

PENERAPAN ALGORITMA K-MEANS CLUSTERING UNTUK MENENTUKAN PERSEDIAAN STOK BARANG

Diana Indriani, Bambang Irawan, Agus Bahtiar

Teknik Informatika STMIK IKMI Cirebon
Jalan Perjuangan No 10 B Cirebon, Indonesia
dianaindriani02@gmail.com

ABSTRAK

Istilah *fashion* umumnya mengacu pada perubahan tren dan gaya berpakaian yang mendominasi pada suatu periode tertentu. Bidang bisnis *fashion* berusaha untuk inovatif dengan menghadirkan produk-produk yang menarik, mengikuti tren terbaru, dan memuaskan selera konsumen. Proses pemeriksaan stok saat ini masih bersifat manual dengan memeriksa ketersediaan barang, sehingga tidak efisien ketika ada permintaan tinggi dan persediaan tidak mencukupi untuk memenuhi kebutuhan konsumen. Untuk mengatasi permasalahan tersebut diperlukan penerapan data *mining* untuk menghasilkan informasi yang dapat digunakan dalam pengelolaan persediaan barang. Proses pengolahan data dilakukan dengan menerapkan teknik *k-means clustering*. Tujuan dari penelitian ini dapat mengetahui barang-barang yang telah terjual dan barang-barang yang masih tersedia dalam stok. Langkah pertama dalam penelitian ini adalah pengumpulan data. Berdasarkan hasil analisis menggunakan *Davies Bouldin Index (DBI)*, ditemukan bahwa konfigurasi optimal untuk pembagian data menjadi *cluster* adalah pada $k=2$, dengan nilai *DBI* mencapai minimum sebesar -0.142. Hasil ini mengindikasikan bahwa pembentukan dua *cluster* ($k=2$) memberikan nilai *DBI* terendah. Setelah mengevaluasi hasil pemrosesan data ditemukan bahwa *Cluster_0* memiliki jumlah penjualan yang tinggi, mencapai 594 *item*. Dari hasil analisis *cluster*, dapat disimpulkan bahwa *Cluster_0* terdiri dari produk yang diminati mengindikasikan *Cluster_0* menjadi fokus utama karena memiliki jumlah *item* yang jauh lebih banyak dan dianggap sebagai produk paling diminati.

Kata kunci : Data Mining, K-Means Clustering, Stok

1. PENDAHULUAN

Dalam era globalisasi, perkembangan teknologi berlangsung dengan cepat. Kemajuan teknologi di seluruh dunia telah memberikan kontribusi besar dalam mempermudah tugas manusia, sehingga membuatnya lebih efisien di berbagai sektor. Bidang bisnis terus bergerak maju seiring berjalannya waktu, salah satu contoh teknologi yang bermanfaat di bisnis *fashion* adalah memfasilitasi pengelolaan persediaan, yang merupakan aspek penting dalam manajemen yang efektif. Dengan demikian, manajemen yang kompeten dalam mengatur stok barang menjadi sangat penting untuk menghindari penumpukan barang yang berlebihan dan barang yang kurang diminati oleh konsumen.

Proses pemeriksaan stok saat ini masih bersifat manual dengan memeriksa ketersediaan barang, sehingga menjadi tidak efisien ketika ada permintaan tinggi dan persediaan tidak mencukupi untuk memenuhi kebutuhan konsumen. Untuk mengatasi permasalahan ini diperlukan penerapan metode data *mining* untuk menghasilkan informasi yang dapat digunakan dalam pengelolaan persediaan barang. Proses pengolahan data dapat dilakukan dengan menerapkan teknik pengelompokan data, seperti *k-means clustering*. *K-Means Clustering* bertujuan untuk mengelompokkan data dengan karakteristik serupa ke dalam kelompok yang sama, sementara data dengan karakteristik yang berbeda akan dikelompokkan ke dalam kelompok yang berbeda.

