

Pengelompokan Persediaan *Fast Moving* Untuk Barang Impor Menggunakan K-Means Di PT Indo Bumi Lestari

Dandi Jaenudin Pratama Nugraha¹, Rangga Sanjaya²

¹Program Studi Teknik Informatika, Universitas Adhirajasa Reswara Sanjaya

²Program Studi Sistem Informasi, Universitas Adhirajasa Reswara Sanjaya

e-mail: ¹dandiwidiya58@gmail.com, ¹rangga@ars.ac.id

Abstrak

Produk *fast moving* adalah kategori dalam perdagangan dengan tingkat perputaran tinggi dalam waktu tertentu. Dalam pengelolaan *inventori*, penting menjaga keseimbangan untuk mencegah kerugian akibat *overstock* atau *stock out*. *Overstock* meningkatkan biaya penyimpanan, risiko kerusakan, dan barang yang tidak laku di pasaran, sementara *stock out* menyebabkan barang tidak tersedia saat permintaan tinggi, sehingga penjualan hilang. Untuk mengatasi masalah ini, perusahaan perlu melakukan *clustering* atau pengelompokan barang agar stok sesuai dengan kebutuhan pasar. PT Indo Bumi Lestari, distributor produk Korea yang bekerja sama dengan YOGYA Group di YOGYA Riau Junction, menghadapi masalah pengelolaan stok. Tanpa metode *clustering* yang efisien, produk sering terbuang karena kedaluwarsa, sementara barang dengan permintaan tinggi justru kosong terlalu lama. Algoritma K-Means *Clustering* cocok diterapkan karena sederhana dan efektif untuk mengelompokkan data penjualan berdasarkan pola tertentu. Dengan pengelompokan menggunakan algoritma ini, perusahaan dapat menentukan produk mana yang harus selalu tersedia dan mana yang cukup sebagai pelengkap. Solusi ini membantu meningkatkan keuntungan sekaligus mengurangi kerugian dengan pengelolaan *inventori* yang lebih terstruktur dan efisien.

Kata kunci— Fast Moving, Algoritma K-Means, Clustering

Abstract

Fast-moving products are a category in trade with a high turnover rate within a certain time. In inventory management, it is important to maintain a balance to prevent losses due to overstock or stock out. Overstock increases storage costs, risk of damage, and unsold goods in the market, while stock out causes goods to be unavailable when demand is high, resulting in lost sales. To overcome this problem, companies need to cluster or group goods so that stock is in accordance with market needs. PT Indo Bumi Lestari, a distributor of Korean products that collaborates with YOGYA Group at YOGYA Riau Junction, faces stock management problems. Without an efficient clustering method, products are often wasted due to expiration, while goods with high demand are empty for too long. The K-Means Clustering algorithm is suitable for application because it is simple and effective for grouping sales data based on certain patterns. By grouping using this algorithm, companies can determine which products should always be available and which are sufficient as complements. This solution helps increase profits while reducing losses with more structured and efficient inventory management.

Keywords— Fast Moving, Algoritma K-Means, Clustering

Corresponding Author:

Rangga Sanjaya,

Email: rangga@ars.ac.id

1. PENDAHULUAN

Di era digitalisasi yang berkembang pesat, persaingan di sektor perdagangan barang impor semakin tinggi. Dengan tingginya persaingan ini para pelaku ekonomi harus memperhatikan salah satu hal penting dalam manajemen yaitu pengendalian persediaan [1]. Pengendalian persediaan dilakukan oleh perusahaan dalam rangka memenuhi kebutuhan konsumen dan menyediakan barang yang dibutuhkan. Pengelompokan stok persediaan *fast moving* digudang merupakan suatu upaya pengendalian yang tujuannya adalah menyediakan segala kebutuhan barang yang dibutuhkan [2]. Perusahaan harus mampu mengelola persediaan barang dengan efisien untuk memastikan ketersediaan produk sesuai dengan permintaan pasar. Fokus utama dalam penelitian ini adalah bagaimana mengelompokkan data persediaan barang *fast moving* secara efektif berdasarkan penjualan. Pengelompokan yang tepat dapat membantu perusahaan dalam mengoptimalkan strategi penjualan dan pengelolaan persediaan barang, juga dapat memberikan pengalaman belanja yang lebih baik kepada pelanggan [3]

Pengelompokan persediaan *fast moving* atau barang yang memiliki tingkat perputaran tinggi sering kali menghadapi tantangan dalam hal ketersediaan stok baik di gudang maupun di area pajang. Dengan pengelompokan data barang yang baik berpengaruh bagi perkembangan dan kemajuan suatu perusahaan atau instansi terutama yang bergerak dalam bidang perdagangan [2]. Jika stok tidak dikelola dengan baik, perusahaan bisa mengalami kehabisan stok yang mengakibatkan penjualan hilang, atau sebaliknya kelebihan stok yang menyebabkan biaya penyimpanan dan kerusakan meningkat [4]. Dengan demikian, pengelolaan yang efisien adalah kunci untuk meningkatkan profitabilitas dan keberlanjutan bisnis.

