

Pengklasifikasian Tingkat Penjualan Sparepart Mobil di Putra Motor Menggunakan Algoritma K-Means Clustering

Siti Nurhalimah^{1*}, Sri Wahyuni², Wiwin Handoko³

^{1,2,3}Fakultas Ilmu Komputer, Program Studi Sistem Informasi, Universitas Royal

^{1*} sitinurhalimah8304@gmail.com, ² sriwahyuni233322@gmail.com,
³ win.van.handoko@gmail.com

Abstrak

K-Means adalah salah satu metode pengelompokan data non-hierarki (partisi) yang bertujuan untuk membagi data ke dalam dua atau lebih kelompok berdasarkan kesamaan tertentu. Dalam penelitian ini, analisis data mining dilakukan menggunakan metode K-Means. Metode ini memungkinkan data yang telah diperoleh dikelompokkan ke dalam beberapa cluster berdasarkan kesamaan atau kemiripan antar data. Data dengan karakteristik yang serupa akan dikelompokkan dalam cluster yang sama, sementara data dengan karakteristik berbeda akan masuk ke cluster lain. Data dengan karakteristik yang serupa akan dikelompokkan dalam cluster yang sama, sementara data dengan karakteristik berbeda akan masuk ke cluster lain. Pembuatan model clustering dilakukan dengan menggunakan K-Means Clustering untuk mengelompokkan tingkat penjualan sparepart mobil pada Putra Motor berdasarkan jumlah total dan total penjualan dari bulan Januari sampai bulan Juni 2024. Algoritma ini diterapkan melalui software Jupyter Notebook, menggunakan data yang telah diambil dari Putra Motor. Proses dimulai dengan menjalankan Jupyter Notebook, kemudian membuat file baru dengan format ipynb dan mengaktifkan fungsi-fungsi yang terdapat didalam python. Hasil analisis menunjukkan bahwa terdapat 62 data dalam kategori penjualan sedang, 31 dalam rendah, dan 17 dalam tinggi. Temuan ini memberikan wawasan penting bagi Putra Motor untuk meningkatkan strategi pemasaran dan manajemen persediaan. Dengan sistem yang lebih efisien, perusahaan dapat mengurangi modal yang terjebak dalam persediaan dan meningkatkan profitabilitas. Penelitian ini juga menekankan pentingnya penerapan teknologi dalam analisis data untuk pengambilan keputusan yang lebih baik di sektor otomotif.

Kata Kunci : Penjualan, K-Means Clustering, Sparepart, Data Mining, Industri Otomotif

Abstract

K-Means is one of the non-hierarchical (partitioning) clustering methods that aims to divide data into two or more groups based on certain similarities. In this study, data mining analysis was conducted using the K-Means method. This method enables the collected data to be grouped into several clusters based on their similarities or similarities among the data. Data with similar characteristics are grouped into the same cluster, while data with different characteristics are placed in other clusters. The clustering model was developed using K-Means clustering to categorize the sales levels of automobile spare parts at Putra Motor, based on the total quantity and total sales from January to June 2024. The algorithm was implemented through Jupyter Notebook software, utilizing data obtained from Putra Motor. The process began by running Jupyter Notebook, creating a new file in .ipynb format, and activating the functions available in Python. The analysis results showed that there were 62 data entries in the medium sales category, 31 in the low category, and 17 in the high category. These findings provide valuable insights for Putra Motor to enhance its marketing strategies and inventory management. With a more efficient system, the company can reduce capital tied up in inventory and increase profitability. This study also highlights the importance of applying technology in data analysis to support better decision-making in the automotive sector.

