

PENERAPAN ALGORITMA K-MEANS CLUSTERING DALAM MENENTUKAN DAERAH PRIORITAS PENANGANAN KEMISKINAN DI WILAYAH JAWA TIMUR

Nanda Sukarno Wijaya, Mohamad Jajuli, Budi Arif Dermawan

Informatika, Universitas Singaperbangsa Karawang

Jalan HS. Ronggo Waluyo Karawang, Indonesia

2010631170154@student.unsika.ac.id

ABSTRAK

Kemiskinan menjadi persoalan yang serius dengan membutuhkan penanganan baik di tingkat nasional maupun global. Di Indonesia terjadi peningkatan serupa sejak september 2022 dengan peningkatan 0,03% dan jumlah masyarakat dinyatakan miskin sebesar 26,36 juta jiwa. Jawa Timur merupakan salah satu wilayah dengan kemiskinan tertinggi di Indonesia. Untuk mendukung upaya penyelesaian kasus yang lebih terstruktur, diperlukan pengendalian dan penekanan kasus kemiskinan dengan menentukan daerah prioritas. Penelitian menggunakan data dari periode tahun 2021-2023. Penelitian ini menggunakan algoritma K-Means dengan metodologi *Knowledge Discovery in Database (KDD)* yaitu *data selection*, *data preprocessing*, *transformation*, *data mining*, dan *interpretation/evaluation*. Algoritma K-Means merupakan sistem untuk pengelompokan data non hierarki yang dapat mempartisi data terdiri dari dua kelompok atau lebih. Penelitian ini bertujuan melakukan clustering menggunakan K-Means dengan bantuan *Silhouette Coefficient*, serta dilakukan visualisasi menggunakan *tools Quantum Geographic Information System (QGIS)* dalam menentukan segmentasi wilayah kemiskinan di Jawa Timur. Berdasarkan data persentase jumlah penduduk miskin dan tingkat pengangguran terbuka. Penelitian ini menghasilkan 2 cluster dengan hasil clustering menggunakan *Davies Bouldin Index* sebesar 0.600 serta score *Silhouette* sebesar 0.550. Hasil visualisasi melalui QGIS menunjukkan cluster 0 masuk ke kategori rendah, dan cluster 1 masuk kategori tinggi.

Kata kunci : Clustering, Jawa Timur, Kemiskinan, K-Means

1. PENDAHULUAN

Kemiskinan menjadi persoalan yang serius dengan membutuhkan penanganan baik di tingkat nasional maupun global, dan termasuk dalam bagian dari agenda tujuan pembangunan berkelanjutan dunia, dengan ini diberikan prioritas yang tinggi [1]. Pada negara Indonesia terjadi peningkatan kemiskinan 0,03 persen dengan sebesar 26,36 juta jiwa pada tahun 2022. Berdasarkan data BPS pulau Jawa menduduki peringkat pertama dalam kasus kemiskinan terbesar di Indonesia sebesar 8,03 di perkotaan dan 5,91 di pedesaan [2].

Provinsi Jawa Timur merupakan salah satu provinsi terbanyak kasus kemiskinan di Indonesia dengan jumlah 1,75 juta penduduk perkotaan dan 2,48 juta penduduk pedesaan. Lalu kemudian diikuti oleh Jawa Barat, Jawa Tengah, dan kota lainnya [3]. Berdasarkan data diatas tingginya angka kemiskinan yang terus menerus menunjukkan bahwa masalah ini belum diatasi secara efektif. Hal ini merupakan meningkatkan persentase kemiskinan dari tahun 2021 ke tahun 2023.

Oleh karena itu, klasterisasi atau *clustering* menjadi bagian penting dari penanganan solusi yang dihadapi oleh provinsi Jawa Timur. Pada teknik *clustering* dilakukan pengelompokan sejumlah data ke dalam *cluster* berdasarkan tingkat kesesuaian antar data pada *cluster* tersebut dan perbedaan pada cluster lainnya [4].

