

PENERAPAN DATA MINING UNTUK ESTIMASI STOK BARANG DENGAN METODE K-MEANS CLUSTERING

Fara Zafira, Bambang Irawan, Agus Bahtiar

Program Studi Teknik Informatika, Sekolah Tinggi Manajemen Informatika Dan Komputer IKMI Cirebon
Jl. Perjuangan No.10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Telp. 0231-490480-490481
farazafira77@gmail.com

ABSTRAK

Toko sembako merupakan suatu usaha jual beli barang yang banyak terjadi transaksi baik itu makanan, minuman maupun barang kebutuhan masyarakat dan lainnya. Toko sembako dapat menggunakan metode *k-means clustering* untuk mengetahui stok barang yang masih ada dalam stok dan yang sudah terjual atau tidak ada. Permasalahan yang dihadapi oleh toko sembako adalah mereka masih belum bisa mengelola persediaan barang mereka. Toko sembako ini sering mengalami permasalahan ketika tingkat persediaan tidak sesuai dengan tingkat penjualan produk. Persediaan barang seringkali kurang dari yang diperlukan, sehingga menyebabkan situasi kehabisan stok. Sebaliknya, situasi serupa terjadi, ketika persediaan terlalu besar dibandingkan dengan penjualan, sehingga persediaan menumpuk. Dengan melakukan melakukan pengelompokan, permasalahan seperti kekurangan atau kelebihan barang yang dijual dapat dihindari. Penelitian ini bertujuan untuk menggunakan metode *k-means clustering* untuk mengelompokkan data stok barang dan juga akan dibahas bagaimana teknik ini diterapkan pada *RapidMiner*. Metode *k-means clustering* bekerja dengan memasukkan data persediaan stok barang yang akan diubah ke dalam dataset. Dengan metode tahapan *Knowledge Discovery in Database (KDD)*, menggunakan *RapidMiner* dapat diketahui stok barang yang masih ada dalam stok sebanyak 466 item dan yang sudah terjual atau tidak ada stok sebanyak 37 items, Berdasarkan hasil analisis yang telah dilakukan maka dapat disimpulkan bahwa *Presentase* setiap *cluster* yaitu *cluster_0* sebanyak 466 *items* atau sebesar 92,6%, *cluster_1* sebanyak 37 *items* atau sebesar 7,4%. Dan $k=2$ dengan nilai *davies boudin* (Dbi) = -0.330.

Kata kunci : Data mining, Persediaan stok, toko sembako, clustering, persediaan stok, k-means

1. PENDAHULUAN

Data mining adalah metode untuk menemukan pola tertentu dari kumpulan data yang berjumlah besar. Meskipun banyak dipelajari pada bidang ilmu komputer dan statistika, data mining adalah metode yang bisa diterapkan dan mempermudah pekerjaan di bidang lainnya juga. Metode *K-Means* sudah menjadi salah satu teknik data mining yang digunakan untuk merancang strategi persediaan barang yang efektif.[1]

Toko sembako merupakan suatu usaha jual beli barang yang banyak terjadi transaksi baik itu makanan, minuman maupun barang kebutuhan masyarakat dan pelanggan lainnya. Dengan banyaknya pelanggan yang keluar masuk setiap harinya, kemungkinan penjualan yang dilakukan di toko sembako sangatlah signifikan. Tentu saja, penjualan yang sangat besar menghasilkan data dalam jumlah yang sangat besar. Jika data yang sangat besar tidak diolah tetapi hanya disimpan maka akan mengakibatkan penumpukan data. Dengan menggunakan data mining untuk mengolah data yang ada diharapkan data penjualan dalam jumlah besar akan memberikan manfaat yang dapat memudahkan pekerjaan manusia.[2] Salah satu manfaat yang akan terlihat jika data penjualan toko sembako diolah dengan menggunakan data mining adalah dapat mengelompokkan produk barang yang sudah terjual dan masih ada dalam stok persediaan.[3]

Toko sembako ini sering mengalami permasalahan ketika tingkat persediaan tidak sesuai dengan tingkat penjualan produk. Persediaan barang seringkali kurang dari yang diperlukan, sehingga

menyebabkan situasi kehabisan stok. Sebaliknya, situasi serupa terjadi, ketika persediaan terlalu besar dibandingkan dengan penjualan, sehingga persediaan menumpuk. Dengan melakukan melakukan pengelompokan, permasalahan seperti kekurangan atau kelebihan barang yang dijual dapat dihindari. Pengolahan persediaan barang penjualan merupakan pengetahuan dasar bisnis yang biasa dikenal dengan stok manajemen.