Keyword : Sales, K-Means Clustering, Spare Parts, Data Mining, Automotive Industry

1. PENDAHULUAN

Pesatnya kemajuan teknologi komputer dan komunikasi telah membawa masyarakat memasuki era informasi, di mana perkembangan database yang semakin canggih dan meningkatnya penggunaan aplikasi berbasis data memainkan peran penting [1]. Saat ini, proses penjualan mengalami kemajuan yang sangat signifikan seiring dengan pesatnya perkembangan teknologi. Beragam cara dan metode diterapkan untuk meningkatkan penjualan produk, termasuk dalam penjualan suku cadang mobil. Dengan kemajuan teknologi yang terus berkembang, para penjual dituntut untuk beradaptasi dan berpartisipasi dalam perubahan tersebut [2]. Penjualan merupakan kegiatan yang melengkapi proses pembelian sehingga memungkinkan terjadinya suatu transaksi [3]. Begitu juga yang dibutuhkan pada Putra Motor

Sparepart adalah barang yang terdiri dari berbagai komponen yang membentuk satu kesatuan dan memiliki fungsi tertentu[4]. Sparepart banyak digunakan pada berbagai jenis kendaraan, sehingga memiliki beragam jenis sesuai dengan kebutuhan dan spesifikasi kendaraan tersebut[5]. Proses produksi mobil yang berkualitas membutuhkan penggunaan suku cadang yang sesuai dan tepat [6]. Sparepart pada kendaraan mobil berperan penting dalam menjaga kinerja dan fungsi kendaraan. Beberapa contoh sparepart mobil meliputi baterai, kampas rem, packing drain, busi, suspensi, filter oli, dan lainnya. Untuk memastikan kendaraan tetap bekerja secara optimal, diperlukan penggunaan sparepart yang berkualitas serta perawatan rutin yang teratur [7].

Putra Motor adalah perusahaan yang bergerak di bidang otomotif, melayani pembelian, penjualan suku cadang mobil, serta menyediakan layanan servis untuk berbagai merek mobil [8]. Namun, perusahaan menghadapi kendala dalam memantau produk yang dijual, mengidentifikasi kebutuhan konsumen, dan mengelola data secara efisien. Hal ini menyebabkan modal yang terlalu tinggi dan penumpukan barang, yang pada akhirnya berdampak pada defisit keuangan perusahaan [9]. Oleh karena itu, diperlukan sebuah sistem yang mampu mendukung pengambilan keputusan dengan cepat dan akurat serta dapat mengklasifikasikan tingkat penjualan menggunakan algoritma K-Means Clustering [10].

K-Means adalah salah satu metode pengelompokan data non-hierarki (partisi) yang bertujuan untuk membagi data ke dalam dua atau lebih kelompok berdasarkan kesamaan tertentu [11]. Dalam penelitian ini, analisis data mining dilakukan menggunakan metode K-Means. Metode ini memungkinkan data yang telah diperoleh dikelompokkan ke dalam beberapa cluster berdasarkan kesamaan atau kemiripan antar data. Data dengan karakteristik yang serupa akan dikelompokkan dalam cluster yang sama, sementara data dengan karakteristik berbeda akan masuk ke cluster lain [12]. Penyesuaian parameter ini dilakukan berdasarkan objek penelitian yang ada, baik secara manual maupun dengan bantuan perangkat lunak [13].

Dari permasalahan yang sudah dijelaskan di atas, maka penulis tertarik melakukan penelitian dengan judul “Pengklasifikasian Tingkat Penjualan Sparepart Mobil di Putra Motor Menggunakan Algoritma K-Means”.

2. METODE

Langkah-langkah penelitian dalam pengelompokkan data penjualan *sparepart* mobil di Putra Motor meliputi beberapa bagian penting sebagai berikut:

a. Teknik Pengumpulan Data (*Data Collecting*)

Pengumpulan data dilakukan untuk memperoleh informasi yang relevan dengan penelitian. Teknik yang digunakan meliputi:

1. Langsung (Observasi): Melakukan pengamatan langsung terhadap proses penjualan *sparepart* mobil untuk mendapatkan data yang akurat.

2. Wawancara (*Interview*): Metode pengumpulan data dilakukan melalui pengamatan langsung terhadap sumber masalah serta berinteraksi secara langsung dengan pihak terkait, yaitu Pemilik Putra Motor.
- b. Studi Kepustakaan (*Study of Literature*)
Melakukan kajian terhadap sumber-sumber literatur, buku, jurnal, dan referensi ilmiah lainnya untuk mendukung teori dan konsep yang digunakan dalam penelitian.
- c. Penerapan Metode *K-Means* dalam pengolahan data
Memanfaatkan algoritma *K-Means* untuk menganalisis data penjualan *sparepart* mobil dan mengelompokkannya berdasarkan pola atau karakteristik tertentu. [14].

