

Penerapan K-Medoids Clustering Pada Data Penerima Program Z-Auto BAZNAS Jawa Timur dengan Metode Evaluasi Elbow

Imro'atus Sholihah^{1*}, Bayu Adhi Nugroho², Noor Wahyudi³, Yunita Ardilla⁴

^{1,2,3,4}UIN Sunan Ampel Surabaya, Surabaya, 60294, Indonesia

imroatussholihah117@gmail.com, n.wahyudi@uinsa.ac.id, yunita.ardilla@uinsa.ac.id

ABSTRAK

Penelitian ini bertujuan untuk menganalisis status penerima Program Z-Auto di BAZNAS Provinsi Jawa Timur dengan menggunakan metode *K-Medoids Clustering*. Program Z-Auto merupakan inisiatif BAZNAS untuk memberdayakan mustahik di bidang otomotif melalui bantuan modal, pelatihan, dan peralatan usaha bengkel. Data penelitian berupa 445 data pendaftar program dengan atribut yang dipilih usia, jenis usaha, pendapatan, dan pengeluaran bulanan. Dengan banyaknya data pendaftar dengan karakteristik yang beragam sehingga sulit untuk dianalisis secara manual, terutama dalam menentukan kelompok mustahik yang memiliki kondisi ekonomi serupa. Tanpa analisis berbasis data, proses seleksi berpotensi tidak objektif dan kurang tepat sasaran. Oleh karena itu, penelitian ini menerapkan *clustering* seperti *K-Medoids* untuk mengelompokkan data penerima secara lebih sistematis dan terukur. Alur penelitian dimulai dari tahapan data selection, data *preprocessing*, data transformation, data mining, dan evaluasi. Proses data mining menggunakan metode *K-medoids* sebagai metode clustering yang kemudian dievaluasi menggunakan *Elbow Method* dan *Silhouette Score* untuk mendapatkan jumlah cluster terbaik. Hasil penelitian menunjukkan bahwa data penerima terbagi menjadi tiga *cluster* utama: (1) kelompok usia lebih tua dengan pendapatan dan pengeluaran rendah, (2) kelompok dengan kondisi ekonomi sedang, dan (3) kelompok usia lebih muda dengan pendapatan serta pengeluaran relatif tinggi. Evaluasi dengan *Silhouette Score* memberikan hasil yang cukup baik (0,7134). Penelitian ini menunjukkan bahwa *K-Medoids clustering* dapat membantu BAZNAS dalam meningkatkan ketepatan distribusi bantuan secara objektif dan transparan.

Kata kunci: Clustering; K-Medoids; Z-Auto; Data Mining; BAZNAS

1. PENDAHULUAN

Perkembangan teknologi informasi telah memberikan peluang besar dalam pemanfaatan data untuk mendukung pengambilan keputusan di berbagai sektor, termasuk lembaga pengelola zakat. Badan Amil Zakat Nasional (BAZNAS) Provinsi Jawa Timur merupakan lembaga yang bertugas menghimpun dan mendistribusikan zakat secara profesional dan transparan. Salah satu inisiatif unggulannya adalah Program Z-Auto, yaitu program pemberdayaan ekonomi mustahik di bidang otomotif melalui bantuan modal usaha, pelatihan, dan peralatan bengkel [1]. Namun, mekanisme penentuan penerima bantuan selama ini masih mengandalkan seleksi administratif dan pengamatan manual, sehingga kurang efektif dalam mengelola data pendaftar yang besar dan beragam.

Berbagai penelitian menunjukkan bahwa metode data mining, khususnya teknik clustering, efektif dalam membantu pengelompokan data penerima bantuan sosial agar lebih objektif dan tepat sasaran. Salah satu algoritma yang populer digunakan adalah K-Medoids Clustering, karena mampu mengelompokkan data dengan cepat berdasarkan kesamaan karakteristik [2]. Beberapa studi sebelumnya telah berhasil menerapkan metode ini dalam konteks pendidikan, bisnis, dan distribusi bantuan sosial, yang membuktikan relevansinya dalam mengidentifikasi pola dan segmen penerima program.