Tanpa penanganan persediaan yang baik, penjual tidak akan dapat memenuhi kebutuhan pelanggan karena kurangnya barang yang dibutuhkan. Jika persediaan terlalu banyak, hal ini dapat mengakibatkan kualitas barang menurun dalam jangka waktu yang lama karena kurang terpakai, dan persediaan barang seperti makanan dan minuman dapat kadaluarsa.

Untuk mengatasi permasalahan tersebut diterapkan suatu metode data mining dengan cara menganalisa data penjualan untuk menghasilkan informasi yang dapat dijadikan sebagai perencanaan dan pengendalian persediaan barang. Adapun pengolahan datanya dapat dilakukan melalui proses *clustering* data dengan menerapkan metode *k-means clustering*. Metode *k-means clustering* bertujuan mengelompokkan data yang mempunyai karakteristik yang sama ke dalam satu klaster yang sama dan data yang mempunyai karakteristik berbeda dikelompokkan ke dalam klaster lain.

Pada penelitian ini akan dilakukan analisis data terhadap data penjualan produk pada toko sembako dengan mengelompokkan produk berdasarkan 2

kelompok yaitu produk barang yang sudah terjual dan masih ada dalam stok. Untuk membantu memberikan informasi persediaan barang dagangan di toko sembako.

Berdasarkan permasalahan diatas peneliti melakukan penelitian dengan judul “Analisis data untuk estimasi stok penjualan menggunakan metode *K-Means Clustering*”. Analisis metode *clustering* tersebut lebih efektif dan efisien sehingga dapat mengetahui tersedianya barang dengan jenis dan jumlah yang sesuai kebutuhan konsumen.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data mining adalah proses pengumpulan data historis untuk menemukan keteraturan dan hubungan dalam kumpulan data yang besar. Tujuan dari data mining adalah untuk menciptakan hubungan yang jelas dan sebelumnya tidak diketahui antara data dan sumber lain sehingga dapat digunakan dan dipahami.[1]

Data mining adalah proses menggunakan statistik, matematika, kecerdasan buatan dan teknik pembelajaran mesin untuk mengekstrak dan mengidentifikasi informasi yang berguna dan pengetahuan terkait dari database besar. Data mining juga merupakan metode yang digunakan dalam pengolahan data skala besar[4]

2.2. Persediaan stok barang

Persediaan barang dapat juga diartikan sebagai simpanan material, seperti bahan mentah, barang dalam proses, dan barang jadi, dan pengendalian persediaan adalah upaya untuk menjaga jumlah persediaan tetap pada tingkat yang diinginkan.[5] Untuk barang jasa, pengendalian persediaan difokuskan pada material, sedangkan untuk barang pasokan, pengendalian persediaan difokuskan pada jasa pasokan, Persediaan membuat proses bisnis lebih mudah dan memberikan solusi kepada pengusaha untuk menentukan jumlah barang yang ideal dan ekonomis untuk dipesan, dapat memberikan solusi kepada pengusaha untuk menentukan kapan harus memesan barang,[6]

2.3. Algoritma K-means Clustering

Algoritma *K-Means* merupakan salah satu metode pengelompokan data non-hierarki yang bertujuan untuk membagi data yang ada ke dalam satu atau lebih *cluster* atau kelompok sehingga data yang memiliki atribut serupa berada dalam satu *cluster* yang sama dan data yang memiliki atribut serupa berada dalam satu *cluster* yang sama. terletak pada kelompok yang berbeda, Algoritma *K-Means* adalah metode yang relatif sederhana untuk mengorganisasikan sejumlah besar objek dengan properti tertentu ke dalam kelompok (*cluster*) yang disebut “K” dan jumlah k cluster.

