

IMPLEMENTASI DATA MINING UNTUK MENENTUKAN POLA PENJUALAN PRODUK MENGGUNAKAN ALGORITMA K-MEANS CLUSTERING

Slamet Pujiono¹, Rini Astuti², Fadhil Muhamad Basysyar³

¹ Teknik Informatika, STMIK IKMI Cirebon

² Sistem Informasi, STMIK LIKMI Bandung

³ Sistem Informasi, STMIK IKMI Cirebon

Jalan Perjuangan No 10 B Cirebon, Indonesia

slametpujiono866@gmail.com

ABSTRAK

Penjualan *fashion* adalah proses pemasaran dan penjualan produk yang berkaitan dengan bisnis *fashion*. Bisnis *fashion* mencakup segala sesuatu yang berhubungan dengan pakaian, aksesoris, dan produk gaya hidup lainnya yang diproduksi dan dijual untuk memenuhi kebutuhan gaya dan penampilan konsumen. Dalam penelitian ini, pihak bersangkutan memilih untuk menjaga kerahasiaan identitasnya dan menggunakan nama samaran merek X. Karena beragamnya produk yang dijual dan tingkat minat konsumen untuk membelinya, mungkin mengalami kesulitan dalam mengisi persediaan produk. Oleh karena itu, dibutuhkan suatu sistem agar dapat membantu membuat keputusan dengan lebih cepat dan akurat. Untuk mengatasi permasalahan ini, analisis dilakukan dengan menerapkan metode *Clustering* menggunakan Algoritma *K-Means*. Penelitian ini bertujuan untuk menyelidiki strategi dan metode yang dapat diterapkan dalam pengelolaan stok barang lebih efisien, mengetahui barang-barang yang diminati guna meningkatkan pendapatan, serta menghindari akumulasi stok berlebihan yang dapat menyebabkan kerugian. Dari hasil penelitian ini diketahui nilai terbaik terdapat pada *cluster 2* dengan nilai *DBI* sebesar -0,310 Terdapat 616 *item* produk yang termasuk dalam kategori paling laris dalam *Cluster_0*, sementara 8 *item* produk kurang laris dalam *Cluster_1*. Hal ini memungkinkan pemangku kepentingan untuk mengambil tindakan yang lebih baik, diantaranya adalah mengefektifkan pengelolaan persediaan, mengamankan produk yang laris, dan memperluas promosi serta memperbaiki produk yang kurang laris.

Kata kunci : Data Mining, K-Means, Clustering, Produk Fashion

1. PENDAHULUAN

Di zaman digital saat ini, perkembangan teknologi telah mengalami perubahan besar dalam sektor bisnis ritel seluruh dunia. Hal ini mengindikasikan bahwa bisnis penjualan secara ritel memiliki potensi yang sangat besar untuk memenuhi beragam kebutuhan konsumen. Oleh karena itu, pemahaman mendalam tentang tren mode terkini, inovasi desain, dan strategi pemasaran yang efektif menjadi kunci kesuksesan dalam menghadapi persaingan.

Dalam konteks ini, perlu diperhatikan bahwa peluang dan tantangan dalam bisnis *fashion* terus berkembang seiring dengan perubahan selera dan preferensi konsumen yang terus berubah. Karena beragamnya produk yang dijual dan tingkat minat konsumen untuk membelinya, mungkin mengalami kesulitan dalam mengisi persediaan produk. Oleh karena itu, dibutuhkan suatu sistem agar dapat membantu membuat keputusan dengan lebih cepat dan akurat. Untuk mengatasi permasalahan ini, analisis dilakukan dengan menerapkan metode *Clustering* menggunakan Algoritma *K-Means*. Salah satu solusi dapat ditemukannya yaitu dengan menggunakan data transaksi sebagai landasan untuk perbaikan ini.

