telah dinormalisasi. Setiap warna menunjukkan cluster yang berbeda, sedangkan centroid ditampilkan sebagai pusat cluster.

**Gambar 15.** Visualisasi Hasil Clustering Scatter Plot

Menampilkan Hasil Akhir Clustering

Tahap ini menampilkan hasil akhir K-Means clustering berupa nama spare parts, jumlah pemesanan, frekuensi pemesanan, hasil normalisasi, dan label cluster yang diperoleh dari proses K-Means Clustering.

```

hasil_akhir = data_cluster[[
    'Material',
    'jumlah_pemesanan',
    'frekuensi_pemesanan',
    'jumlah_norm',
    'frekuensi_norm',
    'cluster'
]]

hasil_akhir.head(188)

```

Gambar 16. Menampilkan Hasil Akhir Clustering

	Material	jumlah_pemesanan	frekuensi_pemesanan	jumlah_norm	frekuensi_norm	cluster
0	AIR PRESSURE CONTROL / GREASE PUMP / PENGATUR ...	1	1	0.000000	0.00	C1 - Jarang Dipesan
1	AVR Merk Stamford	2	1	0.001789	0.00	C1 - Jarang Dipesan
2	Air Circuit Breaker SACE EMAX2 E12N12 1250 A ABB	1	1	0.000000	0.00	C1 - Jarang Dipesan
3	Air Filter	4	2	0.005367	0.25	C2 - Cukup Sering Dipesan
4	Air Filter 25593 Fleetguard	2	1	0.001789	0.00	C1 - Jarang Dipesan
...	...	...	...	...	...	...
183	Wire Rope @285M 26mm 6x36 IWRC RHOL Grade 1960...	1	1	0.000000	0.00	C1 - Jarang Dipesan
184	Wire Rope Boom CC-Nelcon Merk Usha Martin	1	1	0.000000	0.00	C1 - Jarang Dipesan
185	Wirerope 26 mm 6x36WS IWRC x 480 meter RHOL Un...	1	1	0.000000	0.00	C1 - Jarang Dipesan
186	Wirerope Trolley RHOL ø20 mm x 224 mtr IWRC 6 ...	1	1	0.000000	0.00	C1 - Jarang Dipesan
187	Wirerope Trolley RHOL ø20 mm x 330 mtr IWRC 6 ...	1	1	0.000000	0.00	C1 - Jarang Dipesan

Gambar 17. Tampilan Hasil Akhir Clustering

Hasil clustering terhadap 188 item sparepart menunjukkan bahwa Cluster 0 (Jarang Dipesan) merupakan kelompok terbesar dengan 169 item (89,9%), diikuti Cluster 1 (Cukup Sering Dipesan) sebanyak 16 item (8,5%), dan Cluster 2 (Sering Dipesan) sebanyak 3 item (1,6%). Distribusi tersebut menunjukkan bahwa karakteristik pemesanan sparepart pada data yang dianalisis didominasi oleh item dengan jumlah dan frekuensi pemesanan yang relatif rendah, sedangkan hanya sebagian kecil item memiliki karakteristik pemesanan yang lebih tinggi. Dengan demikian, hasil clustering memberikan pemetaan tingkat pemesanan yang lebih terstruktur berdasarkan dua variabel yang digunakan dalam penelitian.

Dari sisi pengelolaan persediaan, hasil pengelompokan tersebut dapat digunakan sebagai informasi pendukung untuk membedakan perhatian terhadap setiap kelompok sparepart. Item pada Cluster 0 dapat dipantau berdasarkan karakteristik pemesanannya yang relatif rendah, sedangkan Cluster 1 memerlukan perhatian lebih karena memiliki tingkat pemesanan yang lebih tinggi. Adapun tiga item pada Cluster 2 menunjukkan karakteristik jumlah dan frekuensi pemesanan tertinggi dibandingkan kelompok lainnya sehingga dapat menjadi perhatian dalam perencanaan ketersediaan sparepart. Meskipun demikian, hasil clustering ini belum dapat secara langsung menentukan jumlah stok, safety stock, reorder point, maupun kebijakan pengadaan karena variabel tersebut tidak dianalisis dalam penelitian.