Berdasarkan masalah tersebut, penelitian ini akan melakukan klasterisasi untuk menentukan daerah prioritas dalam upaya mengurangi kemiskinan di

Kab/kota di Jawa Timur berdasarkan kasus kemiskinan. Dengan menganalisis hasil segmentasi, diharapkan pemerintah Jawa Timur dan lembaga-lembaga yang berorientasi pada kesejahteraan masyarakat dapat menetapkan langkah yang tepat untuk menurunkan tingkat kemiskinan di wilayah tersebut. Algoritma yang digunakan adalah K-Means dengan bantuan *Silhouette Coefficient* dan divisualisasi menggunakan *Quantum Geographic Information System (QGIS)*.

Ada beberapa sumber dan jurnal bahwa penelitian sebelumnya menunjukkan kesuksesan sejumlah peneliti dalam penerapan metode algoritma K-Means clustering kemiskinan di Jawa Timur:

- K-Means clustering* dengan metode elbow untuk pengelompokan kabupaten dan kota di Jawa Timur berdasarkan indikator kemiskinan. Berdasarkan hasil dari penelitian tersebut menggunakan metode K-Means cluster menghasilkan 4 cluster: cluster 1 terdiri dari 10 kabupaten/kota, cluster 2 terdiri dari 20 kabupaten/kota, cluster 3 terdiri dari 4 kabupaten dan kota, dan cluster 4 terdiri dari 4 kabupaten.
- Perbandingan K-Means, Fuzzy C-Means, Fuzzy Gustafson, dan Density-Based Spatial Clustering of Applications with noise untuk pengelompokan desa di Surabaya berdasarkan indikator kemiskinan. Hasil perbandingan metode K-Means adalah terbaik untuk pengelompokan desa di Surabaya berdasarkan nilai total kuadrat dalam cluster. Sementara itu, metode Density-Based

Spatial Clustering of applications with noise menjadi metode terbaik karena mendekati angka 1 dibandingkan dengan metode lainnya.

2. TINJAUAN PUSTAKA

2.1. Kemiskinan

Kemiskinan adalah situasi di mana pendapatan tidak memenuhi kebutuhan dasar dan memastikan keberlangsungan hidup serta keterbatasan hidup yang serba kekurangan sehingga tidak dapat mencukupi kebutuhan dasar atau standar kelayakan hidup [5].

2.2. Data Mining

Data mining adalah sistem pengambilan informasi dari kumpulan data yang besar dan kompleks dalam sebuah database. Mekanisme yang digunakan adalah mengekstraksi data yang sudah ada diubah kedalam bentuk pengetahuan yang berfungsi untuk pengambilan keputusan [6]. *Data mining* dapat diartikan sebagai sistem penggalian informasi yang bermanfaat untuk database yang sangat besar. *Data mining* dikenal luas dalam bagian ilmu komputer sebagai metode untuk menentukan pola tersembunyi untuk menciptakan informasi yang baru dari sekumpulan data [7].

2.3. Knowledge Discovery in Database (KDD)

Knowledge Discovery in Database (KDD) ialah menghasilkan sistem *data mining* (sistem pengestrak kecenderungan sebuah pola data), selanjutnya diubah menjadi hasil yang tepat digunakan untuk menghasilkan informasi yang mudah dimengerti [8].

2.4. Clustering

Metode *clustering* merupakan pembagian pengelompokan data ke dalam kategori yang sama. Meminimalisir variasi fungsi tujuan yang diterapkan dalam proses *cluster* merupakan tujuan dari *clustering* [9]. Klasterisasi digunakan untuk mengorganisir data masuk dalam kelompok atau *cluster* dengan mempertimbangkan kesamaan fitur atau karakteristik. Algoritma *K-Means* dan *hierarchical clustering* menjadi contoh teknik klastering [10].

2.5. K-Means

K-Means ialah bentuk mekanisme *data mining* yang berguna untuk mengelompokkan data tanpa struktur dan hanya mempunyai tingkatan atau cenderung partisi yang terbagi kedalam kelompok-kelompok. Metode ini mampu memisahkan data ke dalam satu kelompok yang memiliki kesamaan dan memisahkan dengan kelompok lain. Teknik itu berlandaskan jarak antar data yang terdapat dalam *cluster* yang telah terbentuk, di mana semakin kecil jaraknya menunjukkan tingkat kemiripan antara data satu dengan data lainnya [11].