Pemanfaatan data mining menjadi semakin penting dalam menganalisis data berukuran besar yang sulit diolah secara manual. Salah satu teknik populer dalam data mining adalah

clustering, yaitu proses pengelompokan data berdasarkan kesamaan karakteristik tertentu tanpa memerlukan label kelas. Metode ini efektif digunakan untuk mengidentifikasi pola tersembunyi dalam data yang kompleks, sehingga dapat mendukung proses pengambilan keputusan berbasis data [3].

Dari berbagai algoritma clustering, K-Medoids merupakan metode yang sering dipilih dalam clustering karena pusat klusternya adalah medoid (yaitu data asli) bukan centroid rata-rata yang dapat terpengaruh oleh outlier, sehingga hasil cluster menjadi lebih robust dan representatif terhadap data yang ekstrem. Meskipun demikian, K-Medoids memiliki kompleksitas komputasi yang lebih tinggi dibanding K-Means, terutama terkait perhitungan jarak antar semua titik serta proses pemilihan medoid yang optimal.[4]

Penelitian terdahulu menunjukkan bahwa K-Medoids telah digunakan dalam berbagai bidang, misalnya di bidang pendidikan `pengelompokan mahasiswa yang berpotensi drop out, yaitu penelitian oleh S. Bahri and D. M. Midyanti, yang berjudul “Penerapan Metode K-Medoids untuk Pengelompokan Mahasiswa Berpotensi Drop Out” pada tahun 2023 [5], bidang bisnis penelitian oleh A. A. D. Sulistyawati and M. Sadikin, yang berjudul “Penerapan Algoritma K-Medoids untuk Menentukan Segmentasi Pelanggan,” pada tahun 2021 [6], hingga analisis untuk penjualan penelitian oleh Y. Diana and F. Hadi, yang berjudul “Analisa Penjualan Menggunakan Algoritma K-Medoids untuk Mengoptimalkan Penjualan Barang,” pada tahun 2023 [7]. Hasilnya membuktikan bahwa metode ini efektif dalam mengidentifikasi kelompok dengan karakteristik berbeda.

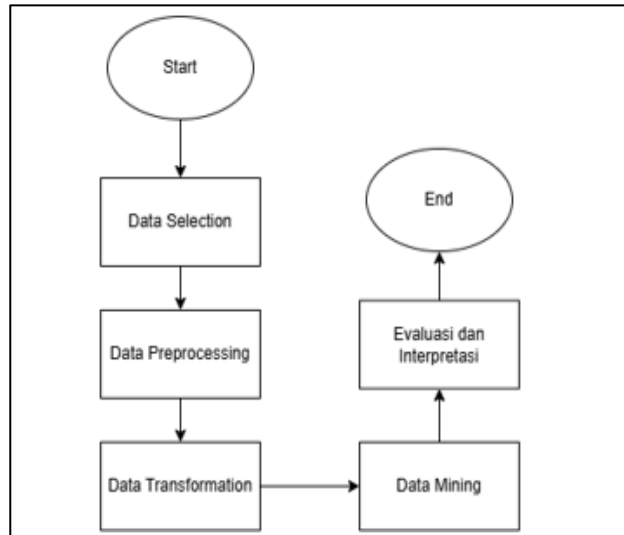
Dalam konteks pengelolaan zakat, BAZNAS sebagai lembaga resmi yang menghimpun dan menyalurkan zakat perlu menerapkan pendekatan berbasis data agar distribusi bantuan lebih tepat sasaran. Program Z-Auto yang berfokus pada pemberdayaan mustahik di sektor otomotif menjadi salah satu contoh program strategis yang memerlukan analisis tersebut. Dengan menerapkan metode K-Medoids, pendaftar program dapat dikelompokkan berdasarkan faktor usia, jenis usaha, pendapatan, dan pengeluaran bulanan. Hal ini diharapkan dapat membantu BAZNAS dalam menentukan prioritas penerima, meningkatkan transparansi, serta memperkuat akuntabilitas dalam pengelolaan dana zakat.