Penelitian sebelumnya dalam bidang *data mining* telah mengidentifikasi sejumlah pola penjualan penting dari pelanggan. Sebagai contoh, studi sebelumnya oleh [1] Perusahaan ini harus mampu memenuhi kebutuhan konsumen setiap harinya dan

mengambil keputusan yang tepat dalam menentukan strategi penjualan. Untuk mencapai hal tersebut, dealer mobil Luxor memerlukan strategi penjualan yang dapat menarik minat pembeli serta meningkatkan keuntungan dan pendapatan perusahaan. Permasalahan utama yang dihadapi Toko Variasi Mobil Luxor Gorontalo adalah sulitnya mengenali kecenderungan pembeli terhadap produk yang tersedia di Toko Variasi Mobil Luxor dan kesulitan dalam memenuhi kebutuhan pembeli yang terus berkembang dan berubah dalam menentukan strategi penjualan. Hasil akhir dari penelitian ini adalah mengelompokkan data tentang produk yang dijual dan mencari data probabilitas atau kecenderungan pelanggan untuk membeli produk tersebut. Hasil ini dapat dijadikan bahan pertimbangan dalam menentukan strategi penjualan. Artinya, dapat mengecualikan produk dengan posisi *cluster* terendah dan fokus pada produk dengan posisi *cluster* tertinggi. Dan menurut [2] CV Terang Jaya merupakan perusahaan yang bergerak di bidang otomotif yang menyediakan jasa dan layanan pembelian dan penjualan suku cadang otomotif berbagai merek otomotif. Namun, verifikasi produk mana yang dibutuhkan konsumen masih kurang, dan penyimpanan data tidak efektif. Untuk mengatasi permasalahan tersebut digunakan analisis *clustering* dengan menggunakan algoritma *K-Means*. Jadi dengan mengelompokkan data tersebut, bisa melihat produk mana yang paling laris, produk mana yang laris, dan produk mana yang kurang laris. Hal ini untuk

mencegah produk menumpuk di gudang. Hasil penelitian ini adalah 15 produk paling laris, 45 produk laris, dan 13 produk kurang laris. Adapun menurut[3] Data penjualan belum dimanfaatkan secara maksimal karena pemilik toko kesulitan membedakan antara produk sepatu yang sering dijual dan produk sepatu yang tidak diminati pelanggan. Tujuan dari penelitian ini adalah agar sistem dapat menggunakan teknik *K-Means* clustering dalam menentukan data untuk mengelompokkan jenis sepatu yang dijual. Tiga *cluster* digunakan yaitu tertinggi, sedang, dan terendah. Data yang akan dianalisis adalah data penjualan Toko Sepatu Ritelindo. Hasil *clustering* diperoleh 4 *item* yang masuk dalam kelompok penjualan sepatu tertinggi, 12 *item* yang masuk ke dalam kelompok penjualan sepatu sedang, dan 9 *item* yang masuk ke dalam kelompok penjualan sepatu yang terendah. Tetapi, meskipun telah banyak penelitian dilakukan, masih perlu untuk melakukan penelitian tambahan dalam merumuskan strategi penjualan yang lebih relevan.

Penelitian ini bertujuan untuk menyelidiki strategi dan metode yang dapat diterapkan dalam pengelolaan stok barang di gudang dengan lebih efisien, Tujuannya juga adalah untuk mengetahui barang-barang yang diminati guna meningkatkan pendapatan perusahaan, serta menghindari akumulasi stok berlebihan yang dapat menyebabkan kerugian. Oleh karena itu, dibutuhkan suatu sistem agar dapat membantu membuat keputusan dengan lebih cepat dan akurat.

Clustering adalah salah satu teknik yang termasuk dalam fungsi *data mining*. Algoritma *clustering* digunakan untuk mengelompokkan sejumlah data ke dalam kelompok-kelompok tertentu, yang disebut *cluster*. *Data mining*, yang sering disebut sebagai *knowledge discovery in database (KDD)*, adalah proses yang mencakup pengumpulan dan penggunaan data historis untuk mengidentifikasi pola, hubungan, atau keteraturan dalam data berukuran besar. Data ini mencakup nama produk, jumlah penjualan, stok awal, stok akhir, dan bulan.