Tahapan algoritma *K-Means* mencakup beberapa hal:

- Menetapkan jumlah cluster (k).
- Menetapkan pusat cluster atau biasa disebut centroid.
- Mengkakulasi jarak pada setiap data terhadap masing masing pusat.

Untuk menghitung jarak rumus yang diterapkan adalah rumus Euclidean (Euclidean Distance) dengan persamaan yaitu:

$$d(a, b) = \sqrt{\sum_{k=1}^n (x_{ak} - x_{bk})^2}$$

Keterangan:

$d(a,b)$: Disatance (jarak) antara objek a dan b

n : Jumlah atribut

x_{ak} : Nilai pusat dari objek a pada dimensi k

x_{bk} : Nilai pusat dari objek b pada dimensi k

- Mengelompokkan data yang didasarkan pada jarak terdekat ke titik pusat.
- Menetapkan nilai pusat yang baru dengan mengkalkulasi rata-rata dari cluster.

2.6. Silhouette Coefficient

Silhouette Coefficient merupakan metode visualisasi yang diaplikasikan untuk mengevaluasi seberapa baik pengelompokan (*clustering*) [12]. Kualitas pengelompokan bisa dilakukan dengan menggunakan rata rata siluet. Dimana peningkatan nilai indeks ini dapat digunakan sebagai petunjuk untuk menentukan jumlah klaster yang optimal [1].

2.7. Davies Bouldin Index

Davies Boludin Index ialah sistem yang biasa digunakan untuk menilai hasil *clustering* yang telah di buat. *Index Davies-Bouldin* menghasilkan proses evaluasi klaster internal. Dimana hasil *clustering* dianggap semakin baik jika DBI yang dihasilkan mendekati 0. Tujuan dari penggunaan *Davies-Bouldin Indeks* adalah guna memperpendek jarak antara titik-titik dalam klaster sambil memperbesar jarak antara klaster [13].

2.8. Pemograman Python

Python merupakan sebuah bahasa pemograman interpetatif dan serbaguna yang difokuskan pada kemudahan kode. *Python* dianggap seperti bahasa yang mengkombinasikan kemampuan, dengan struktur kode yang sangat mudah dipahami, kapabilitas, dan menyediakan beragam fungsionalitas *library* [14].

2.9. Google Colab

Google Colab adalah *platform cloud computing* gratis yang dipersembahkan oleh google. *Google Colab* memungkinkan pengguna untuk menjalankan pemrosesan data dan *machine learning* [15]. *Library* populer seperti TensorFlow, Pytorch, Scikit-learn, dan Numpy sudah tersedia pada *platform google colab*. Ketersedian akses memberikan fleksibilitas bagi pengguna untuk memanfaatkan *library* sesuai kebutuhan [16].

2.10. Quantum Geographic Information System

Sistem Informasi Geografis (SIG) ialah suatu teknik informasi berbasis spasial yang dapat digunakan sebagai *tools* pembelajaran geografis untuk mengenali serta memahami kondisi geografis [17]. Perangkat lunak yang dapat dimanfaatkan untuk membuat peta geografis, analisis spasial, dan memvisualisasikan data geografis [18].

(1)

3. METODE PENELITIAN

Penelitian ini dilakukan dengan metodologi *Knowledge Discovery in Databases* (KDD). *Knowledge Discovery in Databases* adalah langkah-langkah dalam sistem penentu data menjadi kelompok informasi yang mencerminkan yang mencerminkan pola. Informasi ada dalam basis data yang besar dan manfaat yang potensial dari data tersebut masih belum diketahui. Penelitian ini akan dilakukan dengan menggunakan tahapan *data selection*, *data preprocessing*, *transformation*, *data mining*, dan *interpretation/evaluation*.



Gambar 1. Rancangan Penelitian

4. HASIL DAN PEMBAHASAN

Pada tahap *data selction* dilakukan pemilihan data yang relevan yang berada pada dataset, tidak semua data digunakan. *Dataset* yang digunakan adalah data publik yang diambil dari *website* jatim.bps.go.id yang telah dikumpulkan oleh pemerintah Jawa Timur.