[1]

2.4. Clustering

Clustering adalah proses pengelompokan data yang dianalisis berdasarkan kasus ke dalam kelas-kelas yang dianggap serupa, Setiap *cluster* mempunyai data dengan karakteristik yang sama, namun karakteristiknya berbeda dengan data pada *cluster* lainnya.[7] Ada perbedaan antara pengelompokan dan klasifikasi. Dengan kata lain, Jika *clustering* tidak memiliki tujuan, kelas, atau label, maka termasuk dalam kategori *unsupervised learning*. Proses penambahan data sering menggunakan *clustering*.[8]

Clustering adalah teknik penambahan data yang menggabungkan data atau objek ke dalam kelas atau *cluster* berdasarkan karakteristik yang sama. Objektif yang kurang mirip dengan objek-objek dalam kelompok lain diciptakan melalui pengelompokan yang tepat, namun lebih mirip dengan objek-objek dalam kelompok atau *cluster* yang sama. [1]

Untuk mengelompokkan sekumpulan objek ke dalam sebuah cluster. Cluster yang baik adalah cluster yang mempunyai tingkat kemiripan yang tinggi, Dismilaritas yang tinggi antar objek dalam suatu cluster dan ketidakmiripan yang tinggi dengan objek cluster lainnya.[9]

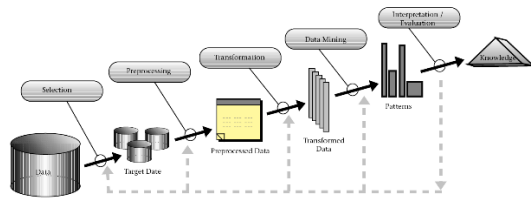
2.5. RapidMiner

RapidMiner adalah aplikasi data mining berbasis *open-source* yang terkemuka dan ternama. Didalamnya terdapat aplikasi yang berdiri sendiri untuk analisis data dan sebagai mesin data mining seperti untuk loading data, transformasi data, pemodelan data, dan metode visualisasi data. *RapidMiner* pertama kali dinamai *Yet Another Learning Environment* atau singkat *YALE*. Pada tahun 2007 akhirnya diganti namanya menjadi *RapidMiner*.[10]

Dengan antarmuka pengguna *grafis (GUI)* *Rapidminer*, pengguna dapat membuat desain pipa analisis. *GUI* membuat *file XML (Extensible Markup Language)* yang berisi prosedur analisis yang diinginkan pengguna, dan *Rapidminer* secara otomatis melakukan analisis pada file ini.[1]

3. METODE PENELITIAN

Dalam penelitian ini akan menggunakan *Knowledge Discovery in Database (KDD)* sebagai bagian dari proses analisis data. proses KDD adalah proses pengumpulan, penggunaan data untuk menentukan aturan atau pola tertentu dalam jumlah data yang sangat besar. KDD secara garis besar meliputi selection, preprocessing/cleansing, transformation, data mining dan evaluasi. Adapun penjelasannya sebagai berikut:



Gambar 1. Knowledge Discovery In Database (KDD)

3.1. Selection

Pada tahap sebelum pengolahan data harus dilakukan seleksi data. Data diperoleh dari data stok barang toko sembako. Data tersebut mencakup atribut dan stok barang. Data dipilih sesuai kebutuhan *Knowledge Discovery in Database (KDD)*. Sehingga nantinya tidak semua atribut yang diperoleh pada saat pengumpulan data digunakan.

3.2. Pre-processing

Pada tahapan ini, hanya fokus pada data yang akan diolah. Ini juga termasuk membersihkan data untuk memperbaiki duplikat, nilai yang hilang, dan memperbaiki kesalahan seperti kesalahan ketik.

3.3. Transformation

Pada tahapan ini data yang dipilih harus dimodifikasi atau disesuaikan agar sesuai dengan data mining Algoritma *K-Means*. Selanjutnya, pemilihan atribut paling sesuai analisis *K-Means*.

3.4. Data Mining

Proses data mining menggunakan metode Algoritma *K-Means Clustering* untuk mengelompokkan barang yang sudah terjual dan masih ada dalam stok berdasarkan karakteristik atau atribut yang sudah ditentukan. Selanjutnya, menentukan jumlah kelompok atau cluster yang optimal untuk dianalisis.