Hasil penelitian ini diharapkan dapat memberikan kontribusi yang berarti dalam pemahaman pola penjualan pelanggan melalui pendekatan *data mining*. Harapannya adalah bahwa hasil penelitian ini akan membantu perusahaan meningkatkan kualitas pengambilan keputusan terkait strategi penjualan. Dengan menggunakan pengelompokan data ini, perusahaan dapat mengidentifikasi barang yang memiliki permintaan tinggi dan barang yang memiliki tingkat penjualan yang rendah, sehingga dapat menghindari penumpukan barang di gudang.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data Mining adalah ekstraksi pengetahuan dari *database*. *Data mining* adalah proses mengekstraksi dan mengidentifikasi informasi berguna dan

pengetahuan terikat dari database besar menggunakan teknik statistik, matematika, kecerdasan buatan, dan pembelajaran mesin. Berdasarkan definisi data *mining* yang dijelaskan diatas, data *mining* adalah pengetahuan yang disembunyikan dalam *database*, dan digunakan untuk mengekstrak dan mengubah informasi pengetahuan dari database dengan menemukan pola dan teknik dalam statistik matematika, kecerdasan buatan, dan pembelajaran mesin.[4]

2.2. Clustering

Proses *clustering* adalah suatu metode pengklasifikasian isyarat-isyarat dalam suatu kelompok dengan menghubungkan hipotesis-hipotesis indeks sedemikian rupa sehingga objek-objek dalam keseluruhan kelompok tidak termasuk objek-objek lainnya sangat mirip satu sama lain dan mempunyai tingkat heterogenitas yang sangat tinggi.[5]

2.3. Algoritma K-Means

K-Means adalah algoritma *clustering* yang digunakan dalam proses data *mining*. *K-Means* bekerja dengan cara mengelompokkan secara partisi, yaitu membagi data ke dalam kelompok-kelompok tertentu dengan meminimalkan jarak rata-rata setiap *cluster* data. Metode ini merupakan metode *clustering* atau pengelompokan yang paling terkenal dan banyak digunakan karena sederhana dan mudah diterapkan, dapat mengelompokkan data dalam jumlah besar, dan mempunyai kompleksitas waktu nol linier (nKT). Di sini, n adalah jumlah dokumen, K adalah jumlah *cluster*, dan T adalah jumlah iterasi.[6]

2.4. Penjualan

Untuk mencapai keuntungan jangka panjang dan menguntungkan baik bagi pembeli maupun penjual, penjual harus memenuhi semua persyaratan pembeli selama proses penjualan. Dalam dunia bisnis, pendapatan juga mengacu pada hasil pemberian suatu jasa.[7]

2.5. Rapidminer

RapidMiner, sebuah platform sumber terbuka, dirancang untuk analisis data, visualisasi, pemodelan prediktif, dan pelaporan. *Rapidminer* menyediakan alat untuk mengumpulkan, mengekstrak, mengubah, dan menganalisis data, serta membangun model prediktif dan menghasilkan laporan analitik. Dalam pekerjaannya, *RapidMiner* memberi pengguna wawasan melalui berbagai metode prediktif dan deskriptif sehingga mereka dapat membuat keputusan terbaik.[7]

2.6. Davies Bouldin Index (DBI)

Davies-Bouldin Index adalah cara untuk mengukur validitas *cluster* dalam metode pengelompokan. Kohesi didefinisikan sebagai jumlah kedekatan data dengan pusat *cluster* dari *cluster* yang dilacak. Pemisahan, sebaliknya, didasarkan pada jarak

