

Penerapan Metode *K-Means Clustering* Untuk Menentukan Prioritas Persediaan Barang

DOI: <http://dx.doi.org/10.35889/jutisi.v15i1.3417>

Creative Commons License 4.0 (CC BY – NC)

Firza Andra Muharam^{1*}, Dien Novita²

Sistem Informasi, Universitas Multi Data Palembang, Palembang, Indonesia

*e-mail Corresponding Author: firzaandra08@mhs.mdp.ac.id

Abstract

Tani Mandiri is a trading company operating in the agricultural sector that provides various types of fertilizers and agricultural equipment. The main problem faced by the company is suboptimal inventory management, which can potentially lead to inventory shortages or excess stock, thereby affecting the effectiveness of business operations. This study aims to determine inventory priorities at the Tani Mandiri Store by applying data mining techniques. The method used is K-means clustering, with attributes including item name and inventory quantity based on data from 2022 to 2025. The research stages follow the Knowledge Discovery in Databases (KDD) process, which includes data selection, data cleaning, data transformation, and clustering. The optimal number of clusters was determined using the Elbow method. The results indicate that the best solution consists of three clusters: low, medium, and high priority. Out of a total of 220 types of items, 139 were classified as low priority, 22 as medium priority, and 59 as high priority.

Keywords: Data Mining; Inventory; K-Means Clustering

Abstrak

Tani Mandiri adalah usaha dagang yang bergerak pada sektor pertanian dengan menyediakan beragam jenis pupuk serta perlengkapan pertanian. Permasalahan utama yang dihadapi perusahaan ini adalah pengelolaan persediaan yang belum dilakukan secara optimal, sehingga berpotensi menimbulkan kondisi kekurangan maupun kelebihan persediaan barang yang berdampak pada efektivitas operasional usaha. Penelitian ini bertujuan untuk mengelompokkan prioritas persediaan barang pada Toko Tani Mandiri melalui penerapan teknik *data mining*. Metode yang digunakan yaitu *k-means clustering* dengan atribut berupa nama barang dan jumlah persediaan berdasarkan data periode 2022 sampai 2025. Tahapan penelitian mengikuti proses KDD yang meliputi seleksi data, pembersihan data, transformasi data, serta proses pengelompokan. Penentuan jumlah *cluster* optimal dilakukan menggunakan metode *elbow*. Hasil penelitian menunjukkan terdapat tiga *cluster* terbaik, yaitu kategori prioritas rendah, sedang, dan tinggi. Dari total 220 jenis barang, diperoleh 139 barang prioritas rendah, 22 barang prioritas sedang, dan 59 barang prioritas tinggi.

Kata kunci: Data Mining; Persediaan Barang; K-Means Clustering

1. Pendahuluan

Persaingan dunia usaha yang semakin ketat mendorong pelaku bisnis untuk mengelola sumber daya secara optimal, salah satunya melalui pemanfaatan teknologi informasi. Pada era digital, data menjadi aset strategis yang dapat diolah menjadi informasi bernilai guna menunjang proses pengambilan keputusan. Pengelolaan data yang efektif, khususnya data persediaan barang, memiliki peran penting dalam menjaga keseimbangan antara ketersediaan persediaan dan kebutuhan pasar. Dengan demikian, penerapan teknologi serta metode analisis data menjadi langkah strategis bagi perusahaan untuk meningkatkan efisiensi operasional dan memperkuat daya saing [1].

Tani Mandiri merupakan sebuah usaha dagang yang bergerak dibidang pertanian dengan menjual berbagai macam pupuk dan alat pertanian. Permasalahan yang dihadapi oleh Tani Mandiri adalah mereka masih belum bisa mengelola persediaan barang mereka dan sering

mengalami permasalahan ketika tingkat persediaan barang tidak sesuai dengan tingkat penjualan produk. Persediaan barang sering kali tidak sesuai dengan kebutuhan, sehingga dapat menimbulkan dua permasalahan utama. Pertama, kekurangan persediaan yang menyebabkan terjadinya kehabisan stok. Kedua, kelebihan persediaan yang menyebabkan penumpukan barang digudang. Untuk mengatasi permasalahan tersebut, diperlukan pengelompokan data persediaan agar ketersediaan barang lebih terkontrol seperti yang juga dilakukan pada penelitian terdahulu [2], [3].

Salah satu solusi yang dapat diterapkan untuk mengatasi permasalahan tersebut adalah dengan memanfaatkan teknik *data mining*, khususnya dengan menggunakan metode *k-means clustering*. Metode ini mampu mengelompokkan data berdasarkan tingkat kemiripan sehingga memudahkan dalam mengidentifikasi pola persediaan barang agar lebih terkontrol. Beberapa penelitian terdahulu [2], dan penelitian [3], menunjukkan bahwa algoritma *k-means* efektif digunakan dalam pengelompokan data persediaan maupun penjualan untuk menentukan kategori barang laris, cukup laris, dan kurang laris. Dengan kemampuannya yang sederhana namun efektif, *k-means* dinilai tepat sebagai pendekatan dalam menentukan prioritas persediaan barang berdasarkan data persediaan yang tersedia.

Penelitian ini berfokus pada analisis data persediaan barang di Toko Tani Mandiri. Dengan memanfaatkan proses *clustering* akan dianalisis lebih lanjut untuk memberikan makna terhadap setiap *cluster* yang terbentuk. Berdasarkan hasil pengelompokan, diperoleh tiga *cluster* utama, yaitu *cluster 1* sebagai kelompok barang dengan prioritas rendah yang jarang dibeli, *cluster 2* sebagai kelompok barang dengan prioritas sedang yang memiliki tingkat pembelian stabil, dan *cluster 3* sebagai kelompok barang dengan prioritas tinggi yang memiliki frekuensi pembelian paling sering. Hasil ini memberikan wawasan bagi pihak Tani Mandiri dalam menentukan strategi pengelolaan persediaan barang. Barang dengan prioritas tinggi dapat dijadikan fokus utama untuk menjaga ketersediaan, sedangkan barang dengan prioritas rendah dapat dievaluasi kembali dalam proses pengadaan agar tidak terjadi penumpukan barang. Dengan demikian, hasil analisis ini dapat menjadi dasar dalam pengambilan keputusan yang lebih efisien dan tepat sasaran dalam pengelolaan persediaan barang di Tani Mandiri.

