

Implementasi Sistem Klasifikasi Inventaris Menggunakan Metode Clustering dengan Algoritma K-Means untuk Pengelolaan Stok Barang di D'Cafe Indramayu

Dina Rima Dhona¹, Yulianingsih², Suranto Saputra³

^{1,2,3}Teknik Informatika, Universitas Indraprasta PGRI, Indonesia

Article Info

Article history:

Received Aug 28, 2024

Revised Dec 20, 2024

Accepted Oct 05, 2025

Keywords:

Inventory management, Java application, K-Means, RFM

ABSTRACT

The primary issue in inventory management lies in the inefficiency of stock classification, which can lead to shortages or surpluses in inventory and consequently affect the effectiveness of marketing strategies. This study aims to develop an inventory classification system using a clustering approach with the k-means algorithm based on the Recency, Frequency, and Monetary (RFM) model to improve the accuracy of product grouping. The research methodology involves collecting data on product process and stock quantities, calculating RFM parameters, performing data normalization, and conducting the clustering process using K-Means algorithm. The finding of the study reveal the information of three main clusters: product with low to medium price ranges, product with medium to high, and product with high prices.

Corresponding Author:

Yulianingsih

Teknik Informatika,

Universitas Indraprasta PGRI,

Jl. Nangka No. 58 C, Tanjung Barat, Jagakarsa, Jakarta Selatan.

Email: yuliaunindra@gmail.com

1. PENDAHULUAN

Manajemen inventaris merupakan aspek krusial dalam pengelolaan operasional bisnis. Ketidakefektifan dalam pengelolaan stok dapat menyebabkan kekurangan atau kelebihan persediaan dan mempengaruhi pengambilan keputusan manajerial. Peningkatan volume inventaris yang tidak berimbang dapat menimbulkan kompleksitas dalam aspek pengelolaan dan dokumentasi stok, yang pada gilirannya dapat mengganggu kelancaran proses penjualan dan pemenuhan kebutuhan konsumen. Kondisi ini berpotensi mengakibatkan hilangnya peluang perolehan profit bagi Perusahaan [1]. Informasi yang akurat mengenai ketersediaan barang dan harga jual memiliki peran vital dalam pengembangan dan optimalisasi strategi pemasaran. Oleh karena itu, diperlukan suatu metode pengolahan data yang tepat dengan mengimplementasikan teknik data mining [2]. Dalam penelitian ini, metode K-Means dipilih sebagai teknik data mining yang diterapkan untuk meningkatkan efisiensi dan efektivitas dalam pengendalian inventaris. Implementasi algoritma K-Means dalam pembentukan *cluster* didasarkan pada model *Recency, Frequency, dan Monetary* (RFM), yang merupakan pendekatan analitis dalam memahami perilaku konsumen dan mendukung manajemen inventaris. Integrasi antara metode RFM dan algoritma K-Means dilakukan dengan menjadikan nilai *Recency, Frequency, dan Monetary* sebagai variabel input dalam proses klusterisasi. Nilai RFM dihitung dari data transaksi penjualan yang mencerminkan perilaku product terhadap frekuensi penjualan, waktu transaksi terakhir, dan nilai total penjualan. Setelah dinormalisasi, ketiga variabel tersebut digunakan sebagai fitur utama dalam proses K-Means untuk mengelompokkan produk dengan karakteristik serupa. Pendekatan ini sejalan dengan penelitian sebelumnya [3][4] yang menunjukkan efektivitas model RFM-KMeans dalam segmentasi pelanggan dan manajemen inventaris.