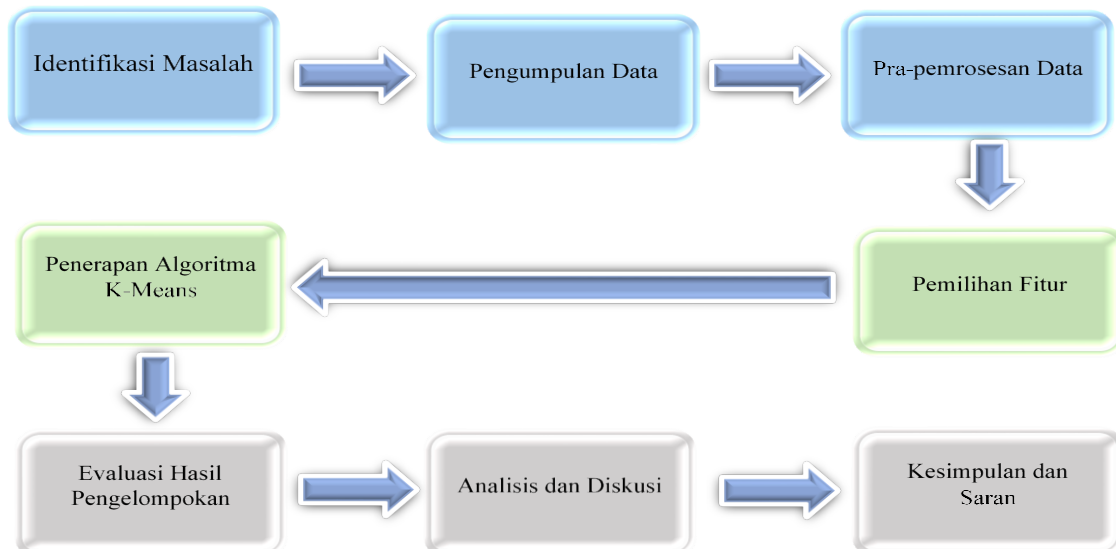
Salah satu metode yang dapat digunakan untuk pengelompokan persediaan adalah algoritma K-Means. Algoritma K-Means adalah salah satu algoritma dengan teknik *clustering* berdasarkan pembagian jarak dalam *data mining* [2], [5]. Keuntungan menggunakan algoritma ini mudah dipahami, diterapkan, serta memiliki efek pengelompokan yang baik [5]. Dengan menggunakan algoritma K-Means, perusahaan dapat mengidentifikasi pola penjualan dan mengelompokkan barang-barang yang memiliki karakteristik penjualan yang serupa dalam satu *cluster* dan data dengan karakteristik yang berbeda akan dikelompokkan kedalam *cluster* lain [6]. Penerapan algoritma Algoritma K-Means *Clustering* dapat digunakan sebagai dasar penentuan strategi pemasaran yang tepat untuk menentukan jumlah objek berdasarkan kriteria kriteria yang telah ditentukan [7]. K-Means ini menjadi krusial pada saat perusahaan berhadapan dengan data penjualan yang besar dan kompleks, yang sulit dianalisis secara manual. Algoritma K-Means dapat diaplikasikan di berbagai bidang, seperti pengelompokan data obat-obatan dan penentuan siswa kelas unggulan penelitian [8], [9]

Penelitian ini dilakukan di PT. Indo Bumi Lestari (MUGUNGHWA), sebuah perusahaan yang bergerak di bidang perdagangan barang impor dari Negara Korea. Studi kasus dilakukan di Cabang Yogya Riau Junction, dimana MUGUNGHWA menghadapi tantangan dalam mengelola persediaan *fast moving*, yang memiliki tingkat penjualan tinggi dan cepat habis. Oleh karena itu, diperlukan suatu metode yang dapat membantu dalam pengelompokan data persediaan berdasarkan penjualan, sehingga pengelolaan persediaan dapat dilakukan dengan lebih efektif dan efisien. Data penjualan MU GUNG HWA di cabang Yogya Riau Junction akan dianalisis menggunakan algoritma K-Means untuk mengelompokkan barang-barang *fast moving* berdasarkan pola penjualannya. Melalui analisis ini, diharapkan di waktu yang akan datang perusahaan dapat mengidentifikasi kelompok barang yang memiliki karakteristik penjualan serupa, sehingga dapat merencanakan pengadaan dan pengelolaan persediaan dengan lebih tepat.

Sebagai upaya untuk mengatasi permasalahan yang telah di uraikan di atas, peneliti tertarik untuk mengembangkan model pengelompokan persediaan *fast moving* untuk barang impor menggunakan K-Means di PT Indo Bumi Lestari. Hasil dari penelitian ini akan diaplikasikan dalam strategi pengelolaan persediaan yang diharapkan dapat memberikan kontribusi signifikan dalam meningkatkan efisiensi operasional dan kinerja penjualan perusahaan MUGUNGHWA di cabang Yogya Riau Junction.

2. METODE PENELITIAN

Untuk penelitian yang berkaitan dengan pengelompokan barang *fast moving* menggunakan algoritma K-Means, metodologi penelitian perlu disusun secara sistematis agar tujuan penelitian tercapai dengan baik. Berikut adalah metodologi yang bisa diikuti



Gambar 1. Tahapan Penelitian

2.1. Identifikasi Masalah

Sebagaimana dinyatakan dalam pendahuluan penelitian ini, algoritma K-Means *Clustering* akan di gunakan untuk mengelompokkan barang *fast moving* berdasarkan pola penjualan barang impor PT Indo Bumi Lestari di Cabang Yogya Riau Junction. Tujuan utama penelitian ini adalah untuk mengelompokkan barang yang memiliki tingkat perputaran tinggi sehingga perusahaan mampu melakukan pengelolaan persediaan barang dengan optimal untuk menghindari kerugian yang di sebabkan *overstock* dan *stock out* di kemudian hari.