2. Tinjauan Pustaka

Hasil penelitian yang telah dilakukan [4], dapat disimpulkan bahwa proses pengelompokan data memberikan kontribusi yang penting bagi PT.BrothersIndo Saudara Sejati dalam pengelolaan persediaan barang. Informasi terkait produk yang memiliki tingkat penjualan tinggi maupun rendah membantu perusahaan dalam mencegah penumpukan persediaan di gudang serta menjaga kerapian dan keteraturan penyimpanan barang. Selain itu, pengolahan data ini berpotensi menjadi dasar dalam merumuskan strategi promosi yang tepat guna meningkatkan daya tarik produk yang kurang diminati. Hasil analisis menunjukkan terbentuknya dua kelompok utama, yaitu kelompok pertama yang terdiri dari 27 jenis produk dengan tingkat penjualan tinggi sehingga memerlukan penambahan stok, serta kelompok kedua yang mencakup 10 jenis produk dengan tingkat penjualan rendah yang perlu dikurangi jumlah persediaannya. Dengan demikian, analisis ini memberikan wawasan yang bernilai bagi perusahaan dalam mengoptimalkan manajemen inventaris guna meningkatkan efisiensi operasional dan keuntungan perusahaan.

Pada penelitian [5]. Dibahas penerapan *data mining* menggunakan metode *k-means clustering* untuk mengelompokkan penjualan sparepart motor di CV. Exa Jaya Motor. Masalah utama perusahaan adalah ketiadaan laporan penjualan yang terstruktur, sehingga menyulitkan dalam menentukan barang yang paling laku, cukup laku, maupun kurang laku. Dengan memanfaatkan data persediaan awal, jumlah barang terjual, dan persediaan akhir, algoritma *k-means* mengelompokkan data ke dalam tiga *cluster*. Hasil pengelompokan menunjukkan adanya 17 barang yang sangat laris, 21 barang yang laris, serta 22 barang yang kurang laku. Hasil tersebut membantu perusahaan dalam pengambilan keputusan, terutama terkait pengelolaan persediaan agar terhindar dari penumpukan atau kekurangan persediaan. Selain itu, implementasi sistem berbasis *web* turut memudahkan proses perhitungan, pelaporan, dan pemantauan penjualan.

Pada penelitian [6]. Disimpulkan bahwa pengelompokan data menggunakan *k-means clustering* dengan metode KDD pada toko sembako mampu memberikan pemahaman yang lebih mendalam terhadap pola penjualan berdasarkan data historis persediaan. Hasil pengelompokan

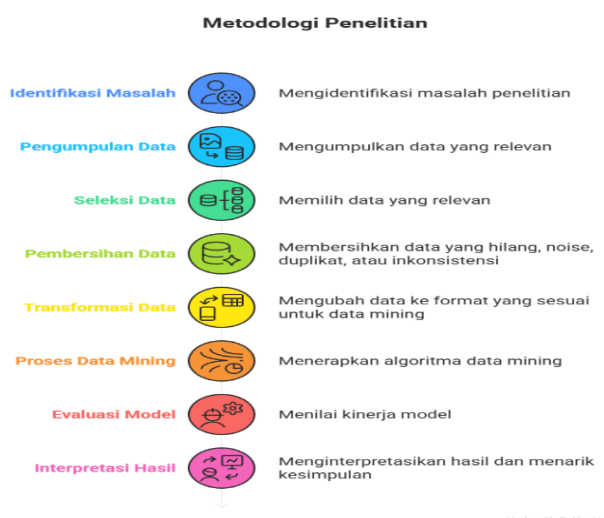
tersebut dapat dijadikan dasar dalam menyusun strategi manajemen persediaan yang lebih efisien dan efektif, sehingga ketersediaan barang di toko sembako dapat terjaga dengan baik.

Hasil penelitian [7]. Disimpulkan bahwa penerapan algoritma *k-means clustering* pada CV. Muria PS yang merupakan sebuah perusahaan dagang yang bergerak di bidang penjualan pakan unggas mampu menghasilkan kelompok persediaan barang yang tergolong laris, cukup laris, dan kurang laris. Hasil pengelompokan tersebut dapat dijadikan acuan bagi pihak manajemen untuk mengatur stok secara lebih efisien, sehingga kebutuhan pelanggan dapat terpenuhi dan mengurangi risiko kekecewaan saat membeli produk di CV. Muria PS.

Berdasarkan beberapa penelitian terdahulu, penelitian ini memperkuat hasil penelitian sebelumnya dengan menghadirkan konteks baru, yaitu pengelompokan persediaan barang berdasarkan data persediaan selama empat tahun terakhir yang mencakup nama barang dan jumlah pembelian setiap tahun, yang masih relatif jarang dikaji dalam literatur *data mining*. Temuan ini menunjukkan bahwa penerapan algoritma *k-means* tidak terbatas pada analisis data transaksi atau penjualan semata, tetapi juga dapat dimanfaatkan untuk mengidentifikasi tingkat pergerakan barang tinggi, sedang, maupun rendah melalui data persediaan atau pembelian barang. Dengan demikian, penelitian ini tidak hanya mendukung temuan sebelumnya, tetapi juga memperluas cakupan pemanfaatan algoritma *k-means* dalam analisis dan pengelolaan data persediaan barang.

3. Metodologi

Metodologi dalam penelitian ini dilakukan dari beberapa tahapan penelitian seperti pada Gambar 1. Tahapan - tahapan dalam metodologi penelitian ini adalah tahapan dalam penelitian secara umum, yaitu dengan mengidentifikasi masalah dan dilanjutkan dengan menggunakan konsep *Cross Industry Standard Process for Data Mining* (CRISP-DM) yang langkah-langkahnya selaras dengan *Knowledge Discovery in Database* (KDD) [8].



Gambar 1. Metodologi Penelitian

1) Identifikasi Masalah

Merupakan tahap awal untuk menemukan serta merumuskan permasalahan yang akan dikaji sehingga tujuan dan arah penelitian menjadi lebih terarah.

2) Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan untuk memperoleh informasi yang relevan mengenai kondisi persediaan barang yang berjalan pada Tani Mandiri. Metode yang digunakan adalah wawancara yang dilakukan dengan kepala toko dan karyawan guna mendapatkan gambaran yang jelas mengenai alur kerja, dan kendala yang sedang dihadapi. Melalui wawancara ini, diperoleh informasi mengenai proses persediaan barang pada Tani Mandiri dan *dataset* persediaan barang pada Tani Mandiri yang berjumlah 2.616 baris serta memiliki 10 kolom yang terlihat pada Gambar 2.