KESIMPULAN

Berdasarkan hasil penelitian, algoritma K-Means Clustering berhasil mengelompokkan 188 data pemesanan spare parts menjadi tiga cluster, yaitu Cluster 0 (Jarang Dipesan) sebanyak 169 item, Cluster 1 (Cukup Sering Dipesan) sebanyak 16 item, dan Cluster 2 (Sering Dipesan) sebanyak 3 item. Hasil pengelompokan tersebut dapat membantu perusahaan dalam mengidentifikasi tingkat kebutuhan spare parts sehingga pengelolaan persediaan dapat dilakukan dengan lebih efektif. Selain itu, hasil evaluasi menunjukkan nilai SSE sebesar 1,670368, DBI sebesar 0,369410, dan SCI sebesar 0,875059 yang menunjukkan bahwa kualitas K-Means clustering yang dihasilkan sudah baik. Penelitian selanjutnya dapat dilakukan dengan menambahkan variabel lain atau membandingkan metode K-Means dengan algoritma K-Means clustering lainnya untuk memperoleh hasil yang lebih optimal.

UCAPAN TERIMA KASIH

Penulis menyampaikan terima kasih kepada Universitas Islam Negeri Sumatera Utara, PT Prima Multi Peralatan Cabang Belawan, serta seluruh pihak yang telah memberikan dukungan, data, informasi, dan bantuan selama proses pelaksanaan penelitian ini. Apresiasi juga disampaikan kepada pihak-pihak yang telah membantu dalam proses pengumpulan, pengolahan, dan penyelesaian penelitian sehingga artikel ini dapat diselesaikan dengan baik.

PERNYATAAN KONTRIBUSI PENULIS

Muhammad Tsaqif Al Mutawakkil Simbolon dan Sriani berkontribusi dalam pelaksanaan penelitian dan penyusunan artikel, meliputi perumusan penelitian, pengumpulan dan pengolahan data, analisis hasil, penyusunan serta revisi naskah, dan persetujuan terhadap versi akhir artikel untuk dipublikasikan. Kedua penulis bertanggung jawab atas keseluruhan isi artikel.