2.2. Pengumpulan Data

Pengumpulan data ini tentunya sangat penting dan menjadi komponen utama untuk melakukan penelitian dalam penggunaan algoritma [10]. Data yang di gunakan merupakan data penjualan perusahaan PT Indo Bumi Lestari di cabang Yogya Riau Junction, data tersebut adalah data penjualan biskuit dan minuman selama 1 bulan penuh pada bulan Mei 2024. Data ini berisi 71 baris data.

2.3. Pra-pemrosesan Data

Pra-pemrosesan data adalah tahap pembersihan data dari hal-hal yang tidak di perlukan selama penggunaan algoritma menjadi data yang dapat diolah menggunakan algoritma yang akan dipakai [11], Ini termasuk normalisasi, menangani data yang hilang, dan penghilangan *outlier*. Tahapan dalam melakukan pra-pemrosesan data antara lain

1. *Data Cleaning*, yaitu proses menghapus data yang tidak lengkap atau salah.
2. *Normalisasi Data*, yaitu proses menyesuaikan skala data agar tidak ada variable yang mendominasi hasil pengelompokan.

2.4. Pemilihan Fitur

Pemilihan fitur adalah proses memilih atribut data yang paling relevan untuk pengelompokan. Misalnya, dalam pengelompokan barang *fast moving*, fitur yang relevan mungkin mencakup jumlah penjualan, harga barang, dan tingkat persediaan. Langkah dalam menentukan fitur antaralain

1. Menentukan variabel penting seperti total penjualan, waktu penyimpanan, atau kecepatan perputaran barang.
2. Mengeliminasi fitur yang tidak relevan atau memiliki korelasi tinggi dengan fitur lain.

2.5. Penerapan Algoritma K-Means

Pada tahap ini, algoritma K-Means diterapkan untuk mengelompokkan barang-barang berdasarkan kemiripan karakteristik. Algoritma ini bekerja dengan meminimalkan jarak antara titik data dan *centroid* (titik pusat) dari setiap kelompok. Langkah-langkah dalam melakukan clustering dengan metode K-Means [12], adalah sebagai berikut

1. *Isialisasi*, Tentukan jumlah *cluster* (k) dan pilih secara acak (k) titik awal sebagai *centroid* awal. Pemilihan titik awal ini dapat memengaruhi hasil akhir *clustering*.
2. *Penugasan Cluster*, Setiap titik data x_i dikelompokkan ke *cluster* yang memiliki *centroid* terdekat berdasarkan jarak *Euclidean*. sebagai ukuran kedekatan. Jarak *Euclidean* antara dua titik data $x=(x_1, x_2, \dots, x_n)$ dan $y=(y_1, y_2, \dots, y_n)$ dihitung dengan formula

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

Untuk setiap titik data x_i *cluster* C_i dipilih berdasarkan jarak terdekat ke *centroid* μ_j

$$C_i = \underset{j}{\operatorname{argmin}} d(x_i, \mu_j)$$

Di mana $d(x_i, \mu_j)$ adalah jarak *Euclidean* antara titik data x_i dan *centroid* μ_j .

3. *Perhitungan Ulang Centroid*, Setelah setiap titik dikelompokkan, hitung ulang *centroid* dari setiap *cluster* sebagai rata-rata dari semua titik yang ada di *cluster* tersebut. Rata-rata ini dihitung menggunakan rumus

$$\mu_j = \frac{1}{|C_j|} \sum_{x_i \in C_j} x_i$$

Di mana $|C_j|$ adalah jumlah titik data dalam *cluster* C_j dan x_i adalah titik data.

4. *Iterasi*, Langkah penugasan *cluster* dan perhitungan ulang *centroid* diulang hingga posisi *centroid* tidak berubah lagi atau perubahannya kecil sekali. Pada saat ini, algoritma dianggap telah mencapai *konvergensi*.
5. *Formula Sum of Squared Errors (SSE)*, bertujuan untuk meminimalkan fungsi, yaitu total jarak kuadrat antara setiap titik data dan *centroid* dalam *cluster* yang sesuai. (*SSE*) dihitung dengan formula