Aplikasi merupakan salah satu subkelas dari perangkat lunak komputer yang secara langsung memanfaatkan kapabilitas komputasional untuk melaksanakan berbagai tugas spesifik sesuai dengan kebutuhan pengguna [5]. Sementara itu, administrasi pada dasarnya berkaitan dengan serangkaian tugas atau

pekerjaan dalam suatu organisasi yang melibatkan peran administrator [6]. Administratif yang baik sangat membantu dalam berbagai usaha pengolahan data. Dalam konteks sistem informasi modern, aplikasi tidak hanya berfungsi sebagai alat bantu administrasi, tetapi juga sebagai sarana mengolah data dan menganalisis data dalam jumlah besar secara efisien. Salah satu pendekatan analisis data yang banyak digunakan dalam pengembangan aplikasi berbasis pengambilan keputusan adalah data mining. Data mining memainkan peran penting sebagai suatu proses identifikasi anomali, pola, atau korelasi dalam himpunan data berskala besar dengan tujuan memprediksi hasil. Fondasi data mining sendiri berakar pada berbagai disiplin ilmu, mencakup statistika, *Artificial Intelligence* (AI), pembelajaran mesin (*machine learning*), dan teknologi basis data, yang bersumber dari pencatatan administratif penjualan yang sistematis. Salah satu metode yang diimplementasikan dalam penelitian ini adalah algoritma K-Means sebagai suatu pendekatan yang mengambil dataset tidak berlabel sebagai input, kemudian membagi dataset tersebut menjadi sejumlah k kluster, dan melakukan iterasi proses tersebut hingga mencapai kluster optimal [7]. Perlu dicatat bahwa nilai k harus ditentukan sebelumnya dalam implementasi algoritma ini. Lebih lanjut, *clustering* sebagai suatu teknik dalam *data mining* yang bertujuan untuk mengelompokkan objek (data) ke dalam beberapa *cluster* atau kelompok, sehingga objek-objek yang memiliki kemiripan dapat diklasifikasikan ke dalam *cluster* yang sama [8]. Dalam konteks penelitian ini, teknik *clustering* diaplikasikan untuk mengembangkan strategi pemasaran berdasarkan variabel harga jual dan tingkat ketersediaan produk. Untuk memperjelas arsitektur sistem yang dikembangkan, penelitian ini menggunakan *Unified Modelling Language* (UML) sebagai suatu metodologi pemodelan yang bertujuan untuk memfasilitasi proses perancangan sistem, sehingga dapat meminimalisasi potensi kesalahan dalam tahap pengembangan program [9].

2. METODE

Metode penelitian ini menggunakan K-Means yang berfokus pada prinsip kedekatan antar data. Langkah-langkah implementasi meliputi pengumpulan data transaksi, perhitungan nilai RFM, normalisasi data menggunakan Min-Max, penentuan jumlah cluster dalam penulisan ini adalah pada nilai $k=3$, dan penerapan algoritma K-Means hingga konvergen [10].

2.1 Prinsip operasional RFM

RFM dihitung dari data transaksi penjualan produk dengan indikator tanggal transaksi, jumlah terjual dan harga satuan. Indikator *Recency* (R) menginterpretasikan semakin kecil nilai maka produk baru saja terjual. *Frequency* (F) semakin besar menunjukkan produk sering terjual dan *Monetary* (M) semakin besar maka produk semakin bernilai tinggi.

2.2 Prinsip operasional K-Means

Prinsip operasional algoritma K-Means melibatkan beberapa tahapan kunci. Pertama, dilakukan penentuan jumlah k *cluster* yang diinginkan, diikuti dengan penetapan titik *centroid* awal secara acak. Selanjutnya, proses alokasi data ke *cluster* terdekat dilakukan, dan tahapan ini diiterasi hingga tercapai kondisi *centroid* yang stabil. Perlu dicatat bahwa *output* dari algoritma K-Means sangat bergantung pada dua faktor utama: penentuan jumlah *cluster* dan pemilihan *centroid* awal yang dilakukan secara acak. Salah satu permasalahan yang sering muncul dalam implementasi algoritma K-Means adalah hasil *centroid* akhir yang tidak sepenuhnya merepresentasikan pusat *cluster* yang sesungguhnya. Untuk mengatasi hal ini, dalam praktiknya, algoritma ini perlu dijalankan secara berulang dengan variasi *centroid* awal yang berbeda-beda guna mendapatkan *centroid* akhir yang dianggap paling optimal [11]. Langkah-langkah implementasi algoritma K-Means secara umum dapat dijabarkan sebagai berikut:

- Inisialisasi *Centroid*: tahap ini melibatkan penetapan titik pusat awal untuk setiap kluster yang telah ditentukan.
- Penugasan *cluster*: pada tahap ini, setiap titik data diberi label sesuai dengan kluster terdekatnya. Penentuan kedekatan ini umumnya didasarkan pada perhitungan jarak Euclidean antara titik data dengan *centroid cluster*, persamaan 1.
Implementasi algoritma K-Means yang efektif memerlukan pemahaman mendalam tentang karakteristik data yang dianalisis serta pertimbangan cermat dalam penentuan parameter-parameter kunci seperti jumlah *cluster* dan metode inisialisasi *cluster*.

$$d(p, q) = \sqrt{\sum_{i=1}^n (p_i - q_i)^2} \quad (1)$$

Keterangan:

p, q : dua titik data dalam ruang n -dimensi.

p_i dan q_i : koordinat ke- i dari titik data p dan q
 n : jumlah dimensi dalam dataset
 $d(p, q)$: jarak *Euclidean*

- c. Pembaharuan *centroid* dimana *centroid* baru dihitung sebagai rata-rata dari semua titik data dalam *cluster*, Persamaan 2.

$$C_j = \frac{1}{n_j} \sum_{i=1}^{n_j} X_i \quad (2)$$

Keterangan:

C_j : *centroid* baru dari cluster j

n_j : jumlah titik data dalam *cluster* j x_i : titik data ke- i dalam *cluster* j

Kriteria konvergensi dan bilaman algoritma K-Means akan berhenti, apabila telah terjadi :
 Evaluasi hasil clustering

Menghitung pusat klaster dengan anggota klaster yang baru.

Proses *cluster* sudah selesai bila pusat klaster tidak berubah, namun jika pusat cluster masih berubah maka ulangi langkah menghitung jarak hingga pusat klaster tidak berubah lagi.

3. HASIL DAN PEMBAHASAN

Algoritma K-Means merupakan suatu metodologi analisis data atau teknik penambangan data (*data mining*) yang melaksanakan proses pemodelan tanpa pengawasan (*unsupervised modeling*). Metode ini termasuk dalam kategori teknik pengelompokan data yang menggunakan sistem partisi. Tujuan utama dari algoritma K-Means adalah melakukan segmentasi sekumpulan data menjadi beberapa kelompok (*cluster*) berdasarkan kemiripan karakteristik atau fitur. Prinsip operasional algoritma K-Means berfokus pada minimalisasi jumlah kuadrat jarak antara titik-titik data dengan pusat cluster atau *centroid*. Efektivitas algoritma ini terletak pada kemampuannya untuk mengidentifikasi pola-pola intrinsik dalam dataset tanpa memerlukan label atau kategori yang telah ditentukan sebelumnya. Dalam konteks penelitian ini, sejumlah data inventaris digunakan sebagai masukan untuk proses clustering. Implementasi algoritma K-Means dalam analisis stok barang melibatkan serangkaian tahapan sistematis, yang dapat dijabarkan sebagai berikut:

- a. Data Stok Barang, berikut adalah contoh data stok barang (10 data) dengan fitur jumlah stok dan harga:

Tabel 1. Data stok barang

ID	Jumlah stok	Harga (Rp)	Recency (hari)	Frequency	Monetary (Rp)
1	50	20.000	3	12	1.200.000
2	20	5.000	7	5	100.000
3	30	7.000	5	7	210.000
4	70	30.000	2	10	3.000.000
5	90	45.000	10	3	1.350.000
6	60	15.000	4	8	1.200.000
7	10	3.000	11	2	60.000
8	80	25.000	0	11	2.750.000
9	40	10.000	8	4	400.000
10	55	18.000	1	9	1.620.000