Berdasarkan hal tersebut, penelitian ini dilakukan dengan tujuan untuk menerapkan metode K-Medoids Clustering dalam menganalisis data pendaftar Program Z-Auto di BAZNAS Provinsi Jawa Timur. Penelitian ini berfokus pada pengelompokan penerima berdasarkan faktor usia, jenis usaha, pendapatan, dan pengeluaran bulanan, sehingga diharapkan dapat memberikan masukan bagi BAZNAS dalam meningkatkan efektivitas dan transparansi distribusi zakat. Dengan adanya analisis berbasis data, program penyaluran bantuan dapat lebih terukur, tepat sasaran, serta mendukung akuntabilitas lembaga zakat di era digital.

2. METODE

Penelitian ini menggunakan pendekatan kuantitatif dengan metode K-Medoids Clustering. Data yang digunakan berupa 445 data pendaftar Program Z-Auto tahun 2025 yang diperoleh dari bagian pendistribusian BAZNAS Provinsi Jawa Timur. Atribut yang

dianalisis meliputi usia, jenis usaha, jenis kelamin, pendapatan, dan pengeluaran bulanan. Tahapan penelitian meliputi Gambar 1.



Gambar 1. Flowchart Alur Penelitian

Gambar 1. menunjukkan tahapan dalam proses penelitian, yaitu:

- 1) *Data Selection*: Tahapan metodologi dimulai dengan pemilihan atribut yang relevan dengan tujuan penelitian.
- 2) *Data Preprocessing*: Setelah data terkumpul, dilakukan proses pra-pemrosesan data (*data preprocessing*). Tahap ini mencakup pembersihan data dari nilai kosong atau tidak valid.
- 3) *Data Transformation*: Tahap selanjutnya adalah penentuan jumlah *cluster* (nilai *k*) yang optimal menggunakan metode Elbow dan Silhouette Score.
- 4) *Data Mining*: Setelah jumlah *cluster* ditentukan, dilakukan proses *clustering* menggunakan algoritma *K-Means*.
- 5) *Evaluasi dan Interpretasi*: Selanjutnya, dilakukan evaluasi dan interpretasi hasil *clustering* menggunakan Silhouette Score untuk mengukur kualitas hasil clustering.

Langkah-langkah *clustering* menggunakan algoritma *k-medoids* [8]:

- 1) Tentukan terlebih dahulu jumlah cluster yang diinginkan (*k*) dan pilih titik awal sebagai pusat cluster.
- 2) Hitung jarak setiap data terhadap pusat cluster terdekat menggunakan rumus Euclidean Distance (1)

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

dengan keterangan:

- $d(x, y)$ = jarak antara data ke-*i* dan data ke-*j*
- x_i = nilai atribut ke-*i* dari suatu data
- y_i = nilai atribut ke-*i* dari data pembanding
- n = jumlah atribut yang digunakan

- 3) Pilih secara acak sebuah objek dari data sebagai kandidat medoid pengganti.
- 4) Hitung jarak setiap data terhadap kandidat medoid tersebut dan kelompokkan ke cluster terdekat.
- 5) Hitung total biaya/penyimpangan (Total Dissimilarity, TD). Jika nilai TD baru lebih kecil dari TD sebelumnya, maka medoid diganti dengan kandidat baru.
- 6) Mengulangi proses penugasan dan pembaruan centroid (langkah 3–5) sampai posisi centroid stabil (konvergen) atau mencapai batas iterasi yang telah ditentukan.

3. HASIL DAN PEMBAHASAN

Hasil klasterisasi pada data penerima Program Z-Auto bertujuan untuk mengelompokkan mustahik berdasarkan kesamaan karakteristik sosial ekonomi, seperti tingkat pendapatan dan pengeluaran yang rendah, menengah, maupun tinggi. Analisis ini memberikan gambaran pola distribusi kelompok yang terbentuk dari data pendaftar. Pembahasan hasil pengelompokan diharapkan dapat memberikan pemahaman yang lebih mendalam mengenai ciri khas tiap cluster, mendukung BAZNAS dalam merancang strategi penyaluran bantuan yang lebih tepat sasaran, serta membantu mengidentifikasi kebutuhan intervensi tambahan bagi kelompok dengan kondisi ekonomi lebih lemah.