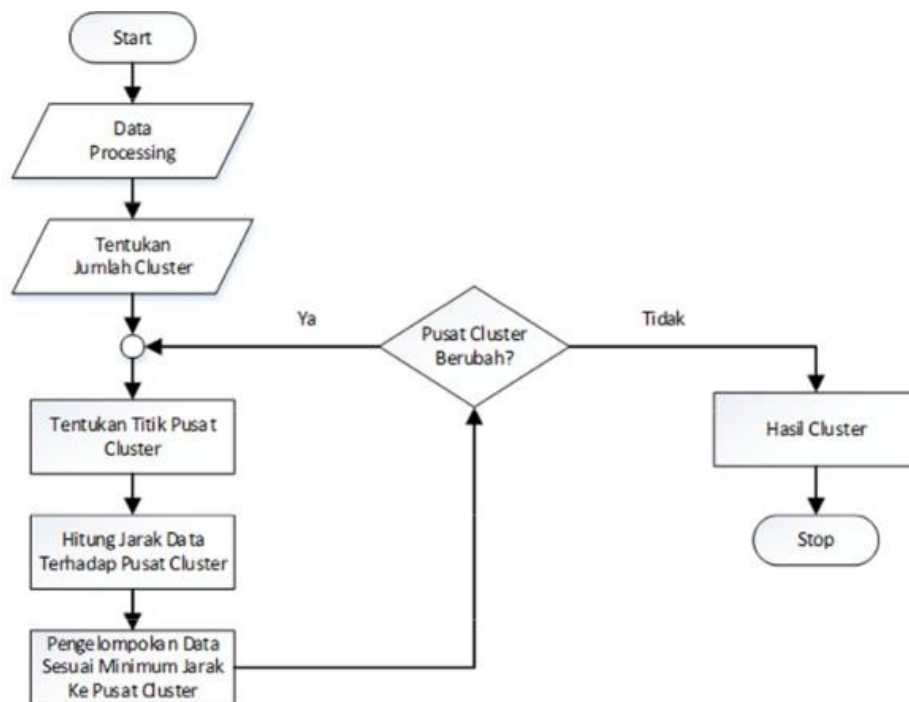
$$J = \sum_{i=1}^k \sum_{x_j \in C_i} ||x_j - \mu_i||^2$$

Dimana

- a. J adalah total kesalahan kuadrat (fungsi biaya),
- b. x_j adalah titik data dalam *cluster* C_i ,
- c. μ_i adalah *centroid cluster* C_i ,
- d. $||x_j - \mu_i||^2$ adalah jarak *Euclidean* kuadrat antara titik data dan *centroid*.

Minimisasi (*SSE*) memastikan bahwa setiap titik data berada sedekat mungkin dengan *centroid*-nya, yang berarti *cluster* lebih kompak dan hasil *clustering* lebih baik.

6. Proses *cluster* sudah selesai bila pusat *cluster* tidak berubah, namun jika pusat *cluster* masih berubah-ubah maka ulangi langkah menghitung jarak hingga pusat *cluster* tidak berubah lagi [13]. Tahapan *clustering* dapat dilihat pada gambar 2.



Gambar 2. Flowchart algoritma K-Means *Clustering*

2.6. Evaluasi Hasil Pengelompokan

Algoritma K-Means memerlukan evaluasi untuk menentukan seberapa baik pengelompokan yang dihasilkan, evaluasi algoritma yang di gunakan adalah *Elbow Method*, Digunakan untuk menentukan jumlah *cluster* (k) yang optimal. Dengan memplot nilai *Within-Cluster Sum of Squares* (WCSS) terhadap berbagai nilai (k), kita dapat melihat titik di mana penurunan WCSS mulai melambat atau dengan sederhananya kita dapat melihat titik yang paling membentuk sudut. Titik ini menunjukkan jumlah *cluster* yang ideal [14].

3. HASIL DAN PEMBAHASAN

3.1. Pengumpulan Data

Penelitian ini menggunakan data penjualan pada perusahaan PT. Indo Bumi Lestari yang bercabang di Yogya Riau Junction selama 1 bulan pada bulan Mei 2023. Data penjualan PT. Indo Bumi Lestari Mei 2023 Sebagai berikut.

Tabel 1. Penjualan PT. Indo Bumi Lestari Cabang Yogya Riau Junction Mei 2023

Penjualan PT. Indo Bumi Lestari Cabang Yogya Riau Junction Mei 2023										
No	keterangan	harga	qty	jumlah	datang barang				Se ll	sto k
					1	2	3	4		
1	Maxim Original 236g	Rp 113,000	19	Rp 2,147,000	0	20	0	0	25	14
2	Maxim Mocha Gold 240g	Rp 110,500	6	Rp 663,000	0	20	0	0	20	6
3	Maxim White Gold 234g	Rp 110,500	22	Rp 2,431,000	0	20	0	0	26	16
4	See's Coffee French Vanila Cappuccino 100g	Rp 45,500	0	Rp -	0	0	0	0	0	0
5	Sempio Coen Tea 500g	Rp 48,200	0	Rp -	0	0	0	0	0	0
6	Sempio Barley Tea 1kg	Rp 94,000	0	Rp -	0	0	0	0	0	0
...			...		...	...	...	...	...	...
69	Mamos Multi Grain	Rp 32,000	0	Rp -	0	0	0	0	0	0
70	Crown Kookhee Choco Sand 70g	Rp 23,500	0	Rp -	0	0	0	0	0	0
71	Choco Free Time Bar	Rp 13,500	30	Rp 405,000	0	0	0	0	8	22

3.2. Pra-pemrosesan Data

Peneliti melakukan pembersihan data dengan menghapus item tanpa data penjualan atau stok serta menghilangkan data yang tidak relevan. Langkah ini bertujuan untuk fokus pada data penjualan *fast moving*, yaitu barang dengan pergerakan tinggi dalam periode tertentu. Barang tanpa penjualan dikategorikan sebagai *slow moving* atau *dead moving* pada periode tersebut. Hasil pembersihan data sebagai berikut.