No	A	B	C	D	E	F	G	H	I	J
	No Nota	Nama Barang	Satuan	Banyaknya	Harga Satuan	jumlah	Tgl Pembelian	Nama Sales	Tempat Supplier	Alamat
1	2314	SEMANGKA INDEN KOTAK	Box	4	1,050,000.00	4,200,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
2	2314	SIDRA UP 1L BOX	Box	1	1,860,000.00	1,860,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
3	2314	HERBI SIDRA UP 5L	Ltr	4	460,000.00	1,840,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
4	2314	KRESNADAN 3R 1KG DUS	Box	5	300,000.00	1,500,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
5	2314	INSEKTISIDA MATARIN 80ML	Ltr	2	750,000.00	1,500,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
6	2314	HERBI KRESTARA 5L DUS	Box	2	740,000.00	1,480,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
7	2314	MATARIN 500ML DUS	Box	2	1,360,000.00	2,720,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
8	2314	MATARIN 250ML DUS	Box	4	700,000.00	2,800,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
9	2314	FUNG MATARIN 80ML DUS	Box	2	750,000.00	1,500,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH
10	2314	SAWI DORA KTK	Box	4	210,000.00	840,000.00	11/01/2022	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI, PRABULIH

Gambar 2. Dataset Persediaan Barang Tani Mandiri

3) Seleksi Data

Dalam proses seleksi data ini, *dataset* yang diperoleh dari Tani Mandiri, yaitu data persediaan barang dari 01 Januari 2022 sampai 31 Oktober 2025 sebanyak 2.616 *record* yang terdiri dari 10 *field* (fitur), yaitu nomor nota, nama barang, satuan, banyaknya, harga satuan, jumlah, tgl pembelian, nama sales, tempat supplier, dan alamat. Dari fitur tersebut, akan dipilih fitur yang sesuai untuk teknik *clustering data mining* persediaan barang pada Tani Mandiri, yaitu 2 fitur terdiri dari nama barang, dan banyaknya seperti yang terlihat pada Tabel 1.

Tabel 1. Hasil Seleksi Data Persediaan Barang

No	Nama Barang	Banyaknya
1	Semangka Inden Kotak	4
2	Sidra Up 1l Box	1
3	Herbi Sidra Up 5l	4
4	Kresnadan 3r 1kg Dus	5
5	Insektisida Matarin 80ml	2
6	Herbi Krestara 5l Dus	2
7	Matarin 500ml Dus	2
8	Matarin 250ml Dus	4
9	Fung Matarin 80ml Dus	2
10	Sawi Dora Ktk	4
...	...	...

4) Pembersihan Data

Di tahapan ini dilakukan proses pembersihan terhadap data yang duplikat, data kosong, ataupun inkonsisten data menggunakan *Google Colab*. Sebelumnya dilakukan perhitungan jumlah baris dan kolom pada *dataset* persediaan barang sebelum melanjutkan tahap penghapusan data yang duplikat. Didapatlah informasi jumlah baris sebanyak 2.616 dan jumlah kolom sebanyak 10. Kemudian dilakukan penghapusan data yang duplikat pada *dataset* persediaan barang dan didapatkan total jumlah baris dan kolom yaitu berjumlah 2.562 baris dan 10 kolom. Selanjutnya data dilakukan pengecekan apakah terdapat data yang kosong ataupun data inkonsisten seperti pada Gambar 3.

```
df=df.drop_duplicates(subset=['Nama Barang','Satuan','Banyaknya','Tgl Pembelian'], keep=False)
df
```

No	No Nota	Nama Barang	Satuan	Banyaknya	Harga Satuan	jumlah	Tgl Pembelian	Nama Sales	Tempat Supplier	Alamat
0	2314	SEMANGKA INDEN KOTAK	Box	4	1,050,000.00	4,200,000.00	2022-01-11	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI...
1	2314	SIDRA UP 1L BOX	Box	1	1,860,000.00	1,860,000.00	2022-01-11	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI...
2	2314	HERBI SIDRA UP 5L	Ltr	4	460,000.00	1,840,000.00	2022-01-11	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI...
3	2314	KRESNADAN 3R 1KG DUS	Box	5	300,000.00	1,500,000.00	2022-01-11	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI...
4	2314	INSEKTISIDA MATARIN 80ML	Ltr	2	750,000.00	1,500,000.00	2022-01-11	RENI	CV. PRABU ARGO SENTOSA	JLN. BUKIT TINGGI RT 001. RW 002 KEL. MAJASARI...

Gambar 3. Proses Penghapusan Data Yang Duplikat

5) Transformasi Data

Langkah selanjutnya setelah dilakukan penghapusan duplikasi data dan mengecek data yang hilang (*Missing Values*) yaitu transformasi data. beberapa atribut seperti nama barang diubah kedalam bentuk angka, untuk mempermudah proses pengelompokkan. Nama barang diperoleh sebanyak 220 jenis barang yang dijual pada Tani Mandiri, sehingga diberikan kode numerik dari 1 sampai 220, sedangkan untuk atribut banyaknya barang tidak diubah karena datanya sudah dalam bentuk angka sehingga didapatkan hasil transformasi data seperti yang terlihat pada Tabel 2.

Tabel 2. Hasil Transformasi Data

No	Nama Barang	Banyaknya Persediaan			
		2022	2023	2024	2025
1	1	40	60	30	30
2	2	20	20	20	40
3	3	0	0	20	0
4	4	40	40	100	120
5	5	10	2	2	14
6	6	8	0	0	0
7	7	20	50	66	80
8	8	260	100	10	160
9	9	90	130	220	250
10	10	50	165	115	120
...	...	...	...	...	...

6) Proses Data Mining

Pada tahap ini dilakukan penerapan algoritma *k-means clustering* untuk mengelompokkan barang berdasarkan atribut yang relevan. Proses ini menghasilkan beberapa *cluster*, yaitu *cluster 1* (kategori prioritas rendah), *cluster 2* (prioritas sedang), dan *cluster 3* (prioritas tinggi), yang membantu toko dalam memahami pola persediaan berdasarkan tingkat pembelian barang. Pemilihan jumlah *cluster* pada penelitian ini ditetapkan sebanyak tiga *cluster*, yaitu *cluster* prioritas rendah, prioritas sedang, dan prioritas tinggi. Penentuan tiga *cluster* dipilih karena pola penjualan dan pergerakan persediaan barang pada umumnya terbagi secara alami ke dalam tiga kategori tersebut. Jika hanya menggunakan dua *cluster*, pengelompokkan akan menjadi terlalu sederhana sehingga tidak mampu membedakan barang dengan tingkat pergerakan persediaan menengah, yang dalam praktiknya memiliki karakteristik berbeda dan memerlukan pengelolaan yang berbeda pula. Untuk menghitung jarak setiap data terhadap *centroid* menggunakan rumus seperti pada formula 1 *euclidean distance* berikut [9].

$$de = \sqrt{(xi - si)^2 + (yi - ti)^2} \quad (1)$$

Untuk menentukan serta menghitung *centroid* baru dengan mengambil nilai rata-rata dari seluruh anggota pada masing-masing *cluster* menggunakan rumus seperti pada formula 2 berikut [9].

$$c_k = \frac{1}{n_k} \sum d_i \quad (2)$$

7) Evaluasi Model

Pada tahap proses evaluasi model *clustering* dalam penelitian ini dilakukan dengan menggunakan dataset persediaan barang pada Tani Mandiri dengan menerapkan metode *elbow* (*Elbow Method*) yang menggunakan nilai *Sum of Squared Errors* (SSE) untuk menentukan jumlah *cluster* yang paling tepat atau optimal. Sehingga didapatkan hasil jumlah *cluster* yang paling optimal pada penelitian ini adalah tiga *cluster*, karena pada titik tersebut model mencapai keseimbangan terbaik antara tingkat kedetailan *cluster* dan efektivitas pengelompokan data [10].