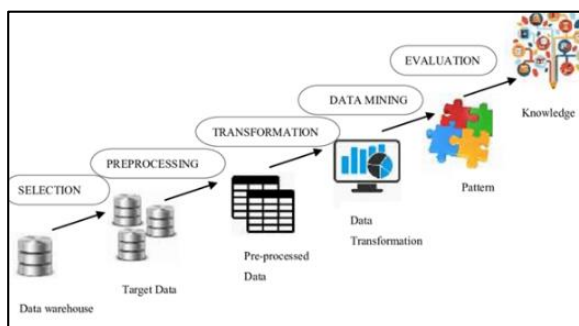
antara pusat *cluster* dan *cluster*. Ukuran yang menggunakan *Davies-Bouldin Index* ini memaksimalkan jarak antara *cluster* C_i dan C_j sambil mencoba meminimalkan jarak antar titik dalam sebuah *cluster*. Apabila jarak antar *cluster* maksimal berarti kesamaan fitur antar masing-masing *cluster* kecil, sehingga perbedaan antar *cluster* terlihat jelas. Nilai *DBI* terendah menunjukkan maka pembagian ke dalam *cluster* tersebut lebih baik atau lebih optimal, dengan batasan antar-*cluster* yang lebih jelas dan homogenitas yang lebih tinggi dalam setiap *cluster*. [8]

2.7. Knowledge Discovery In Database (KDD)

Knowledge Discovery in Database (KDD) merupakan suatu metode pencarian pengetahuan dan informasi yang belum diketahui dari suatu *database*. *KDD* adalah nama lain dari data *mining*. Meskipun kedua istilah tersebut sebenarnya memiliki konsep yang berbeda, namun keduanya saling berkaitan dan data *mining*, salah satu tahapan proses *KDD* secara keseluruhan, merupakan inti dari proses *KDD*. Data *mining* adalah suatu teknik untuk menemukan, mencari, dan mengekstrak informasi dan wawasan baru dari kumpulan data yang sangat besar dengan mengintegrasikan atau menggabungkannya dengan bidang ilmu lain seperti statistik, kecerdasan buatan, dan pembelajaran mesin kemudian menghasilkan informasi berguna. [4]

3. METODE PENELITIAN

Dalam penelitian ini, menggunakan algoritma *K-Means* untuk mengelompokkan data kategorikal. Tujuannya adalah untuk menciptakan *cluster* yang lebih stabil. Algoritma *clustering* digunakan untuk mengelompokkan kumpulan data ke dalam kelompok tertentu yang disebut *cluster*. Data *mining*, juga dikenal sebagai penemuan pengetahuan dalam *database (KDD)*, adalah proses pengumpulan dan penggunaan data historis untuk mengidentifikasi pola, hubungan, atau keteraturan dalam sejumlah besar data.



Gambar 1. Tahapan Knowledge Discovery in Database

3.1. Data Selection

Memilih atribut data yang akan digunakan dan memprioritaskan data untuk penemuan yang disimpan dalam file terpisah dari *database*. [9] Memilih data dan Mengimport dataset dalam format *excel* pada *RapidMiner*. Proses memilih subset atau bagian tertentu dari data yang relevan.

3.2. Data Preprocessing

Membersihkan, mengubah, dan menyesuaikan data yang digunakan maupun tidak.

3.3. Data Transformation

Dengan menggunakan operator nilai nominal ke numerik, data dapat diubah dari data bertipe nominal menjadi data numerik. Memilih metode pengkodean yang tepat saat mengubah variabel kategori menjadi variabel numerik.

3.4. Data Mining

Melakukan pengolahan data menggunakan metode *k-means clustering* untuk mengategorikan barang-barang menjadi 2 kelompok yang laris, dan kurang laris.

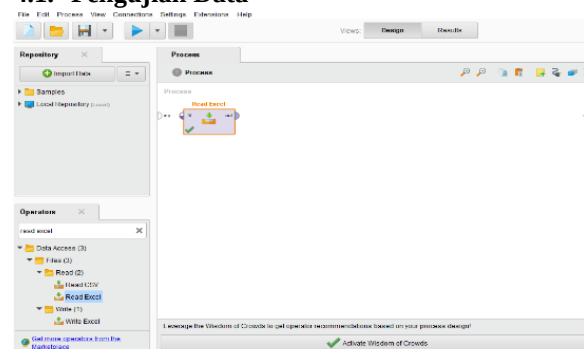
3.5. Interpretasi/Evaluasi

Pada titik ini, model telah dibangun dan diharapkan berkualitas tinggi dari sudut pandang analisis data. [10] Evaluasi hasil dari data yang telah dilakukan menggunakan *RapidMiner*. Mengidentifikasi pola dari kumpulan data yang diproses dan menghasilkan pengetahuan baru yang sesuai untuk penelitian yang diinginkan.