Tabel 2. Data Penjualan PT Indo Bumi Lestari Pada Bulan Mei 2023 Setelah Dilakukan Pra-pemrosesan.

No	Uraian	Penjualan	Stok
1	Maxim Original 236g	25	14
2	Maxim Mocha Gold 240g	20	6
3	Maxim White Gold 234g	26	16
4	Dongsuh Corn tea 300g	4	11
5	Wj Bluebery	9	15
6	Gf Aloe 500ml	31	11
...		...	...
17	Mamos Rice Cracker Banana	19	1
18	Mamos Rice Cracker Choco	20	0
19	Choco Free Time Bar	8	22

3.3. Pemilihan Fitur

Fitur yang digunakan untuk pengelompokan barang *fast moving* meliputi id, nama barang dan total penjualan (Penjualan). Fitur yang di gunakan berfokus pada total penjualan pada bulan mei 2023. Fitur-fitur ini akan diproses dalam analisis *clustering*. Hasil dapat di lihat pada tabel 3.

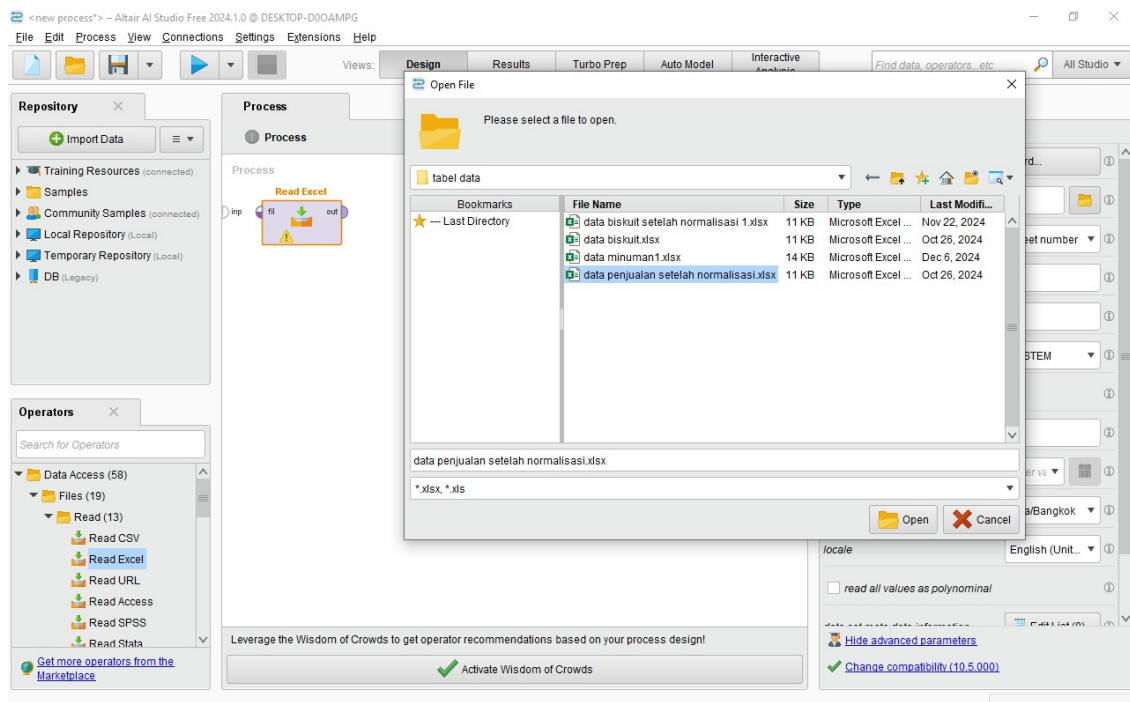
Tabel 3. Pemilihan Fitur Pada Data Penjualan Di Bulan Mei 2023

No	Uraian	Penjualan
1	Maxim Original 236g	25
2	Maxim Mocha Gold 240g	20
3	Maxim White Gold 234g	26
4	Dongsuh Corntea 300g	4
5	Wj Bluebery	9
6	Gf Aloe 500ml	31
...		...
17	Mamos Rice Cracker Banana	19
18	Mamos Rice Cracker Choco	20
19	Choco Free Time Bar	8

3.4. Penerapan Algoritma K-Means

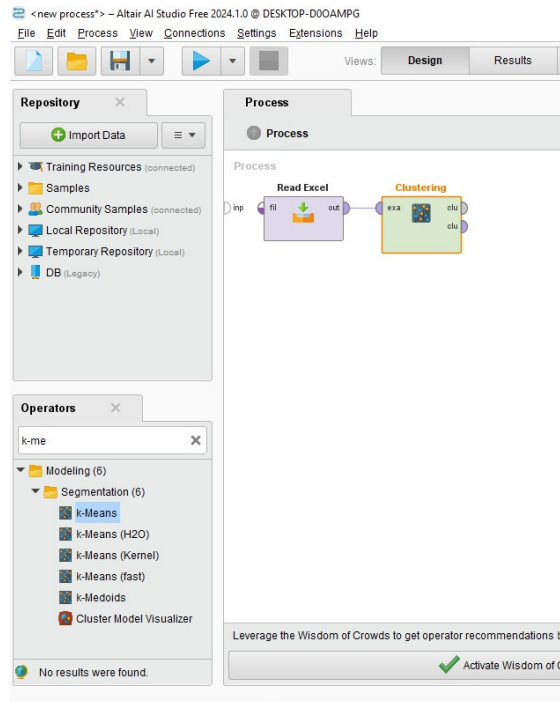
Pada tahapan ini peneliti menggunakan aplikasi Rapidminer AI Studio 2024.1.0 untuk menerapkan Algoritma K-Means *Clustering*

1. *Read Excel*, agar aplikasi Rapidminer mampu membaca data yang telah kita siapkan pada aplikasi Rapidminer kita menggunakan tools pada operator yaitu read excel, karena data yang kita miliki dalam bentuk format Ms.Excel. dengan tampilan sebagai berikut.