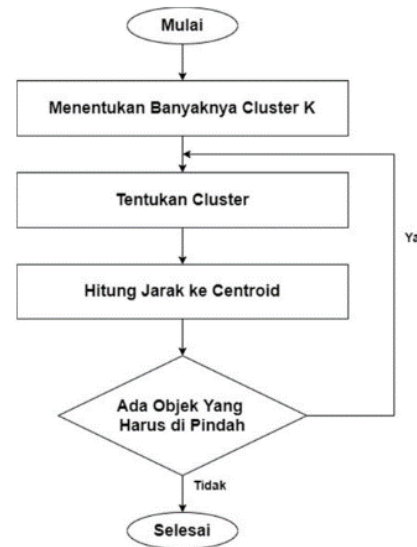
3.5. Interpretation / Evaluation

Tahapan ini merupakan tahap akhir dari proses data mining. Tahap ini dilakukan setelah data diolah dengan algoritma yang dipilih menggunakan *tools rapid miner* dan tujuannya adalah untuk melakukan evaluasi terhadap hasil proses data mining.

4. HASIL DAN PEMBAHASAN

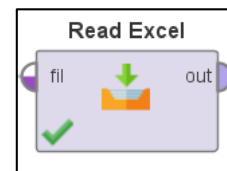
4.1. Hasil Penelitian

Dalam temuan pembahasan ini, proses penggunaan Algoritma *K-Means* untuk mengklasifikasikan atau mengelompokkan dataset stok barang akan dijelaskan. Pengelompokkan ini dilakukan melalui proses pengujian yang dilakukan oleh *software RapidMiner Studio*. Untuk mengelompokkan stok barang yang masih ada dalam stok dan yang sudah terjual atau tidak ada. dengan algoritma *k-means* maka secara umum alur yang digunakan seperti gambar di bawah ini:

Gambar 2. Tahapan Algoritma *K-means Clustering*

4.2. Pengujian Data

Pengujian data dilakukan untuk mengetahui mengapa penelitian ini perlu dilakukan. Pengumpulan data terjadi langsung di toko sembako. Data yang digunakan sebagai bahan penelitian adalah data laporan stok selama 3 bulan. Data ini mencakup total 503 *item* data persediaan seperti nama barang, kode barang, jenis barang, satuan ukuran, stok awal, stok terjual, stok akhir,

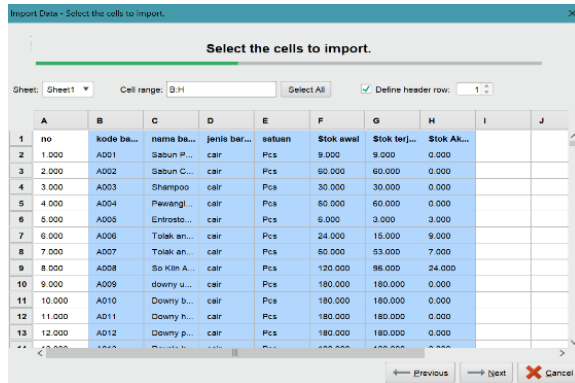
Gambar 3. Read Excel pada *tools Rapidminer*

Tabel 1. Data historis stok barang

No	Kode Barang	Nama Barang	Jenis Barang	Satuan	Stok Awal	Stok terjual	Stok Akhir
1.	A001	Sabun Pakaian	Cair	Pcs	9	9	0
2.	A002	Sabun Cuci Piring	Cair	Pcs	60	60	0
3.	A003	Shampoo	Cair	Pcs	30	30	0
4.	A004	Pewangi Pakaian	Cair	Pcs	60	60	0
5.	A005	Entrost opc air anak	Cair	Pcs	6	3	3
...	...			...	...	...	...
503	A503	Bolu Tape	Rentan Kadalua rsa	Pcs	200	180	20

4.3. Selection Data

Data penjualan toko sembako untuk bulan agustus-hingga oktober 2023 dipilih untuk studi ini. Informasi ini terdiri dari 503 item dan delapan atribut. Karena masing-masing dari keempat atribut memiliki nilai yang sama, satu atribut merupakan hasil seleksi yang tidak digunakan.



Gambar 4. selection data tools rapidminer

4.4. Preprocessing

Pra-pemrosesan data adalah serangkaian Langkah atau proses yang dilakukan terhadap data sebelumnya diolah dan digunakan dalam analisis atau pengolahan proses ini merupakan Langkah awal dalam tahapan pra-pemrosesan data dan memiliki peran penting menentukan kualitas dan relevansi data yang akan digunakan. Untuk meningkatkan akurasi dan kualitas hasil data mining, data selama tiga bulan dikumpulkan dan dipilih untuk penelitian ini. Dilihat berdasarkan data statistik diatas tidak ada missing, yang artinya tidak ada proses *processing*.