Tabel 1. Dataset

Kab/Kota	Penduduk Miskin			Tingkat Pengangguran		
	2021	2022	2023	2021	2022	2023
Pacitan	84,19	76,2	76,2	2,04	3,65	1,83
Ponorogo	89,94	81,8	83,71	4,38	5,51	4,66
Blitar	112,6	101,9	101,9	3,66	5,45	4,91
Kediri	184,4	169,4	171,1	5,15	6,83	5,79
Malang	276,5	252,8	251,3	5,4	6,57	5,7

Berdasarkan Tabel 1 data tersebut merupakan dataset kemiskinan di Jawa Timur berdasarkan persentase penduduk miskin dan tingkat pengangguran pada rentang waktu 2021-2023 dengan mempunyai 10 atribut awal yaitu kode provinsi, provinsi, kota kabupaten kota, kab/kota, penduduk miskin 2021, penduduk miskin 2022, penduduk miskin 2023, TPT 2021, TPT 2022, dan TPT 2023 serta adanya penambahan variabel 2 variabel baru. Data yang dianggap tidak relevan atau tidak terpengaruhi proses analisis akan dihapus. Atribut yang di *selection* yaitu atribut kode provinsi, provinsi, kab/kota dan kode kabupaten kota.

Pada tahap sering disebut juga dengan *data cleaning* atau pembersihan data. Langkah awal yaitu pengecekan *missing value*.

No	Kode_Provinsi	Provinsi	Kode_Kabupaten_Kota	Kab/Kota	Penduduk Miskin_2021	Penduduk Miskin_2022	Penduduk Miskin_2023	TPT_2021	TPT_2022	TPT_2023	Total penduduk miskin	Total TPT
0	False	False	False	False	False	False	False	False	False	False	False	False
1	False	False	False	False	False	False	False	False	False	False	False	False
2	False	False	False	False	False	False	False	False	False	False	False	False
3	False	False	False	False	False	False	False	False	False	False	False	False
4	False	False	False	False	False	False	False	False	False	False	False	False
5	False	False	False	False	False	False	False	False	False	False	False	False
6	False	False	False	False	False	False	False	False	False	False	False	False
7	False	False	False	False	False	False	False	False	False	False	False	False
8	False	False	False	False	False	False	False	False	False	False	False	False
9	False	False	False	False	False	False	False	False	False	False	False	False
10	False	False	False	False	False	False	False	False	False	False	False	False
11	False	False	False	False	False	False	False	False	False	False	False	False
12	False	False	False	False	False	False	False	False	False	False	False	False
13	False	False	False	False	False	False	False	False	False	False	False	False
14	False	False	False	False	False	False	False	False	False	False	False	False

Gambar 2. Missing Value

Pada gambar 2, tidak ada nilai TRUE yang menunjukkan bahwa tidak ada data dengan nilai kosong. Oleh karena itu, langkah selanjutnya adalah melakukan pengecekan terhadap data yang kemungkinan memiliki duplikat atau redundansi.

Setelah melakukan pengecekan *missing value* dan duplikat data. Tahapan berikutnya yaitu proses transformasi data. Pada tahap ini, data yang kompleks diubah menjadi format yang lebih sederhana untuk pengolahan lebih lanjut seperti mengurutkan data, standarisasi data atau dengan melakukan simbolisasi. Transformasi data membantu dalam mengubah format yang lebih bermanfaat atau sesuai untuk analisis atau presentasi. Standarisasi data dilakukan untuk memastikan bahwa variabel yang memiliki skala yang berbeda memiliki rentang nilai yang seragam. Hal ini bertujuan untuk menghindari ketidakseimbangan yang dapat mengakibatkan model klasterisasi menjadi kurang optimal.

Setelah dilakukan transformasi data dengan proses normalisasi. Langkah berikutnya ialah teknik data mining untuk mengidentifikasi pola yang terdapat

dalam data. Teknik yang digunakan dalam tahap ini adalah metode clustering yang bertujuan untuk pengelompokan atau menentukan daerah prioritas data dikelompokkan berdasarkan tingkat kesamaan dalam cluster. Algoritma yang diterapkan ialah algoritma K-Means dengan optimalisasi penentuan jumlah cluster melakukan grafik silhouette digunakan sebagai pertimbangan penentuan jumlah cluster.