8) Interpretasi Hasil

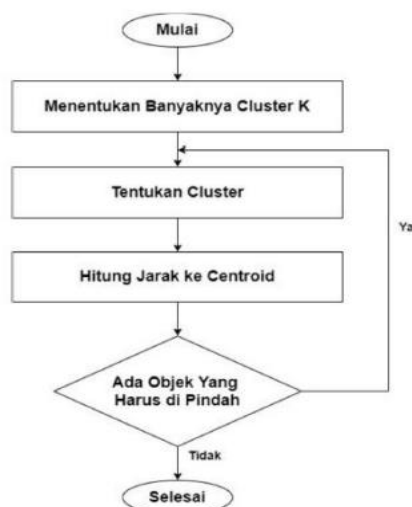
Pada tahap ini hasil dari proses *clustering* dianalisis lebih lanjut untuk memberikan makna terhadap setiap *cluster* yang terbentuk. Berdasarkan hasil pengelompokan, diperoleh tiga *cluster* utama, yaitu *cluster 1* sebagai kelompok barang dengan prioritas rendah yang jarang dibeli, *cluster 2* sebagai kelompok barang dengan prioritas sedang yang memiliki tingkat pembelian stabil, dan *cluster 3* sebagai kelompok barang dengan prioritas tinggi yang memiliki frekuensi pembelian paling sering. Ditahap ini juga dilakukan perancangan *dashboard* aplikasi *data mining clustering* untuk memproses hasil perhitungan *k-means* secara cepat dan efektif. Hasil interpretasi ini memberikan wawasan bagi pihak Tani Mandiri dalam menentukan strategi pengelolaan persediaan barang. Barang dengan prioritas tinggi dapat dijadikan fokus utama untuk menjaga ketersediaan, sedangkan barang dengan prioritas rendah dapat dievaluasi kembali dalam proses pengadaan agar tidak terjadi penumpukan barang [11].

4. Hasil dan Pembahasan

Proses *Knowledge Discovery in Database* (KDD) merupakan rangkaian langkah sistematis yang digunakan untuk mengekstraksi pengetahuan baru dari kumpulan data yang besar. KDD tidak hanya berfokus pada penerapan algoritma *data mining*, tetapi juga mencakup proses awal hingga akhir, mulai dari pengumpulan, pembersihan, pengolahan, hingga penafsiran hasil analisis [12]. Dalam penelitian ini, metode KDD diterapkan pada data persediaan barang Tani Mandiri dengan tujuan menemukan pola persediaan barang melalui penggunaan algoritma *k-means*. Dalam teknik *data mining* dilakukan proses-proses berikut.

4.1 Proses Perhitungan Manual Menggunakan K-Means

Pada tahap ini dilakukan contoh proses perhitungan *k-means* dimulai dari menentukan *cluster* sampai ke *centroid*. Langkah-langkah perhitungan tersebut dapat dilihat pada Gambar 4.



Gambar 4. Proses Perhitungan K-Means [13]

1. Menentukan jumlah *cluster* (*k*) yang diinginkan, yaitu $k = 3$ (yaitu baris 1, 4, dan 7 di Tabel 2).
2. Menentukan titik awal (*centroid*) secara acak sesuai dengan jumlah *cluster*, yaitu C1 (40,60,30,30), C2 (40,40,100,120), dan C3 (20,50,66,80) seperti pada Tabel 3.

Tabel 3. Titik Awal *Centroid*

Baris Centroid	Nama Barang	Banyaknya Persediaan			
		2022	2023	2024	2025
C1	1	40	60	30	30
C2	4	40	40	100	120
C3	7	20	50	66	80

3. Menghitung jarak setiap data terhadap *centroid* menggunakan rumus seperti pada formula 1 *euclidean distance* berikut [9].

$$de = \sqrt{(xi - si)^2 + (yi - ti)^2} \quad (1)$$

Data ke 1 (40,60,30,30)

$$C1 = \sqrt{(40 - 40)^2 + (60 - 60)^2 + (30 - 30)^2 + (30 - 30)^2} = 0$$

$$C2 = \sqrt{(40 - 40)^2 + (60 - 40)^2 + (30 - 100)^2 + (30 - 120)^2} = 115,76$$

$$C3 = \sqrt{(40 - 20)^2 + (60 - 50)^2 + (30 - 66)^2 + (30 - 80)^2} = 65,54$$

Data ke 2 (20,20,20,20)

$$C1 = \sqrt{(20 - 40)^2 + (20 - 60)^2 + (20 - 30)^2 + (20 - 30)^2} = 46,9$$

$$C2 = \sqrt{(20 - 40)^2 + (20 - 40)^2 + (20 - 100)^2 + (20 - 120)^2} = 131,15$$

$$C3 = \sqrt{(20 - 20)^2 + (20 - 50)^2 + (20 - 66)^2 + (20 - 80)^2} = 81,34$$

Data ke 3 (0,0,20,0)

$$C1 = \sqrt{(0 - 40)^2 + (0 - 60)^2 + (20 - 30)^2 + (0 - 30)^2} = 78,74$$

$$C2 = \sqrt{(0 - 40)^2 + (0 - 40)^2 + (20 - 100)^2 + (0 - 120)^2} = 154,92$$

$$C3 = \sqrt{(0 - 20)^2 + (0 - 50)^2 + (20 - 66)^2 + (0 - 80)^2} = 106,85$$

Data ke 4 (40,40,100,120)

$$C1 = \sqrt{(40 - 40)^2 + (40 - 60)^2 + (100 - 30)^2 + (120 - 30)^2} = 115,76$$

$$C2 = \sqrt{(40 - 40)^2 + (40 - 40)^2 + (100 - 100)^2 + (120 - 120)^2} = 0$$

$$C3 = \sqrt{(40 - 20)^2 + (40 - 50)^2 + (100 - 66)^2 + (120 - 80)^2} = 57,06$$

Data ke 5 (10,2,2,14)

$$C1 = \sqrt{(10 - 40)^2 + (2 - 60)^2 + (2 - 30)^2 + (14 - 30)^2} = 72,83$$

$$C2 = \sqrt{(10 - 40)^2 + (2 - 40)^2 + (2 - 100)^2 + (14 - 120)^2} = 152,26$$

$$C3 = \sqrt{(10 - 20)^2 + (2 - 50)^2 + (2 - 66)^2 + (14 - 80)^2} = 104,19$$

Dan diperhitungkan berikutnya menggunakan cara yang sama.

4. Mengelompokkan setiap objek ke dalam *cluster* yang memiliki jarak paling dekat dengan pusatnya. Hasil dari langkah 3 dan 4 dapat dilihat pada Tabel 4.