2.2 Data Mining

Data mining adalah bagian dari tahapan dalam proses *Knowledge Discovery in Database* (KDD). Melalui data mining, kita dapat melakukan pengklasifikasian, prediksi, estimasi, dan memperoleh informasi lain yang berguna dari kumpulan data yang besar. Data mining adalah proses untuk mencari pola atau informasi menarik dalam data yang dipilih menggunakan berbagai teknik atau metode tertentu [15].

2.3 Clustering

Clustering adalah proses pengelompokan record, observasi, atau kelas yang memiliki kesamaan objek. *Clustering* sering digunakan sebagai langkah awal dalam metode data mining. Banyak algoritma *clustering* yang telah digunakan oleh peneliti sebelumnya, di antaranya *K-Means*, *Improved K-Means*, *Fuzzy C-Means*, *DBSCAN*, *K-Medoids (PAM)*, *CLARANS*, dan *Fuzzy Subtractive* [16].

2.4 Metode *K-Means*

Algoritma *K-Means* adalah salah satu metode pengelompokan data non-hierarki (partisi) yang bertujuan untuk membagi data ke dalam dua atau lebih kelompok berdasarkan kesamaan tertentu [11]. Dalam penelitian ini, analisis data mining dilakukan menggunakan metode *K-Means*. Metode ini memungkinkan data yang telah diperoleh dikelompokkan ke dalam beberapa *cluster* berdasarkan kesamaan atau kemiripan antar data. Data dengan karakteristik yang serupa akan dikelompokkan dalam *cluster* yang sama, sementara data dengan karakteristik berbeda akan masuk ke *cluster* lain [12]. Langkah-langkah dalam melakukan *clustering* menggunakan metode *K-Means* adalah sebagai berikut [15]:

1. Pilih jumlah *cluster* k .
2. Inisialisasi ke pusat *cluster* ini bisa dilakukan dengan berbagai cara. Cara yang paling sering dilakukan adalah dengan random atau acak. Pusat-pusat *cluster* diberi dengan nilai awal dengan angka-angka *random*.

$$d(x_i, \mu_j) = \sqrt{\sum (x_i - \mu_j)^2}$$

Dimana:

x_i = data kriteria

μ_j = *centroid* pada *cluster* ke- j s

3. Klasterisasi penjualan setiap data berdasarkan kedekatannya dengan *centroid* atau mencari jarak terkecil.
4. Memperbaharui nilai *centroid* baru, nilai *centroid* baru di peroleh dari rata-rata *cluster* yang bersangkutan dengan menggunakan rumus yaitu :

$$\mu_j(t+1) = \frac{1}{N_{sj}} \sum_{j \in s_j} x_j$$

Keterangan :

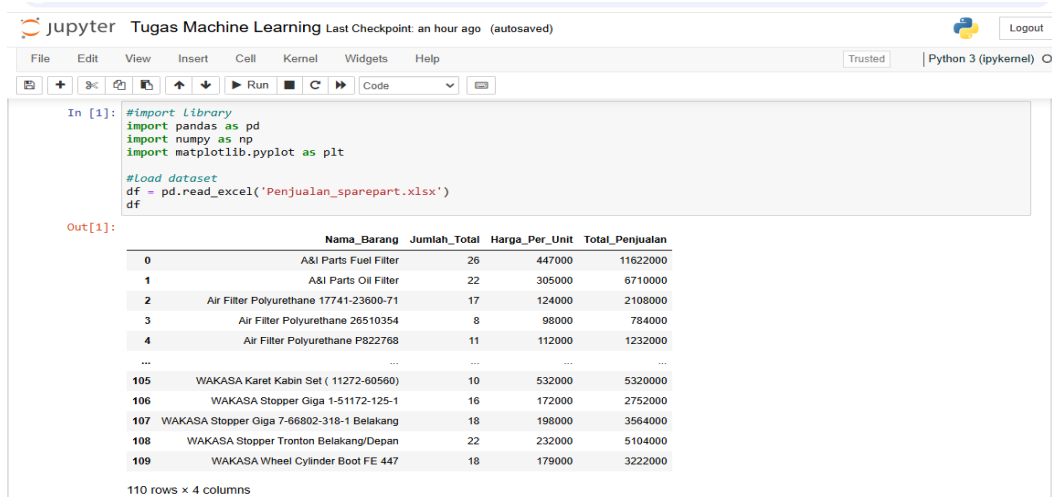
$\mu_j(t+1)$ = *centroid* baru pada iterasi $(t+1)$ N_{sj} = Data pada *cluster* S_j