Dalam studi yang serupa[1] Permasalahan yang dihadapi oleh PT. SAS Group yaitu berkaitan dengan

ketersediaan barang yang merupakan suatu elemen yang sangat penting, sehingga sebagai manajemen yang baik dalam proses mengatur ketersediaan stok barang sangat diperlukan, untuk menghindari penumpukan barang yang sama dan barang yang kurang diminati oleh pembeli dengan menggunakan data *mining*. Tujuan dari penelitian ini adalah memudahkan PT. SAS Group dalam menentukan stok persediaan barang. Hasil dari penelitian ini yaitu setiap *cluster* yaitu *cluster_0* sebanyak 1282 atau sebesar 98,6%, *cluster_1* sebanyak 18 atau sebesar 1,4%. Berdasarkan hal tersebut didapat kelompok stok barang yang berada pada *cluster_0* memiliki rata-rata penjualan yang besar yaitu mencapai 2368 barang. Artinya bahwa kelompok barang pada *cluster_0* perlu menyimpan banyak barang karena merupakan produk yang paling banyak diminati. Menurut[2] Pada kebanyakan minimarket yang ada di jambi saat ini belum mempunyai metode baku yang diterapkan. Pemeriksaan stok hanya dilakukan dengan memeriksa persediaan barang secara manual. Sehingga hal ini kurang efisien jika konsumen membutuhkan barang dalam jumlah banyak dan stok tidak cukup untuk memenuhi kebutuhan konsumen. 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*. Dengan dilakukan pengelompokan data tersebut dapat diketahui barang apa saja yang

sudah terjual dan barang apa saja yang masih ada dalam stok. Pengumpulan data dilakukan pada sebuah minimarket di daerah jambi pada bulan Desember tahun 2021. Barang dalam minimarket dikelompokkan kedalam 20. Data tersebut diolah untuk menciptakan pengetahuan yang dapat digunakan sebagai strategi untuk mengembangkan strategi pengelolaan persediaan barang. Data awal yang diperoleh dari minimarket disimpan sebagai data semua penjualan obat selama 3 bulan terakhir. [3] dari temuan penelitian ini menggunakan data penjualan yang diambil pada toko SS *baby shop* namun permasalahannya, persediaan barang pada Toko SS *Baby Shop* masih dilakukan secara manual, yaitu pemilik hanya melihat barang yang habis saja tanpa mempertimbangkan barang mana yang lebih laku terjual dan barang mana yang lakunya lambat. Padahal mengatur ketersediaan barang sesuai permintaan dan kebutuhan konsumen dapat mengurangi nilai kerugian dari menumpuknya stok barang tidak laku terjual. Sehingga sebagai manajemen yang baik dalam proses mengatur persediaan stok barang sangat diperlukan, untuk menghindari penumpukan barang yang sama dan barang yang kurang diminati oleh konsumen. Untuk mengatasi permasalahan tersebut maka digunakan penerapan *K-Means Clustering*, karena permasalahan yang sering terjadi pada toko mengalami kesulitan pada saat menentukan stok minimum pada tiap barang yang harus dipenuhi berdasarkan keinginan konsumen. Penelitian ini menggunakan data penjualan yang diambil pada tokoh SS *baby shop* dengan jumlah data yang digunakan 203 dan mempunyai 3 atribut yaitu jumlah transaksi jumlah stok dan rata-rata penjualan. Metode yang digunakan dalam penelitian ini adalah *K-Means Clustering* dari perhitungan yang telah dilakukan maka direkomendasikan 5 klaster untuk menentukan persediaan barang yaitu *Cluster 1* yang dikategorikan paling banyak *Cluster 2* yang dikategorikan terjual sedang *Cluster 3* yang dikategorikan terjual paling sedikit *Cluster 4* yang di kategorikan terjual sedikit dan *cluster 5* dikatagorikan terjual banyak, dan hasil perbandingan perhitungan manual dan SPSS berbeda. Hal ini dikarenakan pemilihan pusat *Cluster* nya dipilih secara random hasil. Tetapi, walaupun telah banyak studi yang telah dilakukan, masih ada kebutuhan untuk riset lebih lanjut, dengan penelitian ini akan berfokus pada peningkatan manajemen persediaan stok barang yang lebih tepat.