Kemudian dilakukan normalisasi nilai-nilai RFM sebelum diintegrasikan pada algoritma K-Means menggunakan metode Min-Max, persamaan 3. Dengan hasil sebagaimana ditunjukkan pada tabel 2.

$$X' = \frac{X - X_{min}}{X_{max} - X_{min}} \quad (3)$$

Keterangan

X' : nilai hasil normalisasi

X : data awal
 Xmin : nilai terkecil dalam kolom atau fitur yang sama
 Xmax : nilai terbesar dalam kolom atau fitur yang sama

Tabel 2. Hasil Normalisasi RFM

ID	R	F	M
1	0,27	1,00	0,388
2	0,64	0,30	0,014
3	0,45	0,50	0,051
4	0,18	0,80	1,000
5	0,91	0,10	0,439
6	0,36	0,60	0,388
7	1,00	0,00	0,000
8	0,00	0,90	0,917
9	0,73	0,20	0,116
10	0,09	0,70	0,531

Kemudian melakukan integrasi hasil nilai normalisasi RFM tersebut ke dalam algoritma K-Means dimana ketiga variabel R, F dan M menjadi matriks fitur. Untuk selanjutnya dilakukan tahapan clustering.

- b. Inisialisasi *Centroid*, dalam penelitian ini dipilih $k=3$ (tiga kluster) artinya tiga cluster sudah cukup merepresentasikan variasi data.
- c. *Assign Clusters*, hitung jarak Euclidean dari setiap data poin ke *centroid* dan tetapkan ke *centroid* terdekat terdekat. Kami melakukan perhitungan pada data point ID 2 (0,64, 0,30 dan 0,014):

$$\text{Jarak ke Centroid 1: } \sqrt{(0,64 - 0,27)^2 + (0,30 - 1,00)^2 + (0,014 - 0,388)^2} = 0,825$$

$$\text{Jarak ke Centroid 2: } \sqrt{(0,64 - 0,91)^2 + (0,30 - 0,10)^2 + (0,014 - 0,439)^2} = 0,481$$

$$\text{Jarak ke Centroid 3: } \sqrt{(0,64 - 0,00)^2 + (0,30 - 0,90)^2 + (0,014 - 0,917)^2} = 1,160$$

Pengelompokan awal, karena jarak ke *Centroid* ID 2 paling kecil, data poin ID 2 masuk ke cluster 2. Setelah melakukan perhitungan ini untuk semua data poin, hasil dari perhitungan adalah:

Cluster 1 dengan anggota ID 1,4,6 dan 10

Cluster 2 dengan anggota ID 2,3,5,7 dan 9

Cluster 3 dengan anggota ID 8

- d. *Update Centroid*, Hitung ulang posisi *centroid* sebagai *mean* dari data poin dalam tiap *cluster*
 Pada cluster 1: (ID 1,4,6 dan 10)

$$R_{\text{baru}} = \frac{(0,27+0,18+0,36+0,09)}{4} = 0,225$$

$$F_{\text{baru}} = \frac{(1,00+0,80+0,60+0,70)}{4} = 0,775$$

$$M_{\text{baru}} = \frac{(0,388+1,000+0,388+0,531)}{4} = 0,577$$

- e. Iterasi, mengulangi langkah 2 dan 3 hingga posisi *centroid* tidak berubah atau perubahan sangat kecil.