5. Apabila data setiap *cluster* belum berhenti, lakukan perulangan dari langkah 2 hingga 5, sampai anggota tiap *cluster* tidak ada yang berubah.

Dengan mengikuti langkah-langkah di atas, kita dapat melakukan pengelompokan data menggunakan metode *K-Means* hingga mendapatkan hasil *clustering* yang optimal.

3. HASIL DAN PEMBAHASAN

Pembuatan model clustering dilakukan dengan menggunakan *K-Means Clustering* untuk mengelompokkan tingkat penjualan *sparepart* mobil pada Putra Motor berdasarkan jumlah total dan total penjualan dari bulan januari sampai bulan juni 2024. Algoritma ini diterapkan melalui *software Jupyter Notebook*, menggunakan data yang telah diambil dari Putra Motor. Proses dimulai dengan menjalankan *Jupyter Notebook*, kemudian membuat file baru dengan format *ipnyb* dan mengaktifkan fungsi-fungsi yang terdapat didalam *python* seperti mengaktifkan fungsi *pandas*, *numpy*, *seaborn* dan *matplotlib* pada *python*.



Gambar 1. Impor Data

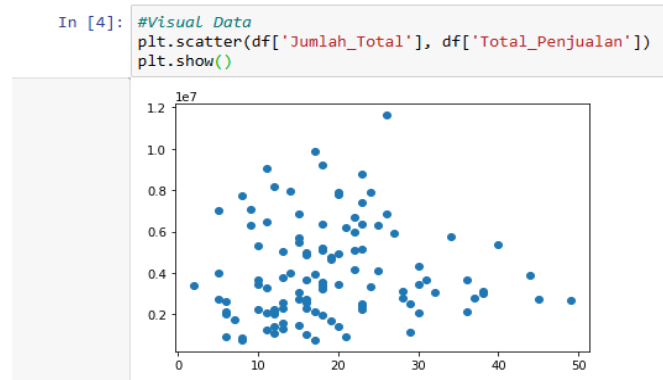
Setelah melakukan import data, masukan fungsi *df.info()* dalam pustaka *pandas* digunakan untuk menampilkan informasi ringkas mengenai struktur *DataFrame*. Informasi ini sangat penting dalam proses eksplorasi data awal (*exploratory data analysis*), khususnya tipe data setiap kolom, serta mendeteksi keberadaan nilai kosong (*null values*).

```
In [2]: #Melihat info data
df.info()

<class 'pandas.core.frame.DataFrame'>
RangeIndex: 110 entries, 0 to 109
Data columns (total 4 columns):
#   Column             Non-Null Count  Dtype  
---  -
0   Nama_Barang         110 non-null    object  
1   Jumlah_Total        110 non-null    int64   
2   Harga_Per_Unit       110 non-null    int64   
3   Total_Penjualan     110 non-null    int64   
dtypes: int64(3), object(1)
memory usage: 3.6+ KB
```

Gambar 2. Melihat Info Data

Selanjutnya untuk melihat visualisasi data bisa menggunakan kode *scatter*, kode tersebut menggunakan fungsi *scatter()* dari pustaka *Matplotlib* untuk membuat *scatter plot* yang memvisualisasikan hubungan antara jumlah total barang yang terjual (*Jumlah_Total*) dan total nilai penjualan (*Total_Penjualan*). *Scatter plot* ini membantu mengidentifikasi pola atau hubungan antara kedua variabel, misalnya apakah terdapat korelasi positif, negatif, atau tidak ada hubungan sama sekali.