Sebagai pemilik usaha, penting untuk merancang strategi manajemen persediaan barang dengan baik untuk mencegah kerugian yang disebabkan oleh penumpukan barang. Oleh karena itu, tujuan dari penelitian ini dapat mengetahui barang-barang yang telah terjual dan barang-barang yang masih tersedia dalam stok. Ketersediaan stok yang tepat dan waktu pengisian kembali yang efisien menjadi kunci keberhasilan bisnis ritel dalam meningkatkan pendapatan dan kepuasan konsumen.

Kumpulan data yang besar memiliki potensi untuk memberikan informasi yang berharga ketika diolah dengan benar. Metode untuk mendapatkan informasi dari data adalah melalui proses data *mining*. Data *mining* adalah proses ekstraksi pengetahuan atau informasi yang bermanfaat dari kumpulan data besar atau dataset. Salah satu metode yang efektif untuk mengolah data tersebut adalah dengan menggunakan *k-means clustering*. Data ini diperoleh melalui survei langsung, wawancara, dan dokumentasi data transaksi penjualan. Data tersebut mencakup rincian mengenai nama barang (artikel), ukuran, jumlah penjualan, stok awal, stok akhir, dan atribut lain yang relevan.

Hasil penelitian ini diharapkan dapat memberikan wawasan yang berharga mengenai persediaan stok barang *fashion*. Dengan menggunakan teknik data *mining K-Means Clustering*, dapat membuat keputusan yang lebih tepat dalam mengelola persediaan stok barang *fashion* yang dapat meningkatkan efisiensi operasional, mengurangi risiko barang yang tidak terjual atau kehabisan stok yang populer serta memaksimalkan keuntungan melalui manajemen persediaan yang lebih baik.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Kata *Mining* merupakan kiasan dari bahasa inggris, *mine*. Jika *mine* berarti menambang sumber daya yang tersembunyi di dalam tanah, maka Data *Mining* merupakan penggalian makna yang tersembunyi dari kumpulan data yang sangat besar.[4]

2.2. Clustering

Clustering adalah metode pengelompokan bagian atau kumpulan data atau objek sehingga objek yang sejenis menjadi *cluster* atau kelompok.[5]

2.3. Algoritma K-Means

K-Means yakni merupakan algoritma klustering dengan menentukan sejumlah data untuk di klaster dalam kesamaan karakteristik dan memaksimalkan perbedaan antar klaster. *K-Means* algoritma salah satu algoritma *clustering* yang paling banyak orang pakai untuk mengkategorikan sebuah produk.[6]

2.4. Persediaan

Persediaan menurut *Stice* dan *Skousen* adalah istilah untuk barang-barang yang dijual selama kegiatan bisnis normal perusahaan, atau secara langsung atau tidak langsung terkandung dalam barang yang diproduksi dan dijual kemudian.[1]

2.5. Fashion

Kata Sansekerta "*bhusana*" adalah asal kata "*fashion*". Namun, konsep berpakaian dan pakaian berbeda. Segala sesuatu yang kita kenakan, dari ujung rambut sampai ujung kaki, adalah pakaian.[7]

2.6. RapidMiner

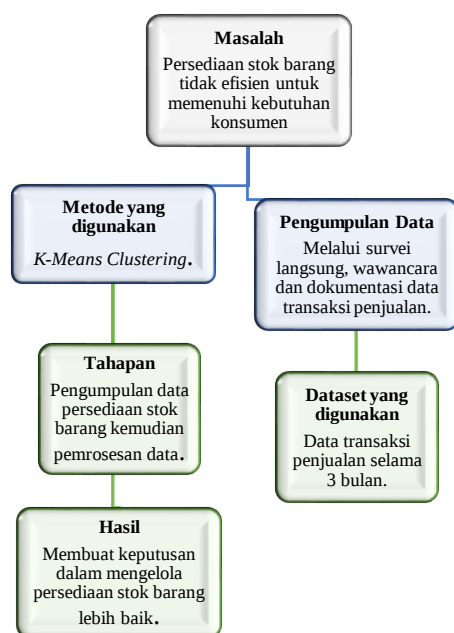
RapidMiner adalah perangkat lunak sumber terbuka. *RapidMiner* adalah solusi untuk penambangan data analitik, penambangan teks, dan analisis prediktif. *RapidMiner* menggunakan berbagai teknik deskriptif dan prediktif untuk memberikan wawasan kepada pengguna sehingga mereka dapat membuat keputusan terbaik.[1]