Tabel 4. Hasil Perhitungan *K-Means*

No	Nama Barang	2022	2023	2024	2025	C1	C2	C3	Jarak Terpendek	Cluster
1	1	40	60	30	30	0	115,76	65,54	0	1
2	2	20	20	20	20	46,9	131,15	81,34	46,9	1
3	3	0	0	20	0	78,74	154,92	106,85	78,74	1
4	4	40	40	100	120	115,76	0	57,06	0	2
5	5	10	2	2	14	72,83	152,26	104,19	72,83	1
6	6	8	0	0	0	80,15	164,39	115,76	80,15	1
7	7	20	50	66	80	65,54	57,06	0	0	3
8	8	260	100	10	160	259,42	248,39	263,89	248,39	2
9	9	90	130	220	250	303,15	204,69	252,82	204,69	2
10	10	50	165	115	120	162,63	126,29	134,63	126,29	2

5. Menentukan serta menghitung *centroid* baru dengan mengambil nilai rata-rata dari seluruh anggota pada masing-masing *cluster* menggunakan rumus seperti pada formula 2 berikut.

$$c_k = \frac{1}{n_k} \sum d_i \quad (2)$$

Ulangi kembali proses mulai dari langkah ke 3 jika nilai rata-rata setiap anggota belum stabil atau belum sama dengan hasil rata-rata pada iterasi sebelumnya [14].

Jika perhitungan dilakukan ke seluruh *dataset* menggunakan Google Colab, hasil perhitungan didapatkanlah pusat *cluster* yaitu.

1. Pada *cluster* kesatu (C1), terdapat seratus tiga puluh sembilan (139) data yang terdiri dari nomor transformasi nama barang 1, 2, 3, 5, 16, 17, 18, 20, 21, 22, 24, 25, 26, 27, 29, 31, 37, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 50, 51, 52, 53, 54, 55, 58, 60, 62, 63, 65, 66, 67, 68, 69, 70, 71, 72, 73, 74, 75, 76, 77, 78, 80, 81, 82, 83, 84, 85, 86, 87, 88, 89, 91, 92, 93, 94, 95, 98, 100, 102, 104, 105, 106, 107, 108, 109, 110, 111, 113, 115, 116, 117, 119, 120, 121, 122, 123, 127, 128, 130, 132, 135, 136, 137, 144, 146, 148, 153, 154, 155, 157, 158, 159, 160, 161, 162, 163, 167, 171, 172, 174, 176, 177, 182, 183, 184, 185, 186, 187, 188, 190, 191, 192, 193, 194, 195, 196, 197, 198, 199, 200, 201, 205, 206, 218, 219, 220 seperti yang terlihat pada Tabel 5.

Tabel 5. Hasil Data Yang Termasuk *Cluster 1*

No	Nama_Barang	Tahun_2022	Tahun_2023	Tahun_2024	Tahun_2025	Cluster
1	ALLY 250G	40	60	30	30	1
2	AMCOMIN 1L	20	20	20	40	1
3	ANDAIL 400ML	0	0	20	0	1
5	BABLAS 1L	10	2	2	14	1
16	BIBIT CABE GENIE KOTAK	0	3	6	0	1
...	...	...	...	...	...	...
205	TERONG REZA 10GR	0	1	1	0	1
206	TERONG RONGGO KOTAK	9	12	11	8	1
218	TYPHOSAT 485SL 20L	4	6	12	4	1
219	TYPHOSAT 485SL 5L	4	8	8	10	1
220	WULUHAN 10GR KTK	32	22	17	56	1

2. Pada *cluster* kedua (C2), terdapat dua puluh dua (22) data yang terdiri dari nomor transformasi nama barang 9, 13, 30, 33, 36, 64, 96, 99, 101, 134, 139, 140, 147, 152, 156, 173, 180, 202, 203, 207, 211, 215 seperti yang terlihat pada Tabel 6.

Tabel 6. Hasil Data Yang Termasuk *Cluster 2*

No	Nama_Barang	Tahun_2022	Tahun_2023	Tahun_2024	Tahun_2025	Cluster
9	BAYAM DOLLY SP	90	130	220	250	2
13	BERSILALANG 559 SL @20 LTR	110	430	100	130	2
30	CABE KERITING HELLBOY SP	150	190	270	320	2
33	CABE RAWIT JAWARA 99 SP	105	130	300	300	2
36	CAISIM GLORY SP	105	120	185	220	2
...	...	...	...	...	...	...
202	TERONG BRUNO SP	305	380	415	455	2

No	Nama_Barang	Tahun_2022	Tahun_2023	Tahun_2024	Tahun_2025	Cluster
203	TERONG JEMY SP	150	170	230	280	2
207	TIMUN MAESTRO SP	130	120	210	220	2
211	TOMAT LIPPO SP	130	170	230	290	2
215	TURMADAN 328 SL @1 L	100	130	140	230	2

3. Pada *cluster* ketiga (C3), terdapat lima puluh sembilan (59) data yang terdiri dari nomor transformasi nama barang 4, 7, 10, 11, 12, 14, 15, 19, 23, 28, 32, 34, 35, 38, 56, 57, 59, 61, 90, 97, 103, 112, 114, 118, 124, 125, 126, 129, 131, 133, 138, 141, 142, 143, 145, 149, 150, 151, 164, 165, 166, 168, 168, 171, 175, 178, 179, 181, 189, 204, 204, 208, 209, 210, 212, 213, 214, 216, 217 seperti yang terlihat pada Tabel 7.

Tabel 7. Hasil Data Yang Termasuk *Cluster* 3

No	Nama_Barang	Tahun_2022	Tahun_2023	Tahun_2024	Tahun_2025	Cluster
4	ASXONE 138 SL @20LTR	40	40	100	120	3
7	BAROZEB BIRU 1KG	20	50	66	80	3
10	BERSIH ALANG 488 SL @20L	50	165	115	120	3
11	BERSILALANG 328 SL @1 LTR	0	200	0	0	3
12	BERSILALANG 559 SL @1 LTR	50	165	115	120	3
...	...	...	...	...	...	...
212	TOMAT TYFANI SP	30	50	90	130	3
213	TOPSAT 1LT	30	40	40	70	3
214	TRISULA 500	50	30	20	90	3
216	TURMADAN 328 SL @5 LTR	80	120	40	150	3
217	TURMADAN 328SL @20 LTR	40	280	0	0	3