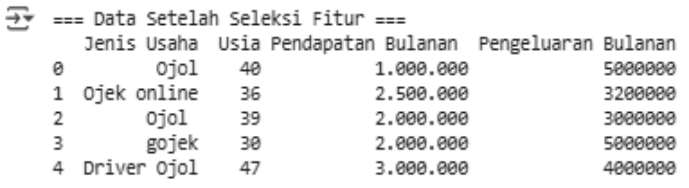
3.1 Data Selection

Tahapan metodologi dimulai dengan pengumpulan data, yaitu data sekunder berupa data excel sebanyak 445 data yang berasal dari bagian pendistribusian BAZNAS Provinsi Jawa Timur yang memuat informasi terkait penerima bantuan program Z-Auto.

	A	B	C	D	E	F	G
1	Jenis Usaha	Jenis Kelamin	Usia	Jenis Motor	Merk Oli	Pendapatan Bulanan	Pengeluaran Bulanan
2	Ojol	Perempuan	40	Beat	Matic Federal	1.000.000	5.000.000
3	Ojek online	Laki-Laki	36	Vario 110	Matic Enduro	2.500.000	3.200.000
4	Ojol	Perempuan	39	Matic, Vario 125 cc	Matic Enduro	2.000.000	3.000.000
5	gojek	Perempuan	30	beat 110	Matic Federal	2.000.000	5.000.000
6	Driver Ojol	Perempuan	47	Manual Revo	Manual Federal	5.000.000	4.000.000
7	Gojek	Laki-Laki	43	Vario 125	Matic Enduro	3.000.000	4.000.000
8	Ojol	Perempuan	30	Matic beat 110 cc	Matic Federal	1.000.000	3.500.000
9	Ojol	Perempuan	32	Revo 2008 110cc	Manual Enduro	2.500.000	5.000.000
10	Driver Online	Laki-Laki	39	Matic, lexii 125	Matic Enduro	10.000.000	5.200.000
11	Ojol	Perempuan	50	Beat 110 cc	Matic Federal	900.000	1.500.000
12	OJEK ONLINE	Perempuan	45	MATIC, BEAT 110 CC	Matic Federal	1.000.000	2.500.000
13	Ojol	Laki-Laki	28	Vario 125 tahun 2013	Matic Federal	2.000.000	3.000.000
14	ojek online	Laki-Laki	37	matic, vario 125 cc	Matic Enduro	1000.000	4.000.000
15	Ojol	Perempuan	43	Scoopy 108CC	Matic Federal	2.000.000	2.000.000
16	Snack kecipir	Perempuan	44	Matic 110	Matic Federal	3.000.000	5.000.000
17	Ojek online	Laki-Laki	38	Honda revo fit 110cc	Manual Enduro	3.000.000	1.500.000
18	Driver on line	Perempuan	49	Manual Yamaha 113 cc	Manual Enduro	5.000.000	4.500.000
19	Ojek online	Perempuan	40	Beat 110cc	Matic Enduro	2.000.000	2.000.000
20	Ojek Online (Gojek)	Laki-Laki	33	BEAT	Matic Enduro	6.500.000	5.000.000
21	Toko kelontong	Perempuan	42	Genio 125 cc	Matic Enduro	1.000.000	1.000.000

Gambar 2. Dataset Penerima Z-Auto

Pada tahap data selection, peneliti memilih data yang relevan sesuai tujuan penelitian. Atribut-atribut yang digunakan difokuskan pada variabel yang mempengaruhi proses clustering, sedangkan data yang tidak relevan atau redundan diabaikan. Selain itu, data disaring agar sesuai dengan ruang lingkup penelitian, sehingga hasil analisis lebih akurat, efisien, dan sesuai dengan kebutuhan.