Name	Type	Missing	Statistics	Filter (6 attributes)	Search for Attributes
nama barang	Polynomial	0	Least Terasi Cap jempol (1)	ABC jeruk nipis (1)	ABC
cluster	Nominal	0	Small cluster_1 (37)	Small cluster_0 (466)	Small clust
satuan = Pcs	Integer	0	Min 1	Max 1	Mean 1
Stok awal	Integer	0	Min 1	Max 576	Mean 62.7
Stok terjual	Integer	0	Min 0	Max 576	Mean 60.0
Stok Akhir	Integer	0	Min 0	Max 97	Mean 2.69

Gambar 5. Preprocessing

4.5. Transformasi Data

Operator *nominal to numerical* digunakan untuk mengubah data *nominal* menjadi *numerik*.

Penjelasan: Hasil proses *rapidminer* operator terdiri dari 503 contoh, dua atribut khusus, empat atribut reguler, termasuk kategori satuan, stok awal, stok terjual, dan stok akhir.

Row No.	nama barang	satuan = Pcs	Stok awal	Stok terjual	Stok Akhir	cluster
1	Sabun Palsu	1	9	9	0	cluster_0
2	Sabun Cuci P...	1	60	60	0	cluster_0
3	Shampoo	1	30	30	0	cluster_0
4	Pewangi Palsu	1	60	60	0	cluster_0
5	Entrostic...	1	6	3	3	cluster_0
6	Tolak angin...	1	24	15	9	cluster_0
7	Tolak angin...	1	60	63	7	cluster_0
8	So Klin A...	1	120	96	24	cluster_0
9	Downy ungu...	1	180	180	0	cluster_0
10	Downy biru...	1	180	180	0	cluster_0
11	Downy Hitam	1	180	180	0	cluster_0
12	Downy putih	1	180	180	0	cluster_0

Gambar 6. Transformasi Data

4.6. Data Mining

tahapan ini dilakukan pemodelan data, metode yang digunakan pada penelitian adalah *K-Means clustering*. Data yang telah dikumpulkan, diseleksi dan ditransformasi akan dianalisis dengan menggunakan software *RapidMiner*. Data stok penjualan memiliki 503 item data, Data mining adalah proses menemukan pola atau informasi yang menarik pada data tertentu dengan menggunakan algoritma *k-means*.

4.6.1. Read Excel pada Tools RapidMiner (pengujian data/selection)

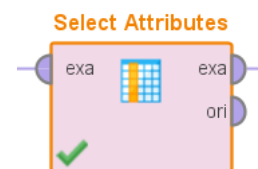


Gambar 7. Read excel pada tools rapidminer

Tahap ini merupakan tahap pengujian data dengan menggunakan *tools rapidminer* dengan ketik *read excel* pada operator lalu masukan dataset dengan klik *import configuration wizard*.

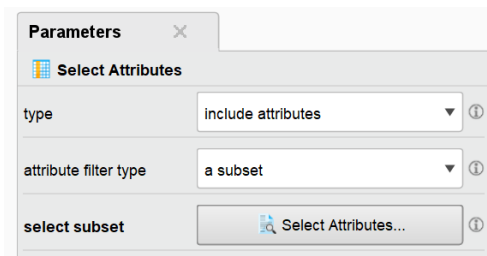
4.6.2. Selection(pre-processing)

Setelah pengujian data pengumpulan data tahapan selanjutnya yaitu seleksi data menggunakan *tools rapid miner* dengan klik *select attribute* pada *operators* lalu klik dua kali, lalu pada parameters di klik *attribute filter type>a subset>select subset>button select attribute*.



Gambar 8. Selection

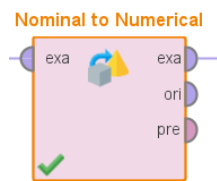
Setelah itu ubah *parameters* yang ada di kanan, *attribute filter type>A subset>Select subset>select attribute* yang tak digunakan.



Gambar 9. parameters select attribute

4.6.3. Nominal To Numerical (Transformasi)

Pada tahapan ini merupakan tahapan transformasi data, dengan mengubah kategori menjadi numerik, dengan klik *nominal to numerical* pada operator.