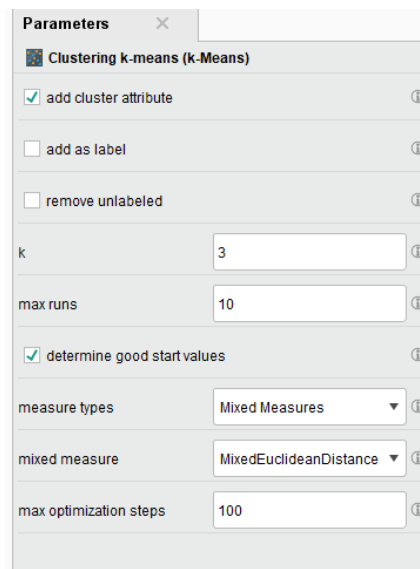


Gambar 3. Memasukan Data Kedalam Aplikasi Rapidminer Menggunakan Tools Read Excel

2. *Memasukkan Algoritma K-Means*, di tahap ini kita memasukkan Algoritma K-Means dapat di lihat pada gambar 4, dan Berdasarkan *Elbow Method*, peneliti akan menentukan jumlah *cluster* yang paling optimal untuk pengelompokan. Data akan dibagi menjadi 3 *cluster* berdasarkan penjualan dapat di lihat pada gambar 5.



Gambar 4. Memasukkan Algoritma K-Means.

Gambar 5. Pengaturan Persiapan Proses *Clustering K-means*

Penjelasan dari gambar 5 di atas:

1. Menentukan jumlah *Cluster (k)*,
Berdasarkan *Elbow Method*, peneliti akan menentukan jumlah *cluster* yang paling optimal untuk pengelompokan. Data akan dibagi menjadi 3 *cluster*, yaitu *fast-moving*, *medium-moving*, dan *slow-moving* berdasarkan penjualan.

2. Menentukan berapa kali Rapidminer melakukan *clustering* K-Means
Rapidminer akan melakukan proses *clustering* K-Means sebanyak 10x secara berulang dan hasil yang akan di tampilkan merupakan hasil terbaik, dalam gambar 5 di atas di tunjukan pada tool (max run).
3. Menentukan *optimization step*
Optimization step atau proses pemindahan titik pusat atau *centroid* dengan metode K-Means. Setiap K-Means beroperasi titik pusat atau *centroid* akan berpindah hingga titik pusat ini berhenti berpindah maka di situlah titik pusat pengelompokan di tentukan atau pusat *centroid* di tentukan. Peneliti pengatur agar penentuan titik pusat ini berulang sebanyak 100x dan di tampilkan hasil yang terbaik.
3. Hasil K-Means Clustering, Setelah pusat data atau *centroid* terus berpindah dan akhirnya berhenti berpindah berdasarkan kedekatan antar data maka di perolehlah hasil K-Means *clustering* pada data penjualan sebagai berikut

Cluster Model

```

Cluster 0: 7 items
Cluster 1: 9 items
Cluster 2: 3 items
Total number of items: 19

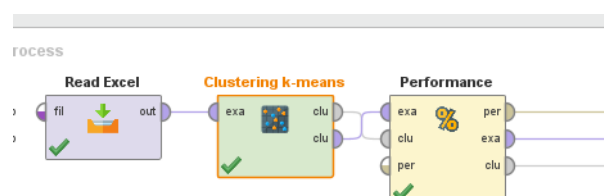
```

Gambar 6. Hasil Dari Pengelompokan Data Menggunakan Algoritma K-Means Clustering

Berdasarkan hasil dari pengelompokan data menggunakan Algoritma K-Means, data sebanyak 19 dibagi menjadi 3 *cluster* berdasarkan penjualan pada bulan mei 2023. *Cluster 0* terdiri dari 7 data. *Cluster 1* terdiri dari 9 data. *Cluster 2* terdiri dari 3 data. *Cluster 0* merupakan kelompok data yang memiliki tingkat perputaran menengah atau *medium moving* sehingga barang jual pada kelompok data ini masih di pertimbangkan dalam melakukan penyediaan stok *inventori*. *Cluster 1* merupakan kelompok data yang memiliki tingkat perputaran rendah atau *slow moving* sehingga barang jual pada kelompok data ini tidak di anjurkan dalam penyediaan stok *inventori*. Sedangkan *Cluster 2* merupakan kelompok data dengan tingkat perputaran tinggi atau *Fast moving* pada periode tersebut, barang jual yang ada pada kelompok data ini harus di pertimbangkan untuk melukan stok lebih pada *inventori*.