Metode Silhouette ialah teknik yang digunakan untuk menilai kualitas klaster dalam analisis klaster. Teknik ini memberikan ukuran seberapa baik setiap titik data berada dalam klasternya dibandingkan dengan klaster lain. Nilai Silhouette memberikan wawasan tentang mirip objek dalam klaster yang sama dibandingkan dengan objek di klaster lain.

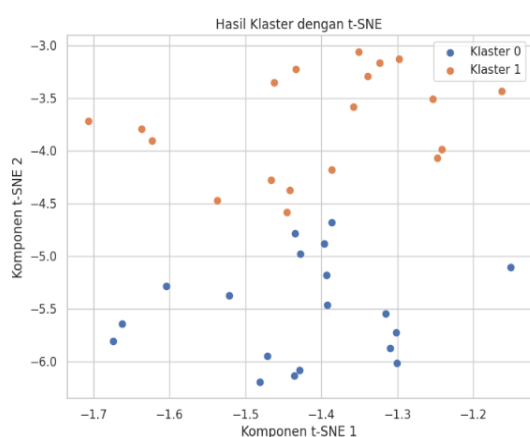


Gambar 3. Grafik Metode Silhouette

Dalam penentuan yang telah didapatkan cluster dianggap memiliki kualitas yang baik saat mendekati nilai maksimum, ialah 1, dalam skala nilai dari -1 hingga 1. Oleh karena itu, berdasarkan grafik Silhouette Coefficient menunjukan titik optimal pada C=2.

Tabel 2. Hasil Indeks Silhouette

Cluster	Indeks Silhouette
2	0.550
3	0.489
4	0.453
5	0.432



Gambar 4. Grafik K-Means Clustering

Nilai indeks Silhouette mengukur sejauh mana objek dalam cluster mendekati cluster mereka sendiri

dibandingkan dengan cluster lainnya. Nilai indeks silhouette berkisar antara -1 hingga 1 pada tabel 1, dimana nilai positif yang lebih tinggi menunjukkan segmentasi yang lebih baik. Berdasarkan tabel ini, cluster dengan angka 2 memiliki indeks silhouette tertinggi 0.550.

Setelah melalui proses tahapan data mining dan telah dilakukan clustering. Mendapatkan hasil cluster terbaik C=2 dan juga pada grafik silhouette coefficient dimana jumlah cluster optimal pertama dengan nilai C=2. Maka dilakukan evaluasi menggunakan nilai Davies Bouldin Index.

Tabel 3. Hasil Indeks Clustering

Jumlah Cluster	DBI
2	0.600
3	0.636
4	0.679
5	0.719

Tabel 4. Karakteristik Cluster

Cluster	Penduduk Miskin			Tingkat Pengangguran		
	2021	2022	2023	2021	2022	2023
C0	1.00	2.12	2.10	2.11	-2.45	2.01
C1	2.90	-2.12	-2.10	2.45	-2.01	-8.77

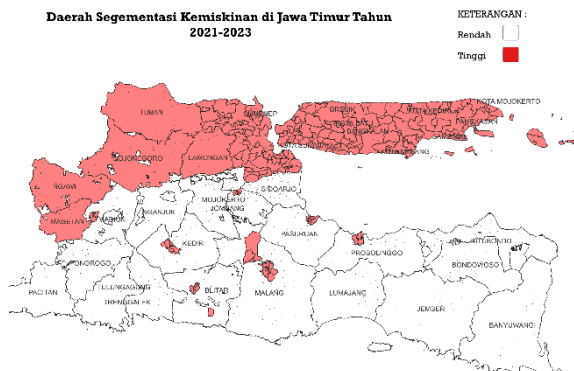
Berdasarkan Davies Bouldin Index (DBI) pengelompokan data dianggap optimal atau semakin baik jika nilai index yang lebih rendah atau mendekati 0. Berdasarkan tabel 2, titik optimal cluster berdasarkan nilai DBI 0.600 berada di C=2. Dan pada tabel 3 proses klasterisasi yang telah dilakukan juga pelabelan pada cluster berdasarkan karakteristik masing-masing cluster yang kemudian menghasilkan 2 label yaitu daerah kemiskinan rendah dan tinggi.