4. HASIL DAN PEMBAHASAN

Penelitian ini mengambil dari file *Excel* yang berisi lebih dari 624 *record* data penjualan bulan April, Mei, Juni, dan Juli. Hasil pembahasan ini memberikan penjelasan bagaimana catatan penjualan dapat dikelompokkan atau dikategorikan dengan menggunakan algoritma *K-means*. Pengelompokan ini dilakukan melalui pembelajaran mesin.

4.1. Pengujian Data



Gambar 2. Read Excel

Setelah proses impor data selesai, operator baru, yaitu *Read Excel*, akan muncul dalam proses utama. Operator ini telah menyertakan file dengan format *.xlsx* yang diimpor langsung dari file *Excel*, sehingga data siap untuk diuji. Selanjutnya, data diimpor dan dipilih dari dataset dalam format *Excel* pada *RapidMiner*.

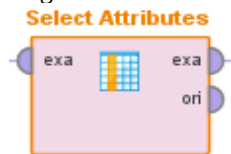
4.2. Data Selection

Seleksi data adalah proses memilih dan mengumpulkan subset atau bagian data tertentu yang relevan dengan tujuan tertentu. Tujuan utama

pemilihan data adalah untuk fokus pada informasi yang paling relevan dan penting, mengurangi kompleksitas dan mempercepat proses analisis. Proses operator *RapidMiner* terdiri dari 624 *item*, 1 atribut spesial, dan 6 atribut reguler, mencakup kategori bulan April, Mei, Juni, Juli, stok awal, stok terjual, dan stok akhir.

4.3. Data Preprocessing

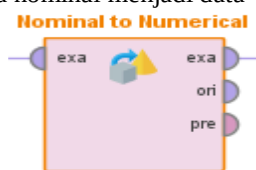
Data preprocessing (pra-pemrosesan data) adalah serangkaian langkah atau teknik yang dilakukan pada data mentah sebelum data tersebut digunakan dalam analisis atau pemodelan. Penelitian ini menggunakan data penjualan pada bulan april hingga juli, ini terdiri dari 624 record dengan 7 atribut.



Gambar 3. Select Attributes

4.4. Data Transformation

Data Transformation adalah proses mengubah atau mengelola data dari format aslinya menjadi format yang lebih sesuai atau berguna untuk tujuan analisis atau pemodelan. Transformasi data dilakukan untuk mengubah data nominal menjadi data numerik dengan menggunakan operator Nominal ke numerik. Transformasi data dilakukan untuk mengubah data nominal menjadi data numerik dengan menggunakan operator nominal ke numerik. dari nominal ke numerik dengan mengatur jenis filter atribut sebagai *subset*, memilih atribut (bulan) yang akan diubah menjadi numerik, dan mengatur atribut. Jenis pengkodean dimasukkan sebagai bilangan bulat unik dan menggunakan operator nominal-ke-numerik untuk mengubah data nominal menjadi data numerik.



Gambar 4. Nominal to Numerical

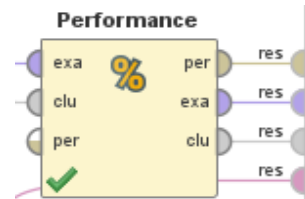
4.5. Data Mining

Data mining adalah ekstraksi pola dan wawasan yang berguna dari sejumlah besar data. Tujuan utama data mining adalah untuk menemukan pola, hubungan, atau informasi tersembunyi yang dapat membantu pengambilan keputusan.