	Jenis Usaha	Usia	Pendapatan Bulanan	Pengeluaran Bulanan
0	Ojol	40	1.000.000	500000
1	Ojek online	36	2.500.000	320000
2	Ojol	39	2.000.000	300000
3	gojek	30	2.000.000	500000
4	Driver Ojol	47	3.000.000	400000

Gambar 3. Data Hasil Seleksi Fitur

3.2 Data Preprocessing

Setelah data terkumpul, dilakukan proses pra-pemrosesan data (data preprocessing). Tahap ini mencakup pembersihan data dari nilai kosong atau tidak valid. Dengan uraian data dari nilai kosong (NaN) per kolom 16 sebanyak 33, serta data tidak valid dari kolom pendapatan bulanan sebanyak 106 dan kolom pengeluaran bulanan sebanyak 116. Normalisasi data untuk menyamakan skala antar atribut, serta seleksi atribut agar hanya fitur yang relevan yang digunakan dalam proses analisis. Normalisasi sangat penting agar algoritma K-Means tidak bias terhadap fitur dengan skala besar.

Pada tahap data preprocessing, dilakukan serangkaian proses untuk memastikan data siap digunakan dalam analisis clustering. Langkah pertama adalah data cleaning, yaitu menghapus data duplikat, mengisi nilai yang hilang dengan pendekatan tertentu, serta memperbaiki inkonsistensi data. Selanjutnya dilakukan penanganan outlier dengan mendeteksi nilai ekstrim yang dapat mengganggu hasil analisis, misalnya data pendapatan yang terlalu rendah atau tinggi tidak wajar.

Tahap berikutnya, di mana variabel kategorikal seperti jenis kelamin diubah menjadi numerik, dan data numerik seperti pendapatan serta jumlah tanggungan dinormalisasi agar berada pada skala yang sama. Jika data berasal dari berbagai sumber, maka dilakukan data integration untuk menyatukan seluruh informasi menjadi dataset yang konsisten. Selain itu, atribut yang tidak relevan dengan analisis dieliminasi melalui data reduction, sehingga hanya variabel penting yang dipertahankan. Dengan demikian, hasil preprocessing menghasilkan dataset yang bersih, konsisten, dan siap untuk dianalisis lebih lanjut menggunakan metode clustering.

3.3 Data Transformation

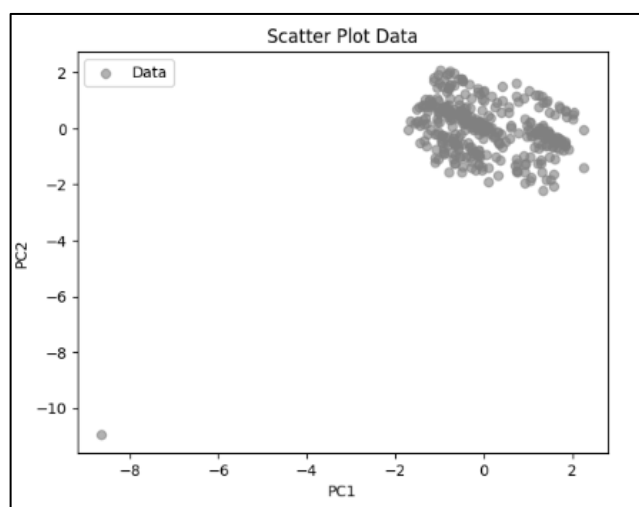
Pada tahap transformasi data, dataset yang tersedia diolah menjadi dataset baru dengan format berbeda dari sebelumnya. Proses ini bertujuan untuk menyesuaikan data agar sesuai dengan kebutuhan analisis maupun model yang akan digunakan [9]. Dalam penelitian ini, dilakukan transformasi data sehingga data dapat diproses lebih lanjut pada tahap berikutnya, yaitu penerapan metode data mining.