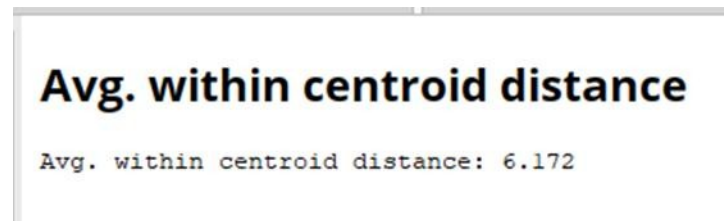
3.5. Evaluasi

Setelah pengelompokan, hasilnya akan dievaluasi menggunakan *Elbow Methode* untuk menentukan kualitas pengelompokan. Namun langkah ini memerlukan pengulangan percobaan yang dilakukan terhadap nilai (k) sebanyak beberapa kali hingga diketahui jarak terbaik. Cara menghitung jarak tersebut yaitu menggunakan bantuan tools yang tersedia di aplikasi Rapidminer *Cluster Performance Distance*. Dapat di lihat pada gambar 7



Gambar 7. Menentukan jumlah *Average Distance*

Hasil jarak antar data yang di hitung setelah pada data yang di bagi kedalam 3 *cluster* adalah 6.172. Hasil dapat di lihat pada gambar 8.



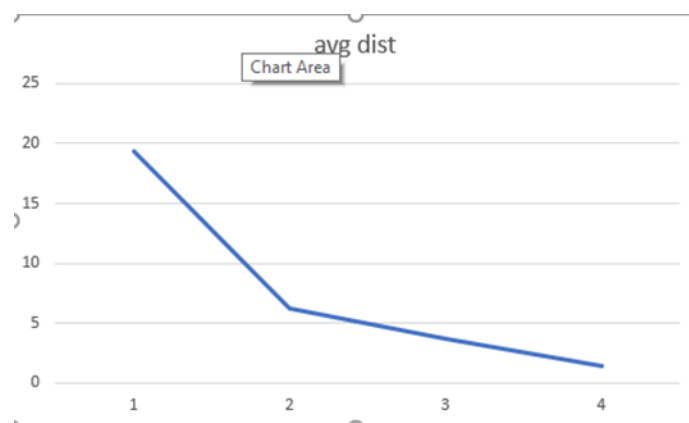
Gambar 8. hasil pengitungan jarak data 3 cluster

Langkah selanjutnya adalah melakukannya berkali kali, peneliti akan melakukan proses ini berulang dengan mengganti jumlah *cluster* dari yang terkecil yaitu dari $k = 2$ sampai $k = 5$, maka hasilnya sebagai berikut

Tabel 4. Hasil Jarak Data Dengan Nilai (k) Yang Berbeda

k	avg dist
2	19.315
3	6.172
4	3.729
5	1.437

Dari gambar di atas di ketahui bahwa jarak antar data satu dan data lainnya jika di bagi kedalam 2 *cluster* adalah 19.315, lalu jika data tersebut di bagi menjadi 3 *cluster* adalah 6.172 begitupun seterusnya hingga data tersebut di bagi menjadi 5 *cluster*, Lalu tahap selanjutnya adalah menganalisis menggunakan diagram x,y . Dapat di lihat pada gambar 9



Gambar 9. Jarak Data Menggunakan Diagram x,y

Dapat di jelaskan dari data di atas bahwa angka 1 pada sumbu x atau pada garis *Horizontal* menandakan 2 *cluster*, angka 2 pada sumbu x atau pada garis *Horizontal* menandakan 3 *cluster* dan seterusnya. Sesuai dengan Namanya, *Elbow Methode* merupakan cara menganalisis data *cluster* dengan cara melihat melalui diagram jarak data mana yang paling dominan membentuk sudut, dilihat dari data di atas bahwa data yang paling membentuk sudut adalah pada angka nomor 2 di garis *horizontal* yaitu pada saat data dibagi menjadi 3 *cluster*.

4. KESIMPULAN

Penelitian ini dilakukan untuk memaksimalkan pengelolaan persediaan barang impor dengan melakukan pengelompokan terhadap barang-barang yang perlu memiliki penanganan khusus dalam pengelolaan persediaan barang dengan cara mengelompokkan barang berdasarkan tingkat penjualan pada periode tertentu menggunakan Algoritma K-means Clustering. Hasilnya bahwa barang yang di jual pada perusahaan PT.Indo Bumi Lestari didominasi oleh barang yang memiliki tingkat penjualan rendah, data menunjukkan terdapat 9 data dari 19 data yang memiliki tingkat penjualan rendah, 7 data dari 19 data memiliki tingkat penjualan menengah sedangkan barang *fast moving* hanya terdapat 3 data saja dari 19 data yang ada. Ini menunjukkan bahwa perusahaan belum memiliki metode klasifikasi yg optimal dalam melakukan pengelolaan persediaan yang banyak menyebabkan kerugian. Algoritma K-Means Clustering bekerja dengan optimal ketika data penjualan dibagi menjadi 3 cluster, hal ini dibuktikan saat evaluasi menggunakan metode elbow method yang menunjukkan ketika data dikelompokkan menjadi 3 cluster jarak antar data sebesar 6.172. Dengan hasil pengelompokan data penjualan menggunakan Algoritma K-Means ini perusahaan dapat merencanakan dan melakukan tindakan yang sesuai terhadap pengelolaan persediaan barang di kemudian hari dengan lebih efektif dan efisien.