Hasil menunjukan bahwa jumlah cluster yang optimal dengan menggunakan DBI yang menghasilkan kesimpulan bahwa cluster optimal pada skenario terletak pada titik C=2. Nilai 1 walau masih tetap masuk kedalam struktur baik. dalam penyajian akhir terdapat 2 cluster yang diberikan label berdasarkan analisis karakteristik dari setiap cluster. Cluster 0 dikelompokkan sebagai daerah dengan tingkat kemiskinan rendah, sementara cluster 1 merupakan daerah dengan tingkat kemiskinan tinggi. Pada tabel 4 hasil clustering bisa dilihat Tabel dan visualisasi peta segmentasi daerah kemiskinan di Jawa Timur dilihat pada Gambar 5.

Tabel 5. Hasil Clustering

C0	C1
Pacitan	Magetan
Ponorogo	Ngawi
Trenggalek	Bojonegoro
Tulungagung	Tuban
Blitar	Lamongan
Kediri	Gresik
Malang	Bangkalan
Lumajang	Sampang
Jember	Pamekasan

C0	C1
Banyuwangi	Sumenep
Bondowoso	Kota Kediri
Situbondo	Kota Blitar
Probolinggo	Kota Malang
Pasuruan	Kota Probolinggo
Sidoarjo	Kota Pasuruan
Mojokerto	Kota Mojokerto
Jombang	Kota Madiun
Nganjuk	Kota Surabaya
Madiun	Kota Batu



Gambar 5. Peta segmentasi daerah kemiskinan di Jawa Timur

5. KESIMPULAN DAN SARAN

Hasil penelitian ini, dapat diperoleh kesimpulan bahwa pada peta segmentasi daerah berdasarkan 38 kab/kota di Jawa Timur sebanyak 19 kab/kota ke dalam daerah kemiskinan pada cluster 0, dan 19 kab/kota masuk kedalam kemiskinan yang tinggi pada cluster 1. Hasil clustering dievaluasi menggunakan nilai *Davies Bouldin Index* (DBI) dengan score 0.600. Score Silhouette dengan nilai 0.550 menunjukan bahwa pola cluster pada data memiliki struktur yang baik artinya data sudah berada di cluster yang tepat. Sebagai rekomendasi untuk pengembangan penelitian, disarankan untuk menggunakan algoritma dan alat analisis yang berbeda serta mempertimbangkan penggunaan variabel yang lebih banyak.

DAFTAR PUSTAKA

- [1] I. Manfaati Nur, M. Rizky, and S. Putri Milasari, "Pengelompokan Tingkat Kemiskinan di Provinsi Jawa Barat dengan Metode K-Means Clustering Info Artikel," vol. 1, no. 2, pp. 51–61, 2023, [Online]. Available: <http://journalnew.unimus.ac.id/index.php/jodi>
- [2] M. Orisa, "Optimasi Cluster pada Algoritma K-Means," *Pros. SENIATI*, vol. 6, no. 2, pp. 430–437, 2022, doi: 10.36040/seniati.v6i2.5034.
- [3] A. Ahdiyat, "10 Provinsi dengan Penduduk Miskin Terbanyak September 2022." [Online]. Available: <https://databoks.katadata.co.id/datapublish/2023/01/17/10-provinsi-dengan-penduduk-miskin-terbanyak-september-2022#:~:text=Provinsi dengan penduduk miskin terbanyak berikutnya>

adalah Jawa Barat%2C Jawa,rincian seperti terlihat pada grafik.