3.4 Data Mining

Data mining merupakan salah satu tahapan dalam proses *Knowledge Discovery in Databases* (KDD) yang meliputi kegiatan pembersihan, integrasi, pemilihan, transformasi, penambahan data, evaluasi pola, hingga penyajian informasi [9]. Pada penelitian ini, data mining dilakukan menggunakan teknik *clustering* dengan algoritma K-Medoids.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 444 entries, 0 to 443
Data columns (total 4 columns):
#   Column                Non-Null Count  Dtype
---  -
0   Jenis Usaha           444 non-null    object
1   Usia                  444 non-null    int64
2   Pendapatan Bulanan    444 non-null    object
3   Pengeluaran Bulanan   444 non-null    object
dtypes: int64(1), object(3)
memory usage: 14.0+ KB
None
```

Gambar 4. Tampilan Informasi Atribut



Gambar 5. Tampilan Scatter Data

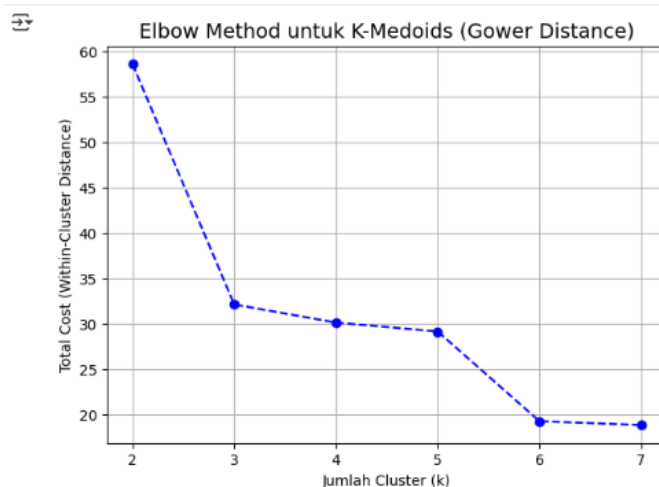
Pada gambar 5. Grafik yang ditampilkan merupakan scatter plot data hasil reduksi dimensi menggunakan Principal Component Analysis (PCA). Pada sumbu horizontal ditunjukkan komponen utama pertama (PC1), sedangkan sumbu vertikal menunjukkan komponen utama kedua (PC2). Setiap titik pada grafik merepresentasikan satu data penerima Program Z-Auto dengan atribut yang telah dipilih (usia, jenis usaha, pendapatan, dan pengeluaran bulanan) yang kemudian diproyeksikan ke dalam ruang dua dimensi.

Sebagian besar data terlihat mengelompok pada area tertentu, khususnya di bagian kanan atas grafik, yang menunjukkan adanya kecenderungan pola dan kemiripan karakteristik antar penerima. Sementara beberapa titik yang terletak agak jauh dari kelompok utama dapat dianggap sebagai data yang berbeda karakteristiknya, bahkan berpotensi sebagai outlier. Visualisasi ini memberikan gambaran awal bahwa data memang memiliki struktur yang bisa dikelompokkan lebih lanjut melalui metode clustering.

3.5 Evaluasi dan Interpretasi

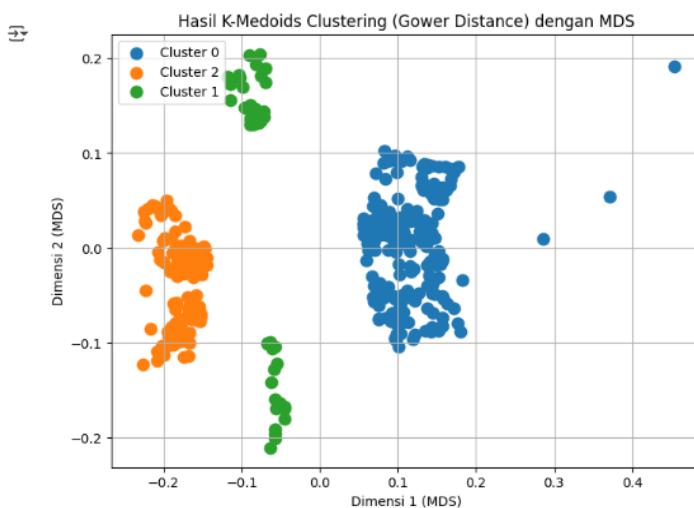
Penentuan jumlah *cluster* optimal (K) pada proses *clustering* menggunakan metode *K-Medoids* dilakukan dengan pendekatan metode *Elbow*. Metode ini bertujuan untuk

menemukan titik optimal di mana penambahan jumlah *cluster* tidak lagi memberikan penurunan yang signifikan terhadap nilai inertia atau *Sum of Squared Errors* (SSE).



Gambar 6. Hasil Evaluasi *Elbow Method*

Dalam visualisasi grafik *Elbow*, nilai inertia menurun drastis dari K=1 hingga K=3, namun setelah K=3, penurunannya mulai melandai. Titik pelandaian ini disebut sebagai "*elbow point*" yang mengindikasikan jumlah *cluster* optimal. Oleh karena itu, dipilih nilai K=3 karena pada titik ini model sudah cukup baik dalam memisahkan data ke dalam *cluster* yang berbeda tanpa menimbulkan *overfitting* atau pembentukan klaster yang terlalu kecil dan mirip.



Gambar 6. Grafik Hasil *Clustering*