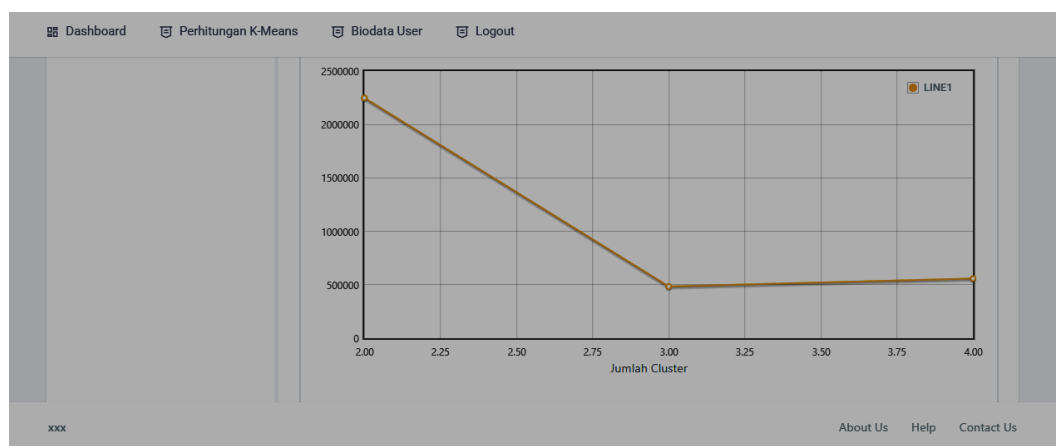
4.2 Hasil Evaluasi Model

Proses evaluasi model *clustering* dalam penelitian ini dilakukan dengan menggunakan dataset persediaan barang pada Tani Mandiri dengan menerapkan metode *elbow (Elbow Method)* yang menggunakan nilai *Sum of Squared Errors (SSE)* untuk menentukan jumlah *cluster* yang paling tepat atau optimal. SSE berfungsi sebagai indikator yang menunjukkan seberapa erat data dalam suatu *cluster* berkumpul di sekitar *centroid*, semakin rendah nilai SSE, semakin baik kualitas pembentukan *cluster* tersebut seperti pada Gambar 5.

Cluster	Tahun_2022	Tahun_2023	Tahun_2024	Tahun_2025	Hasil Penjumlahan
Cluster - 2	(40 - 19.380710659898) ²	(60 - 28.883248730964) ²	(30 - 25.659898477157) ²	(30 - 33.761421319797) ²	1426.392074
Cluster - 2	(20 - 19.380710659898) ²	(20 - 28.883248730964) ²	(20 - 25.659898477157) ²	(40 - 33.761421319797) ²	150.2499420
Cluster - 2	(0 - 19.380710659898) ²	(0 - 28.883248730964) ²	(20 - 25.659898477157) ²	(0 - 33.761421319797) ²	2381.722023

Gambar 5. Proses Elbow Pada *Dataset* Persediaan Barang

Berdasarkan grafik hasil perhitungan SSE yang ditunjukkan pada Gambar 6. Untuk jumlah *cluster* 2, 3, dan 4, terlihat bahwa nilai SSE pada $k = 2$ masih sangat besar, sehingga pengelompokan data belum efektif. Ketika jumlah *cluster* dinaikkan menjadi $k = 3$, terjadi penurunan SSE yang cukup drastis, menandakan bahwa pemisahan data menjadi tiga *cluster* memberikan peningkatan kualitas *clustering* yang signifikan. Sementara itu, pada $k = 4$, penurunan SSE hanya sedikit dan grafik mulai menunjukkan pola mendatar, yang mengindikasikan bahwa penambahan *cluster* tidak lagi memberikan peningkatan berarti. Dari pola tersebut, titik “*elbow*” atau titik belokan grafik berada pada $k = 3$. Dengan melihat pola tersebut, titik “*elbow*” atau titik bahu grafik berada pada $k = 3$. Oleh karena itu, jumlah *cluster* yang paling optimal pada penelitian ini adalah tiga *cluster*, karena pada titik tersebut model mencapai keseimbangan terbaik antara tingkat kedetailan *cluster* dan efektivitas pengelompokan data.



Gambar 6. Hasil Grafik *Elbow*

4.3 Interpretasi Hasil

Menafsirkan hasil penentuan pada *cluster* merupakan langkah penting untuk memperoleh pemahaman yang lebih jelas mengenai makna dari setiap *cluster* yang dihasilkan. Berikut ini adalah analisis yang dapat dilakukan berdasarkan hasil penentuan *cluster* tersebut.

1. *Cluster* C1 (Pembelian Rendah)

a) Ciri – Ciri Anggota *Cluster* C1

Pada *cluster* kesatu (C1), terdapat seratus tiga puluh sembilan (139) data yang terdiri dari nomor transformasi nama barang 1, 2, 3, 5, 16, 17, 18, 20, 21, 22, 24, 25, 26, 27, 29, 31, 37, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 50, 51, 52, 53, 54, 55, 58, 60, 62, 63, 65, 66, 67, 68, 69, 70, 71, 72, 73, 74, 75, 76, 77, 78, 80, 81, 82, 83, 84, 85, 86, 87, 88, 89, 91, 92, 93, 94, 95, 98, 100, 102, 104, 105, 106, 107, 108, 109, 110, 111, 113, 115, 116,

117, 119, 120, 121, 122, 123, 127, 128, 130, 132, 135, 136, 137, 144, 146, 148, 153, 154, 155, 157, 158, 159, 160, 161, 162, 163, 167, 171, 172, 174, 176, 177, 182, 183, 184, 185, 186, 187, 188, 190, 191, 192, 193, 194, 195, 196, 197, 198, 199, 200, 201, 205, 206, 218, 219, 220.

b) Dasar Penentuan *Cluster C1* (Pembelian Rendah)

Barang tersebut kurang diminati menunjukkan bahwa barang tersebut kemungkinan memiliki daya tarik rendah bagi pelanggan atau permintaannya tidak terlalu tinggi. Oleh karena itu, diperlukan peninjauan kembali terhadap *item* dalam *cluster* ini, misalnya dengan mengurangi jumlah persediaan atau mengalihkan sumber daya ke barang yang memiliki permintaan lebih besar.

2. *Cluster C2* (Pembelian Sedang/Stabil)

a) Anggota *Cluster C2*

Pada *cluster* kedua (*C2*), terdapat dua puluh dua (22) data yang terdiri dari nomor transformasi nama barang 9, 13, 30, 33, 36, 64, 96, 99, 101, 134, 139, 140, 147, 152, 156, 173, 180, 202, 203, 207, 211, 215.

b) Dasar Penentuan *Cluster C2* (Pembelian Sedang/Stabil)

Barang-barang ini memiliki tingkat permintaannya lebih rendah dibandingkan *cluster* tiga. Istilah "Pembelian Stabil" menggambarkan bahwa barang ini tidak sepenuhnya laku dan tidak laku, namun permintaan terhadap barang tersebut relatif stabil dibandingkan dengan barang-barang pada kelompok lainnya.

3. *Cluster C3* (Pembelian Tinggi)

a) Anggota *Cluster C3*

Pada *cluster* ketiga (*C3*), terdapat lima puluh sembilan (59) data yang terdiri dari nomor transformasi nama barang 4, 7, 10, 11, 12, 14, 15, 19, 23, 28, 32, 34, 35, 38, 56, 57, 59, 61, 90, 97, 103, 112, 114, 118, 124, 125, 126, 129, 131, 133, 138, 141, 142, 143, 145, 149, 150, 151, 164, 165, 166, 168, 168, 171, 175, 178, 179, 181, 189, 204, 204, 208, 209, 210, 212, 213, 214, 216, 217.

b) Dasar Penentuan *Cluster C3* (Pembelian Tinggi)

Barang - barang ini merupakan pilihan utama pelanggan karena memiliki tingkat pembelian atau persediaan barang yang jauh lebih tinggi dibandingkan *cluster* lainnya, berdasarkan data persediaan barang yang tersedia.