Hasil akhir perhitungan setelah tercapai *centroid* yang stabil, barang-barang diklasifikasikan menjadi tiga cluster berdasarkan karakteristik terdekat:

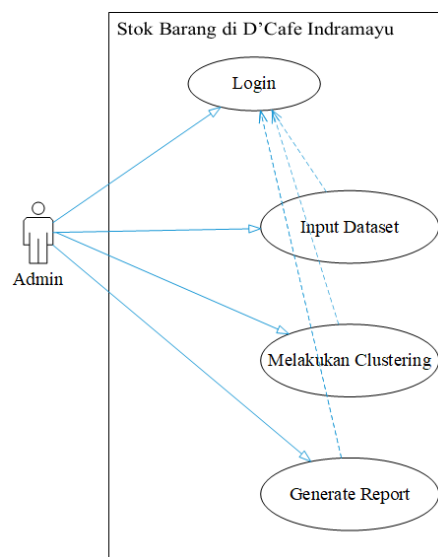
- a. Cluster 1: R rendah, F tinggi, M tinggi maka produk aktif dan bernilai tinggi dengan ID 1,4,6, 8 dan 10
- b. Cluster 2: R sedang, F rendah-sedang, M rendah maka produk dengan pergerakan menengah dengan ID

2,3, 5 dan 9

c. Cluster 3: R tinggi, F rendah, M rendah maka produk jarang terjual atau lambat bergerak dengan ID 7
Algoritma K-Means berhasil mengklasifikasikan barang-barang ke dalam cluster yang berbeda berdasarkan kemiripan fitur (jumlah stok dan harga). Dengan beberapa iterasi, algoritma ini menemukan posisi centroid yang optimal dan mengelompokkan barang-barang secara efektif, sehingga barang dengan karakteristik yang mirip ditempatkan dalam cluster yang sama.

Pemodelan Perangkat Lunak

Diagram *use case* berikut representasi visual yang menggambarkan interaksi antara sistem dan aktor eksternal yang terlibat dalam proses, dalam konteks ini, administrator sistem. Administrator memiliki kapabilitas untuk mengelola seluruh data yang tersimpan dalam basis data sistem, mengidentifikasi dan mendokumentasikan fungsionalitas yang terdapat dalam sistem stok barang di D'Cafe. Implementasi diagram *use case* dalam konteks pengembangan sistem informasi memiliki beberapa manfaat signifikan:



Gambar 1. *Use case* sistem klasifikasi inventaris

Diagram *use case* diatas menggambarkan interaksi antara sistem dan aktor eksternal yang terlibat, Admin dapat mengelola semua data yang disimpan dalam database di dalam sistem. Penjelasan pada gambar dijelaskan pada skenario dibawah ini:

Skenario Use Case Diagram Login

Nama Use Case : Login

Aktor : Admin

Deskripsi : Dalam proses ini admin memasukan *username* dan *password*, kemudian sistem akan Melakukan validasi dengan database, jika data valid maka masuk ke dalam sistem. Jika tidak maka akan tetap berada dalam kondisi untuk login.

Skenario Use Case Diagram Input Dataset

Nama Use Case : Input Dataset

Aktor : Admin

Deskripsi : Dalam menu ini terdapat submenu dataset penjualan. Terdapat kolom yang harus diinput, berupa data kode penjualan, tanggal transaksi, nama barang jumlah dan total harga.

Skenario Use Case Diagram Proses Clustering

Nama Use Case : Proses Clustering

Aktor : Admin

Deskripsi : Dalam menu ini terdapat tabel perhitungan K-Means, berupa Transformasi RFM, Normalisasi RFM, Clustering KMeans Akhir dan Hasil Clustering K-Means.

Skenario Use Case Diagram Generate Report

Implementasi Sistem Klasifikasi Inventaris Menggunakan Metode Clustering dengan....(Dina Rima Dhona)

Nama Use Case : Generate Report
 Aktor : Admin
 Deskripsi : Dalam menu ini terdapat Laporan yang akan dicetak, berupa Laporan Dataset Penjualan, Laporan Transformasi RFM, Laporan Perhitungan Cluster dan Laporan Hasil Perhitungan K-Means

SIMPULAN DAN SARAN

Penelitian ini berhasil mengimplementasikan metode *clustering* dengan algoritma K-Means berbasis model RFM dalam klasifikasi persediaan stok barang. Hasil menunjukkan tiga *cluster* produk yang berbeda berdasarkan karakteristik harga dan ketersediaan stok. Aplikasi yang dikembangkan memberikan dukungan bagi pengambil keputusan dalam manajemen inventaris dan strategi penjualan. Keterbatasan penelitian ini adalah bahwa implementasi belum diuji secara penuh di lingkungan operasional dan belum dibandingkan dengan algoritma *clustering* lainnya. Penelitian selanjutnya disarankan untuk melakukan uji coba implementasi langsung serta mengeksplorasi metode *clustering* lain seperti fuzzy C-Mean untuk perbandingan efektivitas.