- [4] I. N. M. Adiputra, "Clustering Penyakit DBD Pada Rumah Sakit Dharma Kerti Menggunakan Algoritma K-Means," *Inser. Inf. Syst. Emerg. Technol. J.*, vol. 2, no. 2, p. 99, 2022, doi: 10.23887/insert.v2i2.41673.
- [5] J. Suntoro, *Data Mining: Algoritma dan Implementasi dengan pemrograman PHP*. Jakarta: Elex Media Komputido, 2019.
- [6] J. Dongga, A. Sarungallo, N. Koru, and G. Lante, "Implementasi Data Mining Menggunakan Algoritma Apriori Dalam Menentukan Persediaan Barang (Studi Kasus: Toko Swapen Jaya Manokwari)," *G-Tech J. Teknol. Terap.*, vol. 7, no. 1, pp. 119–126, 2023, doi: 10.33379/gtech.v7i1.1938.
- [7] N. I. Febianto and N. Palasara, "Analisa Clustering K-Means Pada Data Informasi Kemiskinan Di Jawa Barat Tahun 2018," *J. Sisfokom (Sistem Inf. dan Komputer)*, vol. 8, no. 2, pp. 130–140, 2019, doi: 10.32736/sisfokom.v8i2.653.
- [8] S. Z. Harahap and A. Nastuti, "Teknik Data Mining Untuk Penentuan Paket Hemat Sembako," *J. Ilm. Fak. Sains dan Teknol.*, vol. 7, no. 3, pp. 111–119, 2019.
- [9] A. Ali, "Clustering Data Antropometri Balita Untuk Menentukan Status Gizi Balita Di Kelurahan Jumpat Rejo Sukodono Sidoarjo," *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 7, no. 3, pp. 395–407, 2020, doi: 10.35957/jatisi.v7i3.530.
- [10] I. G. I. Sudipa, I. G. M. Darmawiguna, I. M. Dendi, and M. Sanjaya, *Buku ajar data mining*, no. January. 2024.
- [11] I. Nuryani and D. Darwis, "Analisis Clustering Pada Pengguna Brand Hp Menggunakan Metode K-Means," *Pros. Semin. Nas. Ilmu Komput.*, vol. 1, no. 1, p. 2021, 2021.
- [12] X. Li, Di. Yan, L. Dong, and R. Wang, "Anti-Forensics of Audio Source Identification Using Generative Adversarial Network," *IEEE Access*, vol. 7, pp. 184332–184339, 2019, doi: 10.1109/ACCESS.2019.2960097.
- [13] J. Ilmiah and W. Pendidikan, "DOI: <https://doi.org/10.5281/zenodo.10081332> p-ISSN: e-ISSN: Accredited by Directorate General of Strengthening for Research and Development Available online at <https://jurnal.peneliti.net/index.php/JIWP>," vol. 9, no. November, pp. 547–558, 2023.
- [14] D. A. Manalu and G. Gunadi, "Implementasi Metode Data Mining K-Means Clustering Terhadap Data Pembayaran Transaksi Menggunakan Bahasa Pemrograman Python Pada Cv Digital Dimensi," *Infotech J. Technol. Inf.*, vol. 8, no. 1, pp. 43–54, 2022, doi: 10.37365/jti.v8i1.131.
- [15] C. Menggunakan and G. Colab, "Manfaat Google

- Colab Aplikasi Google Colab: Cara Install Library Python dan Upload File”.
- [16] M. C. Sitar and R. J. Leary, “Technical note: Colab_zirc_dims: A Google Colab-compatible toolset for automated and semi-Automated measurement of mineral grains in laser ablation-inductively coupled plasma-mass spectrometry images using deep learning models,” *Geochronology*, vol. 5, no. 1, pp. 109–126, 2023, doi: 10.5194/gchron-5-109-2023.
- [17] Nurfitri Andayani, Wimmy Hartawan, and A. Maulana, “Perancangan Sistem Pemetaan Wilayah Calon Pelanggan Dengan Menggunakan Qgis Pada Pt. Indonesia Comnets Plus (Icon+) Sbu Bengkulu,” *J. Inform.*, vol. 1, no. 2, pp. 1–12, 2022, doi: 10.57094/ji.v1i2.357.
- [18] A. Wicaksono and *, Zainul Hidayah, “Pemanfaatan Sistem Informasi Geografis Berbasis Web Dalam Meningkatkan Akurasi Informasi Terkait Rekam Jejak Sumur Minyak Dan Gas Bumi Di Pulau Madura,” *JST (Jurnal Sains dan Teknol.*, vol. 11, no. 2, pp. 362–370, 2022, doi: 10.23887/jstundiksha.v11i2.43553.