2.7. Davies Bouldin Index

Davies-Bouldin Index merupakan metode yang digunakan untuk mengukur validitas *cluster* pada suatu metode pengelompokan, kohesi didefinisikan sebagai jumlah dari kedekatan data terhadap titik pusat *cluster* dari *cluster* yang di ikuti. Sedangkan separasi didasarkan pada jarak antar titik pusat *cluster* terhadap *cluster*nya. Pengukuran dengan *Davies-Bouldin Index* ini memaksimalkan jarak *inter-cluster* antara *cluster* C_i dan C_j dan pada waktu yang sama mencoba untuk meminimalkan jarak antar titik dalam sebuah *cluster*.[8]

3. METODE PENELITIAN

Dalam melaksanakan sebuah penelitian, diperlukan suatu prosedur penelitian agar pelaksanaan penelitian tersebut dapat berjalan secara efektif.



Gambar 1. Tahapan Penelitian

Tahapan penelitian diatas dapat dijelaskan sebagai berikut :

1. Masalah

Dalam penelitian ini permasalahan yang dihadapi yaitu pemeriksaan stok masih bersifat manual dan Persediaan stok barang

tidak efisien untuk memenuhi kebutuhan konsumen.

2. Pengumpulan Data

Yaitu Melalui survei langsung, wawancara dan dokumentasi data transaksi penjualan.

3. Dataset yang digunakan

Data transaksi penjualan selama 3 bulan.

4. Metode yang digunakan

K-Means Clustering dengan menggunakan aplikasi *RapidMiner*.

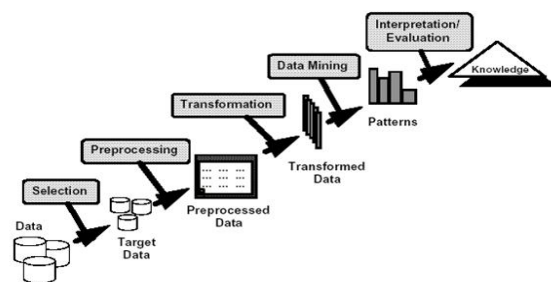
5. Tahapan

Pengumpulan data persediaan stok barang kemudian pemrosesan data.

6. Hasil

Membuat keputusan dalam mengelola persediaan stok barang lebih baik.

Penelitian ini dilakukan dengan menerapkan metode *KDD* (*Knowledge Discovery in Databases*) dan *k-means clustering*. Oleh karena itu melalui penerapan algoritma *k-means clustering* dalam kerangka *Knowledge Discovery in Databases* (*KDD*), untuk memberikan pemahaman yang lebih mendalam tentang struktur dan distribusi stok barang *fashion*. Tahapan *KDD* :



Gambar 2. Tahapan KDD

3.1. Data Selection

Langkah pertama sebelum memulai proses ekstraksi informasi dalam data *mining* adalah memilih (*select*) data yang relevan dari sekumpulan data operasional. Proses data *selection* dapat melibatkan pemilihan kolom tertentu, baris yang memenuhi kriteria tertentu, untuk fokus pada informasi yang paling relevan dalam pengolahan data.

3.2. Data Preprocessing

Membersihkan, merapikan, dan mempersiapkan data agar dapat diolah dengan lebih efektif.

3.3. Data Transformation

Suatu proses yang dilakukan untuk membuat data terpilih sesuai untuk proses data *mining*. Meningkatkan kualitas data, mengatasi masalah spesifik, atau membuat data lebih sesuai.

3.4. Data Mining

Proses menemukan pola dan informasi menarik pada data terpilih dengan menggunakan teknik dan

metode khusus (*k-means clustering*) Mengelompokkan data ke dalam beberapa klaster berdasarkan kemiripan antara titik-titik data.