Gambar 5. K-Means Clustering

Sesuaikan parameter *clustering* yang diperlukan dengan mengatur nilai *k* menjadi 2 dan lihat penjualan mana yang mendominasi di setiap kelompok. *Max* berjalan 10 kali, memungkinkan *loop* untuk mengamati objek lebih detail, sekaligus memungkinkan sistem memeriksa objek lain dengan properti serupa 10 kali agar hasil proses lebih akurat. Jenis pengukuran digunakan adalah perhitungan jarak jenis pengukuran numerik dengan menggunakan *Euclidean distance*. Hal ini dikarenakan ketika mempertimbangkan jarak terpendek suatu objek ke objek lainnya, lebih mudah menerapkan perhitungan dengan jarak *Euclidean*.



Gambar 6. Cluster Distance Performance

Cluster Distance Performance menggunakan parameters *Davies Bouldin Index (DBI)* untuk menghasilkan nilai *Davies Bouldin Index (DBI)*.

Cluster Model

```
Cluster 0: 616 items
Cluster 1: 8 items
Total number of items: 624
```

Gambar 7. Hasil Cluster Model

Dua *cluster* dibentuk dengan menggabungkan data penjualan yang telah diolah dan algoritma *K-Means* ke dalam model *cluster* terbagi 2 *cluster*. Pada *cluster* 0 berisi 616 *item* dan *cluster* 1 berisi 8 *item*, sehingga total 624 *item*.

Tabel 1. Hasil Cluster Optimal

Clustering.k	Davies Bouldin
2	-0.310
3	-0.434
4	-0.492
5	-0.348

Berdasarkan nilai *Davies Bouldin Index (DBI)*, hasil *clustering* terbaik terdapat pada *cluster* 2 dengan nilai *DBI* sebesar -0,310. Nilai *DBI* terendah menunjukkan maka pembagian ke dalam *cluster* tersebut lebih baik atau lebih optimal, dengan batasan antar-*cluster* yang lebih jelas dan homogenitas yang lebih tinggi dalam setiap *cluster*.

Dari hasil *clustering*, dari sini dapat menyimpulkan *Cluster* 0 lebih banyak produk, cenderung lebih diminati oleh konsumen. Oleh karena itu, disarankan untuk mengoptimalkan persediaan stok pada produk-produk dalam *Cluster_0* guna memenuhi permintaan yang tinggi. Langkah-langkah seperti peningkatan produksi, perencanaan persediaan yang efisien, dan manajemen yang baik dapat membantu memastikan ketersediaan produk yang memadai.

Sementara itu, *Cluster_1* mungkin termasuk dalam kategori produk yang kurang laris. Untuk meningkatkan kinerja penjualan produk dalam *Cluster_1* diperlukan strategi promosi, hal ini dapat meningkatkan daya tarik produk bagi konsumen.

Dapat diketahui hasil akhirnya, yaitu terdapat 616 produk yang diidentifikasi sebagai laris dan 8 produk yang diidentifikasi sebagai kurang laris berdasarkan proses analisis *clustering*. Dari hasil tersebut dapat disimpulkan bahwa produk yang laris perlu dilakukan *restock* lebih banyak, sedangkan produk yang kurang laku perlu dilakukan perbaikan atau strategi promosi untuk meningkatkan penjualan. Untuk produk yang diidentifikasi sebagai laris terdapat pada *Cluster_0* perlu dilakukan strategi pengelolaan stok yang lebih efisien dan peningkatan kapasitas persediaan agar dapat memenuhi permintaan konsumen. Pemantauan terus-menerus terhadap tren penjualan dan kebutuhan konsumen akan membantu dalam menyesuaikan strategi *restock* dan menjaga ketersediaan produk. Sementara itu, untuk produk yang diidentifikasi sebagai kurang laris terdapat pada *Cluster_1*, perlu dilakukan analisis mendalam terhadap faktor-faktor yang memengaruhi rendahnya minat konsumen. Hal ini dapat melibatkan strategi penjualan, atau penyesuaian harga. Selain itu, penerapan strategi promosi yang tepat.