4.4 Tampilan Antarmuka (*Interface*) *Clustering Data Barang*

Tampilan antarmuka (*Interface*) merupakan elemen penting dalam sebuah sistem informasi karena berfungsi sebagai media interaksi antara pengguna dan aplikasi. Antarmuka pengguna (*User Interface*) merupakan segala bentuk yang memfasilitasi interaksi antara manusia sebagai pengguna dengan sistem komputer, perangkat lunak, atau perangkat lainnya. UI berfungsi sebagai penghubung yang memungkinkan pengguna berkomunikasi dengan teknologi yang digunakan. Tujuan utama UI adalah meningkatkan kemudahan penggunaan serta mengoptimalkan komunikasi antara pengguna dan sistem [15]. Dalam sistem *k-means clustering* berbasis *web PHP* yang dikembangkan untuk mengelompokkan prioritas persediaan barang pada Tani Mandiri, antarmuka dirancang secara sederhana, responsif, dan mudah digunakan oleh pengguna.

Sistem ini menyediakan empat halaman utama yang mendukung proses pengelompokan data persediaan barang, yaitu:

1. Halaman Perhitungan *K-Means*

- a) Tentukan *Cluster* = Pada *menu* ini ditentukan *centroid* dan jumlah *cluster* yang diinginkan. Pada *dataset* persediaan barang ini dipilih *random centroid* dan memasukkan jumlah *cluster* sebanyak tiga, kemudian klik simpan untuk lanjut tahapan selanjutnya seperti pada Gambar 7.

Gambar 7. Halaman Tentukan Cluster

- b) Proses *K-Means* = Ini merupakan hasil pengelompokan data persediaan barang pada Toko Tani Mandiri, setelah beberapa iterasi hingga terhenti di iterasi ke 6 dan pusat *centroid* telah dinyatakan stabil, hasil akan terlihat barang mana dengan tingkat pembelian tinggi, sedang, dan rendah seperti pada Gambar 8.

Perulangan Ke - 6				
Perulangan 6 - Penentuan Centroid				
Centroid 1	9.726618705036	13.115107913669	11.330935251799	13.122302158273
Centroid 2	121.59090909091	192.5	236.81818181818	302.95454545455
Centroid 3	43.491525423729	68.135593220339	60.525423728814	84.525423728814
Perulangan 6 - Hitung Euclidean Distance				
ALLY 250G	61.221412789859	376.15325422114	63.112613267568	

Gambar 8. Centroid Baru Untuk Iterasi terakhir ke-6

- c) *Clustering* = Pada *menu* ini akan ditampilkan jumlah dari proses *cluster*. Didapatlah hasil jumlah *cluster* pada *dataset* persediaan barang di Tani Mandiri dengan total barang yang dijual berjumlah 220 jenis barang yaitu *cluster* 1 (pembelian rendah) sebanyak 139 jenis barang, *cluster* 2 (pembelian sedang) sebanyak 22 jenis barang, dan *cluster* 3 (pembelian tinggi) berjumlah 59 jenis barang, kemudian pengguna bisa melakukan *export* data hasil *cluster* tadi dengan mengklik tombol *export* seperti pada Gambar 9.

Dashboard					
KANGKUNG WALLET SP	40	75	115	100	3
RIDATOP 5L	20	65	75	60	3
Jumlah Cluster					
Cluster	Jumlah				
1	139				
3	59				
2	22				
Export					

Gambar 9. Halaman Clustering

4.5 Pembahasan

Berdasarkan hasil perhitungan dengan menggunakan algoritma *k-means clustering* didapatlah tiga *cluster* untuk menentukan prioritas persediaan, yaitu *cluster* 1 (prioritas pembelian rendah), *cluster* 2 (prioritas pembelian sedang), dan *cluster* 3 (prioritas pembelian tinggi). Setelah proses perhitungan melalui aplikasi, diperoleh hasil pengelompokan untuk 220 jenis barang yang dijual pada Tani Mandiri, yaitu 139 jenis barang termasuk dalam *cluster* 1 (pembelian rendah), 22 jenis barang dalam *cluster* 2 (pembelian sedang), dan 59 jenis barang masuk ke *cluster* 3 (pembelian tinggi). Dari hasil pengelompokan ini diharapkan dapat memberikan informasi penting bagi pihak Tani Mandiri dalam merumuskan strategi pengelolaan persediaan barang.

Barang dengan prioritas tinggi dapat dijadikan fokus utama dalam menjaga persediaan barangnya agar selalu ada, sementara barang dengan prioritas rendah dapat ditinjau kembali pada proses persediaannya untuk menghindari terjadinya kelebihan persediaan barang yang dapat menyebabkan penumpukan barang digudang. Dengan demikian, hasil penelitian ini dapat dijadikan sebagai dasar pengambilan keputusan yang lebih efektif dan efisien dalam pengelolaan persediaan barang di Toko Tani Mandiri.

Temuan tersebut mendukung hasil penelitian sebelumnya yang dilakukan oleh [5], yang juga menggunakan algoritma *k-means clustering* dalam mengelompokkan persediaan barang berdasarkan data persediaan awal, jumlah barang terjual, dan persediaan akhir algoritma *k-means* mengelompokkan data ke dalam tiga *cluster* yang menunjukkan hasil pengelompokan adanya 17 barang yang sangat laris, 21 barang yang laris, serta 22 barang yang kurang laku. Penelitian lain oleh [6], disimpulkan bahwa pengelompokan data menggunakan *k-means clustering* dengan metode kdd pada toko sembako mampu memberikan pemahaman yang lebih mendalam terhadap pola penjualan berdasarkan data historis persediaan. Hasil pengelompokan tersebut dapat dijadikan dasar dalam menyusun strategi manajemen persediaan yang lebih efisien dan efektif, sehingga ketersediaan barang di toko sembako dapat terjaga dengan baik.

Dalam hal kontribusi, penelitian ini memperkuat penelitian sebelumnya dengan menawarkan konteks baru, yaitu pengelompokkan persediaan barang dengan menggunakan data persediaan barang selama 4 tahun terakhir berdasarkan nama barang dan banyaknya barang yang di beli setiap tahun, yang saat ini jarang diteliti dalam literatur *data mining*. Penemuan ini memperluas penggunaan *k-means* dari sektor penjualan bahwa untuk melihat barang – barang yang memiliki tingkat penjualan tinggi, sedang, maupun rendah tidak hanya menggunakan data transaksi atau penjualan saja tetapi bisa juga menggunakan data persediaan atau pembelian barang. Oleh karena itu, penelitian ini tidak hanya menguatkan temuan – temuan sebelumnya tetapi juga memperluas jangkauan pengetahuan pada penggunaan algoritma *k-means* dalam menganalisis data persediaan barang.