Gambar 3. Visualisasi Data

Selanjutnya melakukan *cluster* dengan menggunakan pustaka *sklearn* untuk mengelompokkan dataset ke dalam tiga *cluster*. Parameter *n_clusters=3* menentukan jumlah *cluster* yang diinginkan, sedangkan *random_state=0* digunakan untuk memastikan konsistensi hasil setiap kali algoritma dijalankan.

```
In [5]: #menjalankan k-means ke dataset
from sklearn.cluster import KMeans

km = KMeans(n_clusters=3, random_state=0)
km

Out[5]: KMeans(n_clusters=3, random_state=0)
```

Gambar 4. Melakukan *Cluster*

Selanjutnya lakukan proses penerapan algoritma *K-Means* untuk mengelompokkan data berdasarkan "Jumlah_Total" dan "Total_Penjualan". Hasil prediksi *cluster* disimpan dalam variabel *y_prediksi* dan ditambahkan sebagai kolom baru bernama "*cluster*" ke dalam *dataframe* *df*. Dengan demikian, setiap data memiliki informasi *cluster* yang mengelompokkan berdasarkan kesamaan pola pada variabel tersebut, yang dapat digunakan untuk analisis lebih lanjut seperti segmentasi atau pengelompokan produk.

```
In [6]: #Memprediksi cluster ke dataset
y_prediksi = km.fit_predict(df[['Jumlah_Total', 'Total_Penjualan']])
y_prediksi

Out[6]: array([0, 0, 1, 1, 1, 1, 1, 2, 1, 2, 1, 1, 1, 2, 1, 1, 1, 1, 1, 0, 2,
                2, 1, 1, 1, 1, 2, 1, 1, 1, 1, 2, 2, 1, 1, 1, 1, 1, 1, 2, 1, 1,
                1, 1, 1, 1, 1, 1, 2, 2, 1, 2, 2, 0, 0, 2, 0, 2, 2, 0, 0, 0, 1,
                1, 0, 1, 2, 1, 2, 1, 1, 1, 1, 2, 2, 1, 1, 1, 1, 1, 1, 1, 2, 1,
                0, 0, 2, 0, 0, 1, 2, 0, 0, 2, 2, 2, 2, 1, 2, 1, 1, 1, 2, 1])
```

Gambar 5. Memprediksi *cluster* ke dataset

Kemudian membuat visualisasi hasil clustering dengan diagram sebar, di mana data dibagi menjadi tiga kluster (0, 1, dan 2) berdasarkan "Jumlah_Total" dan "Total_Penjualan". Setiap kluster ditampilkan dengan warna berbeda untuk menunjukkan pola distribusi dan pengelompokan data.

```
In [7]: df['cluster'] = y_prediksi
df
```

```
Out[7]:
```

	Nama_Barang	Jumlah_Total	Harga_Per_Unit	Total_Penjualan	cluster
0	A&I Parts Fuel Filter	26	447000	11622000	0
1	A&I Parts Oil Filter	22	305000	6710000	0
2	Air Filter Polyurethane 17741-23600-71	17	124000	2108000	1
3	Air Filter Polyurethane 26510354	8	98000	784000	1
4	Air Filter Polyurethane P822768	11	112000	1232000	1
...	...	...	...	...	...
105	WAKASA Karet Kabin Set (11272-80560)	10	532000	5320000	2
106	WAKASA Stopper Giga 1-51172-125-1	16	172000	2752000	1
107	WAKASA Stopper Giga 7-66802-318-1 Belakang	18	198000	3564000	1
108	WAKASA Stopper Tronton Belakang/Depan	22	232000	5104000	2
109	WAKASA Wheel Cylinder Boot FE 447	18	179000	3222000	1