Hasil klasterisasi menunjukkan terbentuknya tiga kelompok utama. *Cluster 0* "green" didominasi oleh peserta dengan usia relatif lebih tua serta memiliki pendapatan dan pengeluaran yang rendah, sehingga dapat dikategorikan sebagai kelompok ekonomi menengah ke bawah. *Cluster 1* "orange" merepresentasikan kelompok dengan karakteristik rata-rata, baik dari segi usia, pendapatan, maupun pengeluaran, sehingga dapat dianggap sebagai gambaran umum populasi. Sementara itu, *Cluster 2* "blue" berisi peserta yang

cenderung lebih muda dengan pendapatan dan pengeluaran lebih tinggi dibandingkan *cluster* lainnya, sehingga mencerminkan kelompok yang lebih sejahtera dan aktif secara ekonomi.

Berdasarkan hasil pengujian, diperoleh nilai *Silhouette Score* sebesar 0,7134. Menurut teori, *Silhouette Score* digunakan untuk mengevaluasi kualitas *cluster* dengan rentang nilai antara -1 hingga +1. Skor yang mendekati +1 menandakan bahwa *cluster* menunjukkan sangat baik karena memiliki kedekatan tinggi dengan *cluster* sendiri dan jauh dari *cluster* lain. Skor mendekati 0 menunjukkan adanya tumpang tindih antar *cluster*, sedangkan skor negatif menandakan objek cenderung salah penempatan. Secara umum, nilai di atas 0,5 dapat dikategorikan baik hingga sangat baik [10]. Dengan demikian, nilai 0,7134 yang diperoleh dalam penelitian ini termasuk kategori sangat baik, yang berarti *cluster* yang terbentuk cukup padat dan terpisah jelas antar kelompok. Hal ini menunjukkan bahwa metode clustering yang digunakan mampu menghasilkan pengelompokan yang representatif serta dapat digunakan untuk analisis lanjutan.

4. KESIMPULAN

Penelitian ini menunjukkan bahwa metode K-Medoids Clustering efektif digunakan untuk menganalisis data penerima Program Z-Auto BAZNAS Jawa Timur. Dengan 445 data pendaftar dan atribut usia, jenis usaha, pendapatan, serta pengeluaran bulanan, hasil analisis membagi mustahik ke dalam tiga *cluster* utama dengan kualitas pengelompokan yang cukup baik (*Silhouette Score* 0,7134). Temuan ini membuktikan bahwa K-Medoids dapat membantu BAZNAS dalam memahami karakteristik penerima dan meningkatkan ketepatan distribusi bantuan secara lebih objektif dan transparan. Dengan demikian, hasil analisis ini dapat menjadi bahan pertimbangan dalam perumusan strategi pemberdayaan yang lebih tepat sasaran. Sebagai tindak lanjut, penelitian berikutnya disarankan untuk mencoba algoritma clustering lain guna membandingkan performa serta meningkatkan kualitas hasil yang diperoleh. Selain itu, pemanfaatan dataset yang lebih luas dan beragam juga diharapkan mampu menghasilkan analisis yang lebih komprehensif dan akurat.

DAFTAR PUSTAKA

- [1] Humas BAZNAS Jatim. (2024). Dorong Kesejahteraan Mustahik, BAZNAS RI Luncurkan Program ZAuto di Jawa Timur. <https://jatim.baznas.go.id/news-show/peluncuranzautojawatimur2024/10911?utm>