3.5. Interpretasi/Evaluasi

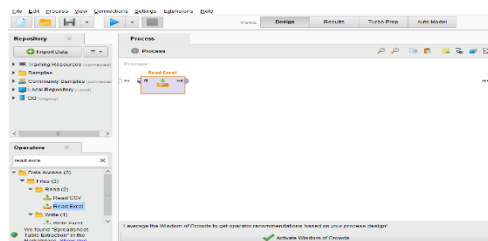
Evaluasi kinerja algoritma, Memahami, mengevaluasi, dan memberikan wawasan terhadap aktivitas yang dijelaskan. Hasil pola informasi yang dihasilkan dari proses data mining perlu disajikan dalam format yang dapat dipahami oleh pihak-pihak yang berkepentingan.[9]

4. HASIL DAN PEMBAHASAN

Proses pengelompokan dataset penjualan *fashion* dengan menggunakan algoritma *K-Means* dijelaskan pada hasil pembahasan kali ini. Pengelompokan ini dilakukan melalui metode pengujian pembelajaran mesin yang disebut *RapidMiner*. Beberapa tujuan yang mendasari analisis *cluster* terhadap data inventaris menggunakan *K-means* menemukan barang terjual dan masih tersedia, ketersediaan inventaris yang memadai, dan waktu pengisian ulang yang efisien.

a. Pengujian Data

Pada pengujian ini menggunakan aplikasi *RapidMiner Studio* versi 10.1, carilah operator *Read Excel*, lalu klik untuk mengimpor data yang berada dalam format *.xlsx*, yang diambil langsung dari *file Excel*. Data tersebut telah siap untuk dilakukan pengujian. Langkah pertama sebelum memulai proses ekstraksi informasi dalam data *mining* adalah memilih (*select*) data yang relevan dari sekumpulan data operasional.



Gambar 3. Read Excel

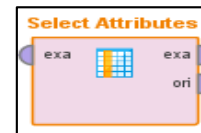
Seperti terlihat pada gambar di atas, ini adalah impor data yang berhasil karena tanda seru berwarna kuning pada tombol operator *Read Excel* telah hilang (operator telah siap mengolah data).

4.1. Data Selection

Data selection merujuk pada proses memilih *subset* atau bagian tertentu dari data yang akan digunakan dalam suatu analisis atau pengolahan. Proses ini merupakan langkah awal dalam tahapan pra-pemrosesan data dan memiliki peran penting dalam 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. Informasi ini mencakup 596 *item* dan 7 atribut.

4.2. Data Preprocessing

Pra-pemrosesan data adalah serangkaian langkah atau proses yang dilakukan terhadap data sebelum diolah atau digunakan dalam analisis atau pemodelan, terutama dalam konteks data ilmiah, statistik, dan pembelajaran mesin. Tahap pertama dari proses ini adalah pembersihan data, dimana variabel atau fitur yang sesuai dipilih dari kumpulan data untuk diproses lebih lanjut. Seperti pada gambar dibawah dengan menambahkan operator *Select Attributes*.

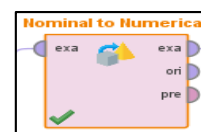


Gambar 4. Select Attributes

4.3. Data Transformation

Data transformation adalah proses mengubah format, struktur, atau nilai-nilai dalam dataset untuk membuatnya lebih sesuai atau berguna untuk analisis data. Tujuan utama dari *data transformation* adalah untuk meningkatkan kualitas data, memudahkan pemahaman, dan mendukung proses analisis lebih lanjut.

Dengan menambahkan operator *Nominal to Numerical* yaitu Proses mengubah format, struktur, atau nilai-nilai dalam dataset untuk membuatnya lebih sesuai atau berguna untuk analisis data.

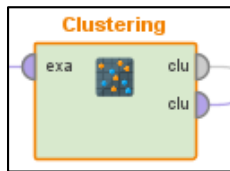


Gambar 5. Nominal to Numerical

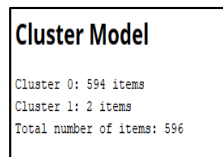
4.4. Data Mining

Data mining adalah proses penggalian data dari informasi yang sangat penting. *Data mining* juga merupakan proses eksplorasi pola dari data. Pola diperoleh dari berbagai jenis *database*, seperti *database* relasional, data *warehouse*, data transaksional, dan data berorientasi objek. Penggunaan data *mining* dapat membantu pengusaha mengambil keputusan dengan cepat dan akurat.[10] Tujuan utama dari data *mining* adalah untuk mengungkap hubungan-hubungan yang belum diketahui dalam data dan memberikan wawasan yang dapat mendukung pengambilan keputusan. Metode yang diterapkan dalam penelitian ini adalah *K-means clustering*.