DAFTAR PUSTAKA

- [1] P. W. Nofiani, C. Mursid, and M. C. Mursid, "PENTINGNYA PERILAKU ORGANISASI DAN STRATEGI PEMASARAN DALAM MENGHADAPI PERSAINGAN BISNIS DI ERA DIGITAL," *Jurnal Logistik Bisnis*, vol. 11, no. 02, 2021, [Online]. Available: <https://ejurnal.poltekpos.ac.id/index.php/logistik/index>
- [2] 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 : Jurnal Ilmiah Manajemen Informatika dan Komputer*, vol. 4, no. 2, pp. 58–64, Jun. 2022, doi: 10.32485/kopertip.v4i2.120.
- [3] D. Putri Adilah Asih, B. Irawan, and A. Bahtiar, "PENGELOMPOKAN DATA TRANSAKSI DALAM MENENTUKAN STRATEGI PENJUALAN MENGGUNAKAN ALGORITMA K-MEANS," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 1, pp. 141–147, Feb. 2024, doi: 10.36040/jati.v8i1.8321.
- [4] N. Syahfitri, E. Budianita, A. Nazir, and I. Afrianty, "KLIK: Kajian Ilmiah Informatika dan Komputer Pengelompokan Produk Berdasarkan Data Persediaan Barang Menggunakan Metode Elbow dan K-Medoid," *Media Online*, vol. 4, no. 3, pp. 1668–1675, 2023, doi: 10.30865/klik.v4i3.1525.
- [5] Sekar Setyaningtyas, B. Indarmawan Nugroho, and Z. Arif, "TINJAUAN PUSTAKA SISTEMATIS: PENERAPAN DATA MINING TEKNIK CLUSTERING ALGORITMA K-MEANS," *Jurnal Teknoif Teknik Informatika Institut Teknologi Padang*, vol. 10, no. 2, pp. 52–61, Oct. 2022, doi: 10.21063/jtif.2022.v10.2.52-61.
- [6] N. F. Adani, A. F. Boy, S. Kom, M. Kom, and R. Syahputra, "Implementasi Data Mining Untuk Pengelompokan Data Penjualan Berdasarkan Pola Pembelian Menggunakan Algoritma K-Means Clustering Pada Toko Syihan STMIK Triguna Dharma ** Program Studi Sistem Informasi, STMIK Triguna Dharma *** Program Studi Sistem Informasi, STMIK Triguna Dharma," *Jurnal CyberTech*, vol. x. No.x, [Online]. Available: <https://ojs.trigunadharma.ac.id>
- [7] A. A. Rismayadi, N. N. Fatonah, and E. Junianto, "ALGORITMA K-MEANS CLUSTERING UNTUK MENENTUKAN STRATEGI PEMASARAN DI CV. INTEGRIT KONSTRUKSI," *JURNAL RESPONSIF*, vol. 3, no. 1, pp. 30–36, 2021, [Online]. Available: <http://ejurnal.ars.ac.id/index.php/jti>

- [8] G. Gustientiedina, M. H. Adiya, and Y. Desnelita, "Penerapan Algoritma K-Means Untuk Clustering Data Obat-Obatan," *Jurnal Nasional Teknologi dan Sistem Informasi*, vol. 5, no. 1, pp. 17–24, Apr. 2019, doi: 10.25077/TEKNOSI.v5i1.2019.17-24.
- [9] A. Sulistiyawati and E. Supriyanto, "Implementasi Algoritma K-means Clustering dalam Penentuan Siswa Kelas Unggulan," *Jurnal Tekno Kompak*, vol. 15, no. 2, p. 25, Aug. 2021, doi: 10.33365/jtk.v15i2.1162.
- [10] B. Rizki, N. Hidayat, and R. Sanjaya, "Penerapan Text Mining Dengan Algoritma Random Forest Menganalisis Sentimen Ulasan SATUSEHAT Mobile," vol. 5, no. 2, 2024.
- [11] F. Nasari, C. Jhony, and M. Sianturi, "Penerapan Algoritma K-Means Clustering... 108 Penerapan Algoritma K-Means Clustering Untuk Pengelompokan Penyebaran Diare Di Kabupaten Langkat."
- [12] R. Supardi and I. Kanedi, "IMPLEMENTASI METODE ALGORITMA K-MEANS CLUSTERING PADA TOKO EIDELWEIS," *Jurnal Teknologi Informasi*, vol. 4, no. 2, 2020.
- [13] J. Hutagalung, Y. Hendro Syahputra, Z. Pertiwi Tanjung, S. Triguna Dharma, and J. I. Pintu Air, "Pemetaan Siswa Kelas Unggulan Menggunakan Algoritma K-Means Clustering," *Hal AH Nasution*, vol. 9, no. 1, 2022, [Online]. Available: <http://jurnal.mdp.ac.id>
- [14] D. Amelia, T. N. Padilah, and A. Jamaludin, "Optimasi Algoritma K-Means Menggunakan Metode Elbow dalam Pengelompokan Penyakit Demam Berdarah Dengue (DBD) di Jawa Barat," *Jurnal Ilmiah Wahana Pendidikan*, vol. 8, no. 11, pp. 207–215, 2022, doi: 10.5281/zenodo.6831380.