110 rows x 5 columns

Gambar 6. Membuat Diagram sebar

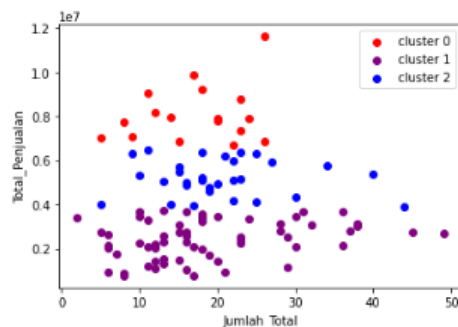
Untuk membuat visualisasi hasil *clustering* dengan diagram sebar, dan membagi data ke dalam tiga kluster (0, 1, dan 2) berdasarkan "Jumlah_Total" dan "Total_Penjualan". Setiap *cluster* diberi warna berbeda (merah, ungu, biru) untuk menunjukkan pola distribusi dan pengelompokan data.

```
In [8]: #visualisasi hasil cluster

df0 = df[df.cluster == 0]
df1 = df[df.cluster == 1]
df2 = df[df.cluster == 2]

plt.scatter(df0.Jumlah_Total, df0['Total_Penjualan'], color='red', label = 'cluster 0')
plt.scatter(df1.Jumlah_Total, df1['Total_Penjualan'], color='purple', label = 'cluster 1')
plt.scatter(df2.Jumlah_Total, df2['Total_Penjualan'], color='blue', label = 'cluster 2')

plt.xlabel('Jumlah_Total')
plt.ylabel('Total_Penjualan')
plt.legend()
plt.show()
```



Gambar 7. Pemberian warna pada setiap *cluster*

Selanjutnya setiap *cluster* dapat diubah dengan pilihan “Penjualan Tinggi”, “Penjualan Sedang” dan “Penjualan Rendah”.

```
In [9]: condition = [
        (df['cluster']==0),
        (df['cluster']==1),
        (df['cluster']==2)]
choice = ['Penjualan Tinggi', 'Penjualan Sedang', 'Penjualan Rendah']
df['cluster'] = np.select(condition, choice)
df
```

Out[9]:

	Nama_Barang	Jumlah_Total	Harga_Per_Unit	Total_Penjualan	cluster
0	A&I Parts Fuel Filter	26	447000	11622000	Penjualan Tinggi
1	A&I Parts Oil Filter	22	305000	6710000	Penjualan Tinggi
2	Air Filter Polyurethane 17741-23800-71	17	124000	2108000	Penjualan Sedang
3	Air Filter Polyurethane 26510354	8	98000	784000	Penjualan Sedang
4	Air Filter Polyurethane P822768	11	112000	1232000	Penjualan Sedang
...	...	...	...	...	...
105	WAKASA Karet Kabin Set (11272-80580)	10	532000	5320000	Penjualan Rendah
106	WAKASA Stopper Giga 1-51172-125-1	16	172000	2752000	Penjualan Sedang
107	WAKASA Stopper Giga 7-66802-318-1 Belakang	18	198000	3564000	Penjualan Sedang
108	WAKASA Stopper Tronton Belakang/Depan	22	232000	5104000	Penjualan Rendah
109	WAKASA Wheel Cylinder Boot FE 447	18	179000	3222000	Penjualan Sedang

110 rows x 5 columns

Gambar 8. Penentuan jenis *cluster*

Output ini menunjukkan distribusi data dalam kolom '*cluster*', yang mengkategorikan data berdasarkan tingkat penjualan (Sedang, Rendah, Tinggi) dan menghitung jumlah masing-masing *cluster*.

```
In [10]: df['cluster'].value_counts()

Out[10]: Penjualan Sedang    62
          Penjualan Rendah   31
          Penjualan Tinggi   17
          Name: cluster, dtype: int64
```

Gambar 9. Menghitung jumlah setiap *cluster*

Dengan kata lain, dari data yang dianalisis, terdapat 62 data dengan kategori "Penjualan Sedang", 31 data dengan kategori "Penjualan Rendah", dan 17 data dengan kategori "Penjualan Tinggi". Informasi ini sangat berguna untuk memahami komposisi data dan melihat proporsi setiap kategori.