DAFTAR PUSTAKA

- Abdelmeguid, A., & Corsini, L. (2026). Barriers And Enablers Of Circular Economy In Electrical And Electronic Equipment: A Systematic Literature Review. *Business Strategy And The Environment*.
- Alvianatinova, V., Ali, I., Rahaningsih, N., & Bahtiar, A. (2024). Penerapan Algoritma K-Means K-Means Clustering Dalam Pengelompokan Data Penjualan Supermarket Berdasarkan Cabang (Branch). *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(2), 1529–1535.
- Andrian, D. (2026). Analisis Proses Perawatan Preventif Beberapa Komponen Dalam Sebuah Mesin Dan Pengaruhnya Terhadap Kapasitas Produksi (Studi Kasus: Mesin Kiln Putar Di Industri Ubin Keramik). *Science Tech: Jurnal Ilmu Pengetahuan Dan Teknologi*, 12(2), 182–197.
- Aprilia, A., Kurniati, D., Rayesa, N. F., Arifiyanti, N., Jarlis, N., Hardana, A. E., Berlian, G. O., & Setyaningsih, W. L. (2026). *Manajemen Risiko Agribisnis Terintegrasi: Dari Usahatani, Rantai Pasok, Hingga Resiliensi Sistem Pangan*. PT Ghani Press Group.
- Arizqi, M. R., & Putri, E. P. (2025). Penjadwalan Perawatan Dan Optimalisasi Persediaan Sparepart Mesin Press Batu Tahan Api Guna Memitigasi Akibat Breakdown Pada PT XYZ. *Jurnal Surya Teknika*, 12(1), 127–138. <https://doi.org/10.37859/Jst.V12i1.9202>
- Astrina, F. J., Siswanto, S., & Sartika, D. (2025). Expert System For Diagnosing Eczema At The Pulau Panggung Community Health Centre Using The Forward Chaining Method. *Jurnal Komputer Indonesia*, 4(1), 1–10.
- Ayu, F., Destiwati, F., & Dja'far, H. I. (2026). Implementasi Metode K-Means Clustering Dalam Pengelompokan Produk Terlaris Untuk Optimalisasi Stok Pada Toko Ridho. *Jurnal Rekayasa Komputasi Terapan*, 6(02), 147–153.
- Darmi, Y. D., & Setiawan, A. (2016). Penerapan Metode Clustering K-Means Dalam Pengelompokan Penjualan Produk. *Jurnal Media Infotama*, 12(2).
- Fahri, F., Harahap, B., & Suliawati, S. (2025). Analisis Efektivitas Preventive Maintenance Dengan Metode Periodic Inspection Untuk Meningkatkan Kinerja Pada Unit Wa800-3. *Blend Sains Jurnal Teknik*, 3(3), 246–267.
- Rahayu, Z. K. I., Nurhasanah, D., Haigh, R., Amaratunga, D., H. P., & Wahdiny, I. I. (2024). Unveiling Transboundary Challenges In River Flood Risk Management: Learning From The Ciliwung River Basin. *Natural Hazards And Earth System Sciences*, 24(6), 2045–2064. <https://doi.org/10.5194/Nhess-24-2045-2024>
- Ramadani, A., Zainuri, F., & Yusyama, A. Y. (2025). Perencanaan Persediaan Stok Sparepart Menggunakan Metode Min–Max Untuk Pemeliharaan Berkala Di PT X. *Prosiding Seminar Nasional Teknik Mesin*, 1, 786–790.
- Ramdhan, D., Dwilestari, G., Dana, R. D., & Ajiz, A. (2022). Clustering Data Persediaan Barang Dengan Menggunakan Metode K-Means. *Means (Media Informasi Analisa Dan Sistem)*, 1–9.
- Shepyantoni, F., Kanedi, I., & Suryana, E. (2024). Penerapan Metode K-Means Clustering Dalam Pengelompokan Data Pasien Rawat Inap Peserta BPJS Di Rumah Sakit Umum Daerah Kabupaten Kaur. *Jurnal Media Infotama*, 20(2), 493–500.
- Sitinjak, J. M., Sari, K., & Yetri, M. (2024). Penerapan Data Mining Dalam Penjualan Sparepart Motor Dengan Menggunakan Metode K-Means Clustering. *Jurnal Sistem Informasi TGD*, 3(6), 862–

871.

- Suryadi, U. T., & Supriatna, Y. (2019). Sistem Clustering Tindak Kejahatan Pencurian Di Wilayah Jawa Barat Menggunakan Algoritma K-Means. *Jurnal Teknologi Informasi Dan Komunikasi*, 14(1), 15–27.
- Wananda, T. T., & Nugraha, J. (2023). Pengelompokan Peserta Program Merdeka Belajar-Kampus Merdeka (MBKM) Berdasarkan Perguruan Tinggi Dengan Analisis Cluster Non Hierarki Metode K-Means. *Emerging Statistics And Data Science Journal*, 1(2), 260–267.
- Wibowo, R. P., & Elisa, E. (2026). Penerapan Algoritma K-Means Untuk Pengelompokan Pada Data Supermarket Menggunakan Pemograman Python. *Computer And Science Industrial Engineering (COMASIE)*, 14(1), 90–100.
- Widari, N. S., Andrian, D., & Tentua, A. P. D. (2025). Penentuan Interval Waktu Optimal Perawatan Preventif Komponen Kritis Di Mesin Kiln Putar Industri Ubin Keramik. *Journal Of Industrial View*, 7(1), 76–83. <https://doi.org/10.26905/Jiv.V7i1.15056>
- Wohon, K. G., Purba, A. A., & Endrawati, B. F. (2023). Penjadwalan Perawatan Sparepart Mesin Dengan Pendekatan Reliability Centered Maintenance Dan Failure Mode Effect Analysis Di PT ABC. *Jurnal Teknik Industri*, 13(3), 183–188.
- Zhang, S., Huang, K., & Yuan, Y. (2021). Spare Parts Inventory Management: A Literature Review. *Sustainability*, 13(5), 2460. <https://doi.org/10.3390/Su13052460>