DAFTAR PUSTAKA

- [1] M. Rasyidan and Z. Zaenuddin, "Perancangan Sistem Informasi Persediaan Barang Menggunakan Metode Average (Studi Kasus Toko Nazar Banjarmasin)," *Technol. J. Ilm.*, vol. 11, no. 4, p. 191, 2020, doi: 10.31602/tji.v11i4.3638.
- [2] I. Saputra, P. Irfansyah, E. D. Sirait, D. D. Apriyani, and M. Sonny, "Comparison of the Performance of the k-Nearest Neighbor, Naïve Bayes Classifier and Support Vector Machine Algorithm With SMOTE for Classification of Bully Behavior on the WhatsApp Messenger Application," in *Proceedings of the 1st International Conference on Folklore, Language, Education and Exhibition (ICOFLEX 2019)*, Atlantis Press, 2020, pp. 143–149. doi: <https://doi.org/10.2991/assehr.k.201230.028>.
- [3] D. S. Nugaha, I. Thoib, N. Sururi, F. B. F. and B. P. Candra, "Segmentasi Pelanggan Berbasis Analisis RFM Menggunakan Algoritma K- Means Clustering," vol. 4, no. 2, pp. 1361–1369, 2025.
- [4] Y. Syahra, A. Fadlil, and H. Yuliansyah, "Customer Segmentation Using RFM and K-Means Clustering to Support CRM in Retail Industry," *Sinkron*, vol. 9, no. 3, pp. 1120–1131, 2025, doi: 10.33395/sinkron.v9i3.14907.
- [5] M. Yusril Helmi Setyawan and A. S. Munari, *Panduan Lengkap Membangun Sistem Monitoring Kinerja Mahasiswa Intership Berbasis Web dan Global Positioning System*. Kreatif Industri Nusantara, 2020.
- [6] L. Marliani, "Definisi Administrasi Dalam Berbagai Sudut Pandang," *J. Fak. Ilmu Sos. dan Ilmu Polit. Univ. Galuh*, vol. 5, no. 4, pp. 17–18, 2018.
- [7] Trivusi, "K-Means Clustering: Pengertian, Cara Kerja, Kelebihan, dan Kekurangannya."
- [8] A. Kurnia, "Perbandingan Algoritma K-Means dan Fuzzy C-Means Untuk Clustering Puskesmas Berdasarkan Gizi Balita Surabaya," *J. Process.*, vol. 18, no. 1, pp. 83–88, 2023, doi: 10.33998/processor.2023.18.1.696.
- [9] S. Nabila, A. R. Putri, A. Hafizhah, F. H. Rahmah, and R. Muslikhah, "Pemodelan Diagram UML Pada Perancangan Sistem Aplikasi Konsultasi Hewan Peliharaan Berbasis Android (Studi Kasus: Alopét)," *J. Ilmu Komput. dan Bisnis*, vol. 12, no. 2, pp. 130–139, 2021, doi: 10.47927/jikb.v12i2.150.
- [10] A. Maftukhah, A. Fadlil, and S. Sunardi, "Butterfly Image Classification using Convolution Neural Network with AlexNet Architecture," *J. Infotel*, vol. 16, no. 1, pp. 82–95, 2024, doi: 10.20895/infotel.v16i1.1004.
- [11] M. Orisa, "Optimasi Cluster pada Algoritma K-Means," *Pros. SENIATI*, vol. 6, no. 2, pp. 430–437, 2022, doi: 10.36040/seniati.v6i2.5034.