Table 1: Jumlah Setiap *Cluster*

<i>Jenis Cluster</i>	<i>Jumlah Cluster</i>
Penjualan Tinggi	17
Penjualan Sedang	62
Penjualan Rendah	31

Kemudian menampilkan isi dari *DataFrame* *df_tinggi*, yang berisi data penjualan untuk barang-barang dengan kategori "Penjualan Tinggi". Ini adalah teknik umum dalam analisis data untuk menyaring data berdasarkan kriteria tertentu.

```
In [11]: #menampilkan cluster penjualan
df_tinggi = df[df["cluster"] == "Penjualan Tinggi"]
df_tinggi
```

	Nama_Barang	Jumlah_Total	Harga_Per_Unit	Total_Penjualan	cluster
0	A&I Parts Fuel Filter	26	447000	11622000	Penjualan Tinggi
1	A&I Parts Oil Filter	22	305000	6710000	Penjualan Tinggi
20	Compressor FN 527 MH 031250 44cm	24	330000	7920000	Penjualan Tinggi
56	ISUZU Shock Abs Front (for Isuzu NMR71)	23	381000	8763000	Penjualan Tinggi
57	ISUZU Zippo Lighter 3	15	457000	6855000	Penjualan Tinggi
59	Kampas Rem Depan AN-549WKI	20	390000	7800000	Penjualan Tinggi
62	Medium Filter MPF-287/592/48-8GN2G	14	570000	7980000	Penjualan Tinggi
63	Medium Filter MPF-592/592/48-8GN2G	11	825000	9075000	Penjualan Tinggi
64	Medium Filter PF-592/592/150-6GNNG	5	1400000	7000000	Penjualan Tinggi
67	Mur As Roda M45 H58	8	985000	7720000	Penjualan Tinggi
88	Ring Gardan 8DC 9 MC 814670	18	512000	9216000	Penjualan Tinggi
89	Seal Oil Trunion (1-51389-005-0)	26	263000	6838000	Penjualan Tinggi

Gambar 10. Penjualan Tinggi

4. KESIMPULAN

Penelitian ini bertujuan untuk mengelompokkan data penjualan dari Januari hingga Juni 2024 ke dalam tiga kategori, yaitu penjualan tinggi, penjualan sedang, dan penjualan rendah. Dengan menggunakan metode K-Means Clustering, data diolah untuk mengidentifikasi pola penjualan berdasarkan jumlah barang terjual dan nilai penjualan.

Hasil analisis menunjukkan bahwa terdapat 62 data dalam kategori penjualan sedang, 31 data dalam kategori penjualan rendah, dan 17 data dalam kategori penjualan tinggi. Temuan ini memberikan wawasan penting bagi Putra Motor dalam meningkatkan strategi pemasaran dan manajemen persediaan. Dengan sistem yang lebih efisien, perusahaan dapat mengurangi modal yang terikat dalam persediaan serta meningkatkan profitabilitas.

Selain itu, penelitian ini juga menekankan pentingnya penerapan teknologi dalam analisis data guna mendukung pengambilan keputusan yang lebih baik di sektor otomotif.

5. DAFTAR PUSTAKA

- [1] N. Ananda and R. Amalia Aras, "Seminar Nasional Sains dan Teknologi Informasi (SENSASI) Clustering Pengeluaran Tahunan Berbagai Macam Produk Menggunakan Metode K-Means," Agustus, pp. 143-147, 2021.
- [2] A. R. Sinaga and G. D. Pranata, "SisInfo Clustering Data Penjualan Produk pada Toko Yudha dengan Algoritma K-Means Kata Kunci: Data Mining, Naive Bayes, Prestasi. SisInfo," vol. 3, no. 02, pp. 135-139, 2021.
- [3] Y. Anggraeni and P. Handayani, "Penerapan Metode K-Means Untuk Menentukan Penjualan Tiket Renang Pada Splash Swimming Pool & Gym," J. Inform. dan Teknol. Inf., vol. 2, no. 1, pp. 173-181, 2023, doi: 10.56854/jt.v2i1.167.
- [4] J. M. Sitinjak, K. Sari, and M. Yetri, "Penerapan Data Mining Dalam Penjualan Sparepart Motor Dengan Menggunakan Metode K-Means Clustering," vol. 3, no. November, pp. 862-871, 2024.
- [5] B. Supriyadi, "Perancangan Program Penjualan Sparepart Mobil," IMTechno J. Ind. Manag. Technol., vol. 1, no. 1, pp. 9-18, 2020, [Online]. Available: <http://jurnal.bsi.ac.id/index.php/imtechno>
- [6] S. Abdy, E. R. Br Gultom, S. Ramadhany, and A. Afifudin, "Prediksi Penjualan Sparepart Mobil Terlaris Menggunakan Metode K-Nearest Neighbor," JURIKOM (Jurnal Ris. Komputer), vol. 9, no. 6, p. 2003, 2022, doi: 10.30865/jurikom.v9i6.5189.