Proses ekstraksi pengetahuan yang berguna, pola-pola tersembunyi, dan informasi berharga dari dataset yang besar atau kompleks. Selanjutnya, tambahkan operator pengelompokan *k-means*. Atur parameter *k* ke nilai 2, *max run* ke nilai 10, jenis pengukuran ke nilai *mixedmeasure*, *mixed Measure* ke nilai *Mixed EuclideanDistance*, dan langkah optimalisasi maksimum ke nilai 100.



Gambar 6. K-Means Clustering



Gambar 7. Hasil Cluster Model

Model *cluster* ditentukan dari hasil pengolahan operator *k-means clustering* dengan mengatur parameter $k = 2$, $max\ runs = 10$, $measure\ types = BregmanDivergences$ sehingga hasilnya dapat seperti gambar diatas, *cluster model* merupakan hasil pengelompokan data melalui algoritma *K-Means* yang diolah dan dibagi menjadi dua kelompok, sehingga menghasilkan *cluster 0* yang berjumlah 594 *item* dan *cluster 1* yang berjumlah 2 *item*.

Tabel 1. Cluster Optimal pada Davies Bouldin

Iterasi	Clustering.k	Davies Bouldin
1	2	-0.142
2	3	-0.330
3	4	-0.216
4	5	-0.353
5	6	-0.391
6	7	-0.527
7	8	-0.615
8	9	-0.579
9	10	-0.537

Pada hasil *Cluster Optimal* pada *Davies Bouldin Index (DBI)* terdapat pada $k=2$ dengan nilai *Davies Bouldin Index (DBI)* -0.142.

Langkah pertama menentukan jumlah *cluster* yang ingin dibentuk.[11] Melalui proses pemodelan ini, ditentukan dari hasil pengolahan operator *k-means clustering* dengan mengatur parameter $k = 2$, $max\ runs = 10$, $measure\ types = BregmanDivergences$ sehingga hasilnya berhasil mengelompokkan dataset yang terdiri dari total 596 data menjadi 2 kelompok. sehingga menghasilkan *cluster 0* yang berjumlah 594 *item* dan *cluster 1* yang berjumlah 2 *item*. Pada hasil *Cluster Optimal* pada *Davies Bouldin Index (DBI)*, ditemukan bahwa nilai *DBI* mencapai minimum pada $k=2$ dengan nilai -0.142. Hasil ini menunjukkan bahwa pembagian data menjadi dua *cluster* ($k=2$) memberikan nilai *DBI* terendah, yang dapat diinterpretasikan sebagai pembentukan *cluster* yang lebih baik dan lebih terpisah. Nilai *DBI* yang mendekati nol menunjukkan bahwa kedua *cluster* memiliki sejumlah besar perbedaan antara satu sama lain, dan $k=2$ dianggap sebagai konfigurasi optimal berdasarkan metrik ini.

kelompok persediaan *Cluster_1* tidak memerlukan penyimpanan persediaan dalam jumlah besar karena persediaan produknya berlebih. Dari hasil tersebut, dapat disimpulkan bahwa *Cluster_0* membutuhkan penyimpanan yang cukup besar karena termasuk produk yang paling diminati. Berdasarkan analisis clustering ini, dapat memanfaatkannya untuk secara efektif memeriksa persediaan produk guna memaksimalkan manajemen persediaan stok barang.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil analisis menggunakan *Davies Bouldin Index (DBI)*, ditemukan bahwa konfigurasi optimal untuk pembagian data menjadi *cluster* adalah pada $k=2$, dengan nilai *DBI* mencapai minimum sebesar -0.142. Hasil ini mengindikasikan bahwa pembentukan dua *cluster* ($k=2$) memberikan nilai *DBI* terendah. Setelah mengevaluasi hasil pemrosesan data ditemukan bahwa *Cluster_0* memiliki jumlah penjualan yang tinggi, mencapai 594 *item*. Dari hasil analisis *cluster*, dapat disimpulkan bahwa *Cluster_0* terdiri dari produk yang diminati mengindikasikan *Cluster_0* menjadi fokus utama karena memiliki jumlah *item* yang jauh lebih banyak dan dianggap sebagai produk paling diminati. Untuk penyelidikan lebih lanjut harus mempertimbangkan untuk menambahkan parameter atau variabel ke proses pengelompokan. Sebaiknya penelitian selanjutnya menggunakan dataset yang lebih besar untuk memperoleh data yang lebih akurat dan memperluas penggunaan atribut agar penelitian lebih komprehensif. Selain itu, penelitian di masa depan diharapkan dapat menambah jumlah data dan menggabungkannya dengan pendekatan yang berbeda serta harus mempertimbangkan penggunaan atribut tambahan untuk meningkatkan relevansi dan juga harus mencoba metode analisis selain pengelompokan *K-means*.