- [2] Saputra, E. A., & Nataliani, Y. (2021). Analisis Pengelompokan Data Nilai Siswa untuk Menentukan Siswa-Siswi Unggulan Menggunakan Metode K-Medoids Clustering. *Journal of Information Systems and Informatics*, 3(3), 426-435..
- [3] F. Handayani, “Aplikasi Data Mining Menggunakan Algoritma K-Medoids Clustering untuk Mengelompokkan Mahasiswa Berdasarkan Gaya Belajar,” *Jurnal Teknologi dan Informasi (JATI)*, vol. 12, no. 1, pp. 46–62, 2022.
- [4] Herman, E., Zsido, K.-E., & Fenyves, V. (2022). *Cluster Analysis with K-Means versus K-Medoid in Financial Performance Evaluation*. *Applied Sciences*, 12(16), 7985. DOI:10.3390/app12167985



- [5] S. Bahri and D. M. Midyanti, “Penerapan Metode K-Medoids untuk Pengelompokan Mahasiswa Berpotensi Drop Out,” *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 10, no. 1, pp. 165–172, Feb. 2023.
- [6] A. A. D. Sulistyawati and M. Sadikin, “Penerapan Algoritma K-Medoids untuk Menentukan Segmentasi Pelanggan,” *SISTEMASI: Jurnal Sistem Informasi*, vol. 10, no. 3, pp. 516–526, 2021.
- [7] Y. Diana and F. Hadi, “Analisa Penjualan Menggunakan Algoritma K-Medoids untuk Mengoptimalkan Penjualan Barang,” *JOISIE: Journal of Information System and Informatics Engineering*, vol. 7, no. 1, pp. 97–103, Jun. 2023.
- [8] A. Setiawan and R. Andriani, “Penerapan Algoritma K-Medoids untuk Clustering Data Transaksi Penjualan,” *Jurnal Media Informatika Budidarma*, vol. 6, no. 1, pp. 245–252, Jan. 2022.
- [9] Iqbal Tawakal, M. Makmun Effendi, dan Annisa Maulana Majid, “Analisis Tingkat Kemiskinan dengan Algoritma K-Medoids Menggunakan RapidMiner di Tingkat Kota Kabupaten di Jawa Tengah,” *Journal of Information System Management (JOISM)*, vol. 7, no. 1, hlm. 112–119, 2025.
- [10] D. A. I. C. Dewi and D. A. K. Pramita, “Analisis Perbandingan Metode Elbow dan Silhouette pada Algoritma Clustering K-Medoids dalam Pengelompokan Produksi Kerajinan Bali,” *Jurnal Matrix*, vol. 9, no. 3, pp. 102–109, Nov. 2019.
- [11] D. N. P. Sari dan Y. L. Sukestiyarno, “Analisis Cluster Metode K-Medoids pada Persebaran Kasus COVID-19 Berdasarkan Provinsi di Indonesia,” *Prosiding PRISMA 4, Seminar Nasional Matematika, Universitas Negeri Semarang*, pp. 35–42, 2021.
- [12] M. Ardan, “Perbandingan K-Medoids dan Fuzzy C-Means pada Data Pendidikan Jakarta,” *Jurnal Komputer dan Informatika (JKI)*, vol. 22, no. 1, pp. 45–54, 2024.
- [13] W. M. P. Dhuhita, “Clustering Menggunakan Metode K-Means,” *Jurnal Informatika*, vol. 9, no. 1, pp. 15–21, 2015.
- [14] F. Alghifari dan D. Juardi, “Penerapan Data Mining Pada Penjualan Makanan Dan Minuman Menggunakan Metode Algoritma Naive Bayes,” 2021.
- [15] D.N. Yoliadi, “Data Mining Dalam Analisis Tingkat Penjualan Barang Elektronik Menggunakan Algoritma K-Medoids,” 2023.