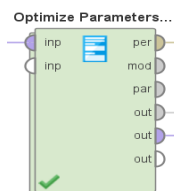


Gambar 10. transformasi data

4.6.4. Data Mining

Pada tahapan ini menggunakan 3 operators yaitu *optimize parameters (Grid)*, *K-means Clustering*, *cluster distance performance*.

4.6.4. Optimize parameters (Grid) digunakan untuk mengoptimalkan nilai Dbi.



Gambar 9. Optimize parameters

Tabel 2. Parameter optimize

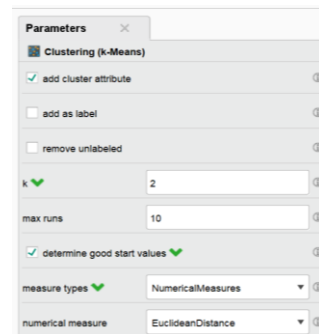
No	Operator	Parameters
1	Clustering	<i>K, numerical_measure</i>
	K	Min=2.0 Max=10.0 Steps=10
2	Performance (cluster distance performance)	<i>Eaclidean distance manhattan distance</i>

4.6.5. K-means Clustering

digunakan untuk menentukan nilai *k(cluster)* dari data yang dipilih dengan ketik *k-means clustering* lalu drag, dengan parameters *k=2*, *max run=10*, *measure type>numericalmeasure*, *numerical measure>eaclidean distance*.



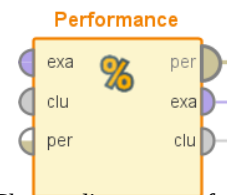
Gambar 1. K-means Clustering



Gambar 11. Parameters k-means clustering

4.6.6. Cluster distance performance

digunakan untuk menampilkan nilai *davies boouldin index (Dbi)*. *Cluster distance performance>main criterion>davies bouldin*.



Gambar 12. Cluster distance performance

4.7. Interpretasi/Evaluasi

Dari 503 data stok barang, dua klaster telah diidentifikasi menggunakan temuan *Software RapidMiner*. 2 cluster diperoleh isi pada *attribute* (satuan=pcs, stok awal, stok terjual, stok akhir), lalu pada *cluster_0* satuan=pcs merupakan satuan barang, dan stok awal senilai 37.766, dan stok terjualnya senilai 35.279, stok akhirnya 2.487, sedangkan pada *cluster_1* satuan =pcs merupakan satuan barang, dan stok awal senilai 377.676, stok terjualnya senilai 372.324, stok akhirnya senilai 5.351. dengan nilai *Davies Bouldin Index (Dbi)* senilai -0.330. maka bisa kita liat bahwa nilai terbesar berada pada *cluster_1* dengan nilai stok akhirnya 5.351

Attribute	cluster_0	cluster_1
satuan = Pcs	1	1
Stok awal	37.766	377.676
Stok terjual	35.279	372.324
Stok Akhir	2.487	5.351

Gambar 13. Centroid Table

Davies Bouldin

Davies Bouldin: -0.330

Gambar 14. Hasil Davies Bouldin Index (Dbi)

5. KESIMPULAN DAN SARAN

Berdasarkan hasil pengelompokkan data menggunakan metode *K-Means clustering* dengan menerapkan metode KDD pada toko sembako tersebut, ditemukan dua *cluster* utama, yaitu *cluster_0* dengan 466 *item* dan *cluster_1* dengan 37 *item*. *Cluster_0* kemungkinan besar mencakup barang-barang yang memiliki pola penjualan serupa, sementara *cluster_1* mungkin mencakup barang-barang dengan karakteristik penjualan yang berbeda. Dalam evaluasi kualitas pengelompokkan, *Davies Bouldin index(Dbi)* yang dihasilkan sebesar 0,330 mendekati nilai 1, menunjukkan bahwa pengelompokkan ini cukup baik. Hasil ini mengindikasikan bahwa perbedaan antara pusat kluster relatif kecil dibandingkan dengan ukuran dan varietas dalam setiap kluster, sehingga *cluster* yang dihasilkan dapat dianggap sama. Dengan demikian, dapat disimpulkan bahwa penerapan metode *K-Means clustering* dengan metode KDD pada data historis stok penjualan toko sembako memberikan pemahaman yang lebih baik terhadap pola penjualan barang. Pengelompokkan ini dapat menjadi landasan untuk pengembangan strategi manajemen stok yang lebih efisien dan efektif guna meningkatkan ketersediaan barang di toko sembako tersebut. Diperlukan tindakan lanjutan berupa perbaikan manajemen stok dan penerapan strategi yang lebih terarah berdasarkan temuan dari hasil pengelompokkan ini guna meningkatkan kinerja toko sembako secara keseluruhan.