DAFTAR PUSTAKA

- [1] Ika Anikah, Agus Surip, Nela Puji Rahayu, Muhammad Harun Al- Musa, and Edi Tohidi, "Pengelompokan Data Barang Dengan Menggunakan Metode K-Means Untuk Menentukan Stok Persediaan Barang," *KOPERTIP J. Ilm. Manaj. Inform. dan Komput.*, vol. 4, no. 2, pp. 58–64, 2022, doi: 10.32485/kopertip.v4i2.120.
- [2] H. Pratiwi, Jeny Pricilia, and 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, doi: 10.33998/jakakom.2022.2.1.34.
- [3] J. Informatika et al., "Implementasi Metode K-Means Clustering Untuk Menentukan Persediaan Barang Pada Toko SS BabyShop Jurnal Informatika Dan Rekayasa Komputer (JAKAKOM)," vol. 3, no. September, pp. 712–

- 718, 2023.
- [4] P. P. Tjaya, R. Rino, and H. Wijaya, "Implementasi Metode Clustering K-Means Untuk Rekomendasi Pengadaan Stock Lampu Pada Pt Global Lighting Indonesia," *Algor*, vol. 3, no. 1, pp. 50–59, 2021, doi: 10.31253/algor.v3i1.628.
 - [5] D. A. Ramadhanty, R. Syafitri, E. Raswir, and D. Meisak, "Implementasi Data Mining Untuk Menentukan Persediaan Stok Obat Jurnal Informatika Dan Rekayasa Komputer (JAKAKOM)," *J. Inform. dan Rekayasa Komput.*, vol. 1, no. 2, pp. 155–160, 2022.
 - [6] S. I. Wahyudi and A. Wibowo, "Implementasi Metode K-Means Clustering Untuk Pengelompokan Data Stok Produk Toko Online Perdagangan Kaos," *Semin. Nas. Mhs. Fak. Teknol. Inf.*, no. September, pp. 470–478, 2022, [Online]. Available: <https://senafiti.budiluhur.ac.id/index.php>
 - [7] I. Pii, N. Suarna, and N. Rahaningsih, "PENERAPAN DATA MINING PADA PENJUALAN PRODUK PAKAIAN DAMEYRA FASHION MENGGUNAKAN METODE K-MEANS CLUSTERING," *JATI (Jurnal Mhs. Tek. ...)*, 2023, [Online]. Available: <https://ejournal.itn.ac.id/index.php/jati/article/view/6336>
 - [8] I. Safira, R. Salkiawati, and W. Priatna, "Penerapan Algoritma K-Means untuk Mengetahui Pola Persediaan Barang pada Toko Raja Bekasi," *J. Inform. ...*, 2022, [Online]. Available: <https://ejurnal.ubharajaya.ac.id/index.php/jiforty/article/view/1253>
 - [9] T. B. Pamungkas, S. Maesaroh, and 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/sainstek.v7i2.2697.
 - [10] G. Aprilianur and E. L. Hadisaputro, "Penerapan Data Mining Menggunakan Metode K-Means Clustering Untuk Analisa Penjualan Toko Myam Hijab Penajam," *JUPITER (Jurnal Penelit. Ilmu ...)*, 2022, [Online]. Available: <https://jurnal.polsri.ac.id/index.php/jupiter/article/view/4684>
 - [11] . 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.