- [7] S. Usman, "Predictive Sparepart Maintenance Menggunakan Algoritma Machine Learning Extreme Gradient Boosting Regressor," *J. Syst. Comput. Eng.*, vol. 5, no. 2, pp. 249–258, 2024, doi: 10.61628/jsce.v5i2.1418.
- [8] A. O. Br Ginting, "Penerapan Data Mining Korelasi Penjualan Spare Part Mobil Menggunakan Metode Algoritma Apriori (Studi Kasus: CV. Citra Kencana Mobil)," *J. Inf. Technol.*, vol. 1, no. 2, pp. 70–77, 2021, doi: 10.32938/jitu.v1i2.1472.
- [9] A. Julianto and S. Andayani, "Penerapan Data Mining Untuk Klasifikasi Produk Terlaris Menggunakan Algoritma Naive Bayes Pada Bengkel Motor," *JuSiTik J. Sist. dan Teknol. Inf. Komun.*, vol. 7, no. 2, pp. 50–58, 2024, doi: 10.32524/jusitik.v7i2.1148.
- [10] F. F. T. Kesuma and S. P. Tamba, "Penerapan Data Mining Untuk Menentukan Penjualan Sparepart Toyota Dengan Metode K-Means Clustering," *J. Sist. Inf. dan Ilmu Komput. Prima (JUSIKOM PRIMA)*, vol. 2, no. 2, pp. 67–72, 2020, doi: 10.34012/jusikom.v2i2.376.
- [11] A. Torence, M. Ramadhan, and E. F. Ginting, "Penerapan Data Mining Menggunakan Algoritma K-Means Clustering Dalam Pengelompokkan Data Penerima Vaksinasi Covid-19," *J. Sist. Inf. Triguna Dharma (JURSI TGD)*, vol. 2, no. 3, p. 482, 2023, doi: 10.53513/jursi.v2i3.6829.
- [12] S. Nurani, Y. Syahra, and A. Calam, "Penerapan Data Mining Dalam Clustering Pencapaian Target Penjualan Menggunakan Algoritma K-Means," *J. Sist. Inf. Triguna Dharma (JURSI TGD)*, vol. 2, no. 3, p. 355, 2023, doi: 10.53513/jursi.v2i3.6552.
- [13] F. Indriyani and E. Irfiani, "Clustering Data Penjualan pada Toko Perlengkapan Outdoor Menggunakan Metode K-Means," *JUITA J. Inform.*, vol. 7, no. 2, p. 109, 2021, doi: 10.30595/juita.v7i2.5529.
- [14] N. Ghina et al., "Implementasi Algoritma K-Means Clustering Untuk Mengetahui Minat Pembeli di Agen Buah Melon Yudi," *J. Inform. dan Rekayasa Komputer (JAKAKOM)*, vol. 2, no. 2, pp. 254–262, 2022, doi: 10.33998/jakakom.2022.2.2.116.
- [15] C. Alvian, F. Taufik, and A. Alhafiz, "Penerapan Data Mining Untuk Evaluasi Data Penjualan Sparepart Motor Menggunakan Metode K-Means Pada Bengkel Usaha Baru," vol. 3, pp. 472–482, 2024.
- [16] P. Metode, K. Penjualan, and P. S. Com, "Penerapan Metode K-Means dalam Penjualan Produk Souq.Com," vol. 5, no. September 2021, 2022.