DAFTAR PUSTAKA

- [1] I. Pii, N. Suarna, and N. Rahaningsih, "PENERAPAN DATA MINING PADA PENJUALAN PRODUK PAKAIAN DAMEYRA FASHION MENGGUNAKAN METODE K-MEANS CLUSTERING," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 1. LPPM ITN Malang, pp. 423–430, 2023. doi: 10.36040/jati.v7i1.6336.
- [2] H. Prastiwi, J. Pricilia, and E. Rasywir, "Implementasi Data Mining Untuk Menentukn Persediaan Stok Barang Di Mini Market Menggunakan Metode K-Means Clustering," *Jurnal Informatika Dan Rekayasa Komputer(JAKAKOM)*, vol. 2, no. 1. LPPM STIKOM Dinamika Bangsa Jambi, pp. 141–148, 2022. doi: 10.33998/jakakom.2022.2.1.34.
- [3] F. Indriyani and E. Irfiani, "Clustering Data Penjualan pada Toko Perlengkapan Outdoor Menggunakan Metode K-Means," *JUITA : Jurnal Informatika*, vol. 7, no. 2. Lembaga Publikasi Ilmiah dan Penerbitan Universitas Muhammadiyah Purwokerto, p. 109, 2019. doi: 10.30595/juita.v7i2.5529.
- [4] P. Rahayu, I. Anikah, D. B. Saputra, T. Anelia, and Martanto, "Penerapan Data Mining Metode K-Means Clustering Untuk Analisa Penjualan Rotan," *KOPERTIP : Jurnal Ilmiah Manajemen Informatika dan Komputer*, vol. 4, no. 2. Puslitbang KOPERTIP Indonesia, pp. 42–50, 2020. doi: 10.32485/kopertip.v4i2.118.
- [5] M. R. Sahputra, E. Rahayu, and N. Nurjamiyah, "Penerapan Metode Reorder Point pada Persediaan Stok Barang Berbasis Website," *JITEKH*, vol. 10, no. 2. Universitas Harapan Medan, pp. 68–74, 2022. doi: 10.35447/jitekh.v10i2.579.
- [6] Puji Rahayu, Ika Anikah, Dias Bayu Saputra, Tri Anelia, and Martanto, "Penerapan Data Mining Metode K-Means Clustering Untuk Analisa Penjualan Rotan," *KOPERTIP J. Ilm. Manaj. Inform. dan Komput.*, vol. 4, no. 2, pp. 42–50, 2020, doi: 10.32485/kopertip.v4i2.118.
- [7] N. Noviati, M. Mulyawan, D. A. Kurnia, and A. R. Rinaldi, "Clustering Data Penjualan Produk Makanan pada Toko Toserba Yogya Siliwangi dengan Menggunakan Metode K-Means," *MEANS (Media Informasi Analisa dan Sistem)*. Universitas Katolik Santo Thomas, pp. 77–84, 2022. doi: 10.54367/means.v7i1.1850.
- [8] A. Chintia Devi, H. Zulfia Zahro', and N. Vendyansyah, "Penerapan Metode K-Means Clustering Untuk Pengelompokan Data Barang Penjualan Berbasis Web Pada Koperasi Pt. X," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 4, no. 1, pp. 319–324, 2020, doi: 10.36040/jati.v4i1.2337.
- [9] A. MARDIANA and B. Budiman, "IMPLEMENTASI METODE K-MEANS UNTUK CLUSTERING PRODUK YANG KURANG DIMINATI BERDASARKAN DATA PENJUALAN," *SEMINAR TEKNOLOGI MAJALENGKA (STIMA)*, vol. 6. Universitas Majalengka, pp. 198–210, 2022. doi: 10.31949/stima.v6i0.697.
- [10] R. Nofitri and N. Irawati, "Integrasi Metode Neive Bayes Dan Software Rapidminer Dalam Analisis Hasil Usaha Perusahaan Dagang," *JURTEKSI (Jurnal Teknol. dan Sist. ...)*, 2019, [Online]. Available: <https://jurnal.stmikroyal.ac.id/index.php/jurteks/article/view/393>