5. Simpulan

Penelitian ini memanfaatkan data persediaan barang dari Toko Tani Mandiri dengan total 2.562 baris data dan dua atribut, yaitu nama barang dan jumlahnya. Metode yang digunakan adalah *k-means clustering*. Berdasarkan perhitungan yang dilakukan, direkomendasikan tiga *cluster* untuk menentukan prioritas persediaan, yaitu *cluster* 1 (prioritas pembelian rendah), *cluster* 2 (prioritas pembelian sedang), dan *cluster* 3 (prioritas pembelian tinggi). Perbedaan hasil antara perhitungan manual dan perhitungan menggunakan aplikasi disebabkan oleh penentuan pusat *cluster* yang bersifat acak. Setelah proses perhitungan melalui aplikasi, diperoleh hasil pengelompokan untuk 220 jenis barang yang dijual di Tani Mandiri, yaitu 139 jenis barang termasuk dalam *cluster* 1 (pembelian rendah), 22 jenis barang dalam *cluster* 2 (pembelian sedang), dan 59 jenis barang masuk ke *cluster* 3 (pembelian tinggi).

Penelitian ini diharapkan dapat memberikan pemahaman yang lebih baik bagi Tani Mandiri dalam menentukan strategi pengelolaan persediaan. Barang dengan prioritas pembelian tinggi dan sedang dapat dijadikan perhatian utama untuk memastikan stok tetap tersedia, sementara barang dengan prioritas rendah dapat ditinjau kembali dalam proses pengadaannya guna mencegah penumpukan persediaan barang di gudang.

Daftar Referensi

- [1] M. G. Saparudin dan S. Sholihin, "Penggunaan Data Mining untuk Analisis Pola Pembelian Pelanggan Menggunakan Metode Association Rule Algoritma Apriori (Studi Kasus di Toko Waspada)," *J. Teknol. Sist. Inf. dan Apl.*, vol. 6, no. 1, pp. 27–33, 2023, doi: 10.32493/jtsi.v6i1.26927.
- [2] A. Azis dan Sutisna, "Penerapan Data Mining untuk Menentukan Ketersediaan Stok Barang Berdasarkan Permintaan Konsumen di PT Indonesia Thai Summit Plastech Menggunakan K-Means Clustering," vol. 5, no. 3, pp. 3099–3106, 2024.
- [3] H. Prastiwi, Jeny Pricilia, dan Errissya Rasywir, "Implementasi Data Mining Untuk Menentukan Persediaan Stok Barang Di Mini Market Menggunakan Metode K-Means Clustering," *J. Inform. Dan Rekayasa Komputer(JAKAKOM)*, vol. 2, no. 1, pp. 141–148, 2022.
- [4] R. K. Serli, B. Wijonarko, dan M. S. Nurhayati, "Analisis Persediaan Barang Menggunakan Clustering K-Means Pada PT. Brothersindo Saudara Sejati," *J. Tek.*, vol. 17, no. 2, pp. 569–

- 580, 2023.
- [5] J. M. Sitinjak, K. Sari, dan M. Yetri, "Penerapan Data Mining Dalam Penjualan Sparepart Motor Dengan Menggunakan Metode K-Means Clustering," vol. 3, no. November, pp. 862–871, 2024.
 - [6] F. Zafira, B. Irawan, dan A. Bahtiar, "Penerapan Data Mining Untuk Estimasi Stok Barang Dengan Metode K-Means Clustering," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 156–161, 2024, doi: 10.36040/jati.v8i1.8319.
 - [7] F. Nurdiyansyah dan I. Akbar, "Implementasi Algoritma K-Means untuk Menentukan Persediaan Barang pada Poultry Shop," *J. Teknol. dan Manaj. Inform.*, vol. 7, no. 2, pp. 86–94, 2021, doi: 10.26905/jtmi.v7i2.6377.
 - [8] N. Hidayati, H. Widi Nugroho, dan Nurjoko, "Penerapan Data Mining Untuk Menghasilkan Pola Pembelian Roti Menggunakan Algoritma Apriori," *Semin. Nas. Has. Penelit. dan Pengabd. Masy.*, pp. 246–254, 2021.
 - [9] D. Ramdhan, G. Dwilestari, R. D. Dana, A. Ajiz, dan K. Kaslani, "Clustering Data Persediaan Barang Dengan Menggunakan Metode K-Means," *MEANS (Media Inf. Anal. dan Sist.*, vol. 7, no. 1, pp. 1–9, 2022, doi: 10.54367/means.v7i1.1826.
 - [10] Muhammad Raqib Syahkur, D. Hartama, dan S. Solikhun, "Evaluasi Jumlah Cluster pada Algoritma K-Means++ Menggunakan Silhouette dan Elbow dengan Validasi Nilai DBI dalam Mengelompokkan Gizi Balita," *JST (Jurnal Sains dan Teknol.*, vol. 13, no. 3, pp. 487–496, 2024, doi: 10.23887/jstundiksha.v13i3.86419.
 - [11] F.R. Ani, R. Kurniawan, & T. Suprpti, "Penerapan Algoritma K-Means Clustering Untuk Perencanaan Persediaan Stok Sepatu Di Toko Diomclothing. JATI (Jurnal Mahasiswa Teknik Informatika), vol. 7, no. 6, pp. 3699-3707, 2023.
 - [12] A. A. Ningtiyas, N. B. Nugroho, dan M. Syaifuddin, "Penerapan Data Mining Untuk Menentukan Persediaan Barang Pada PT. Deli Food Menggunakan Metode K-Means," vol. 4, no. 7, pp. 1-12, 2021, [Daring]. Tersedia pada: <https://ojs.trigunadharma.ac.id/>
 - [13] D. A. Ramadhanty, R. Syafitri, E. Raswir, dan D. Meisak, "Implementasi Data Mining Untuk Menentukan Persediaan Stok Obat Di Apotek K-24 Menggunakan Metode K-Means Clustering," *J. Inform. dan Rekayasa Komput.*, vol. 1, no. 2, pp. 155–160, 2022.
 - [14] T. B. Pamungkas, S. Maesaroh, dan P. Ardiansyah, "Implementasi Data Mining Pada Stok Penggunaan Barang Di Gmf Aeroasia Menggunakan Algoritma K-Means Clustering," *J. Ilm. Sains dan Teknol.*, vol. 7, no. 2, pp. 112–123, 2023, doi: 10.47080/saintek.v7i2.2697.
 - [15] Y. Prihastomo dan W. Winanti, "Tren Perkembangan Graphical User Interface Melalui Permohonan Desain Industri Di Indonesia," *J. Commun. Educ.*, vol. 18, no. 1, pp. 93–100, 2024, doi: 10.58217/joce-ip.v18i1.390.