Untuk *cluster k=2* peneliti mendapatkan nilai *Davies-Bouldin* : -0.310. Jika nilai *Davies Bouldin Index (DBI)* tinggi, itu menandakan bahwa proses evaluasi *cluster* kurang memadai atau tidak sesuai dengan harapan, sedangkan Dapat diketahui hasil akhirnya, yaitu terdapat 616 produk yang diidentifikasi sebagai laris dan 8 produk yang diidentifikasi sebagai kurang laris berdasarkan proses analisis *clustering*. Dari hasil tersebut dapat disimpulkan bahwa produk yang laris perlu dilakukan *restock* lebih banyak, sedangkan produk yang kurang laku perlu dilakukan perbaikan atau strategi promosi untuk meningkatkan penjualan Nilai yang diperoleh secara proporsional lebih baik dengan nilai *Davies-Bouldin Index (DBI)* yang lebih rendah.

5. KESIMPULAN DAN SARAN

Dengan menggabungkan data ini, dapat mengidentifikasi produk yang laris dan tidak laris, menghindari penumpukan produk yang tidak efisien. Hasil penelitian dari hasil *clustering* terbaik terdapat pada *cluster 2* dengan nilai *DBI* sebesar -0,310. Terdapat 616 *item* produk yang termasuk dalam kategori paling laris dalam *Cluster_0*, sementara 8 *item* produk diidentifikasi sebagai kurang laris dalam *Cluster_1*. Pengolahan data diharapkan memberikan solusi untuk memahami tren penjualan produk. Proses penerapan *K-means* melibatkan pengelompokan data persediaan penjualan produk menggunakan *RapidMiner*. *Cluster_0*, yang mencakup barang yang diminati, memerlukan perencanaan persediaan yang efisien untuk memenuhi kebutuhan konsumen dan mencegah kekurangan stok. Selain itu, *Cluster_1*,

yang mungkin mencakup barang yang kurang diminati, memerlukan strategi promosi yang efektif untuk meningkatkan minat konsumen dan mencegah penumpukan barang yang dapat mengurangi pendapatan. Mengusulkan peneliti lain untuk memanfaatkan data penjualan toko lain dalam penelitian dan mencari metode alternatif untuk meningkatkan kualitas penelitian. Jika ingin mengelompokkan lebih lanjut data penjualan, mungkin ingin mempertimbangkan metode selain pengelompokan *K-means* sebagai pengujian alternatif. Selain itu, menyarankan penggunaan kumpulan data yang lebih besar untuk mencapai hasil yang lebih tepat, yang memungkinkan penelitian di masa depan mempertimbangkan metode alternatif untuk mengatur atau mengkategorikan data.

DAFTAR PUSTAKA

- [1] Haditsah Annur, "Penerapan Data Mining Menentukan Strategi Penjualan Variasi Mobil Menggunakan Metode K-Means Clustering (Studi Kasus Toko Luxor Variasi Gorontalo)," *J. Inform. Upgris*, vol. 5, no. 1, 2019.
- [2] . 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.
- [3] F. A. Bramasta and R. Halilintar, "Penerapan Data Mining Untuk Menentukan Strategi Penjualan Toko Sepatu," *Pros. SEMNAS INOTEK ...*, pp. 236–241, 2021, [Online]. Available: <https://proceeding.unpkediri.ac.id/index.php/inotek/article/view/1135%0Ahttps://proceeding.unpkediri.ac.id/index.php/inotek/article/download/1135/736>
- [4] 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.
- [5] 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.
- [6] G. Triyandana, L. A. Putri, and ..., "Penerapan Data Mining Pengelompokan Menu Makanan dan Minuman Berdasarkan Tingkat Penjualan Menggunakan Metode K-Means," *J. Appl. ...*, 2022, [Online]. Available: <https://jurnal.polibatam.ac.id/index.php/JAIC/article/view/3824>
- [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. Inf. Secur.*, vol. 3, no. 1, pp. 99–110, 2022, doi: 10.31599/jiforty.v3i1.1253.
- [9] D. Anggarwati, O. Nurdiawan, I. Ali, and D. A. Kurnia, “Penerapan Algoritma K-Means Dalam Prediksi Penjualan Karoseri,” *J. Data Sci. Inform.*, vol. 1, no. 2, pp. 58–62, 2021.
- [10] 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.