

PENERAPAN ALGORITMA K-MEANS UNTUK MENGANALISIS PENGANGGURAN BERDASARKAN FAKTOR PENDIDIKAN DI JAWA BARAT

Ahmad Wahyudiana, Baenil Huda, Elfina Novalia, Tukino

Sistem Informasi, Fakultas Ilmu Komputer, Universitas Buana Perjuangan Karawang
Jalan Ronggo Waluyo Sirnabaya, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361
Si22.ahmadwahyudiana@mhs.ubpkarawang.ac.id

ABSTRAK

Tingkat pengangguran merupakan salah satu indikator penting dalam menilai kondisi ketenagakerjaan suatu wilayah. Di Provinsi Jawa Barat, perbedaan tingkat pendidikan dapat memengaruhi variasi tingkat pengangguran pada setiap kabupaten/kota. Permasalahan yang muncul adalah bagaimana mengidentifikasi pola pengangguran berdasarkan jenjang pendidikan agar dapat memberikan gambaran kondisi ketenagakerjaan antarwilayah secara lebih terstruktur. Penelitian ini bertujuan untuk menganalisis pola pengangguran berdasarkan jenjang pendidikan di kabupaten/kota Provinsi Jawa Barat dengan memanfaatkan teknik *data mining* menggunakan algoritma *K-Means*. Data yang digunakan merupakan data sekunder periode 2011–2024 yang diperoleh dari Portal Open Data Jabar. Proses analisis dilakukan melalui tahapan pemilihan data, pra-pemrosesan, transformasi data, serta proses clustering menggunakan algoritma *K-Means*. Penentuan jumlah kluster optimal dilakukan dengan metode Elbow dan Silhouette Score. Hasil penelitian menunjukkan bahwa data pengangguran dapat dikelompokkan menjadi empat kluster dengan karakteristik yang berbeda pada setiap kelompok wilayah. Pengelompokan tersebut memberikan gambaran yang lebih jelas mengenai pola pengangguran berdasarkan faktor pendidikan sehingga dapat menjadi bahan pertimbangan dalam penyusunan kebijakan ketenagakerjaan yang lebih tepat sasaran.

Kata kunci: Pengangguran, Pendidikan, Data Mining, K-Means, Clustering

1. PENDAHULUAN

Pengangguran merupakan kondisi ketika seseorang atau sekelompok individu tidak memiliki pekerjaan maupun sumber penghasilan tetap [1]. Tingginya angka pengangguran masih menjadi tantangan serius bagi negara-negara berkembang, termasuk Indonesia, terutama dalam aspek ketenagakerjaan. Kondisi ini menunjukkan bahwa permasalahan penyediaan lapangan kerja belum sepenuhnya teratasi, sehingga tingkat pengangguran di Indonesia masih berada pada angka yang cukup signifikan dibandingkan dengan beberapa negara di kawasan sekitarnya [2], [3]. Penelitian ini mengambil fokus pada Provinsi Jawa Barat mengingat wilayah tersebut termasuk salah satu provinsi dengan tingkat pengangguran tertinggi di Indonesia [4].

Pendidikan menjadi variabel penting yang berkorelasi dengan kondisi pengangguran, di mana peningkatan jenjang pendidikan biasanya diikuti oleh peluang kerja yang lebih besar. Namun demikian, kondisi pengangguran juga dipengaruhi oleh faktor ekonomi dan kebijakan ketenagakerjaan. Selain itu, rendahnya kemampuan penduduk usia kerja dalam mengikuti perkembangan teknologi turut berkontribusi terhadap peningkatan angka pengangguran [5]. Dalam upaya mengatasi permasalahan pengangguran, diperlukan pengambilan keputusan yang didasarkan pada data dan bukti empiris yang kuat. Penerapan algoritma *K-Means* dapat membantu dalam menyediakan dasar analisis yang lebih sistematis untuk mendukung perumusan kebijakan serta perencanaan strategi penanganan pengangguran [6], [7].

Pemahaman pola pengangguran dapat dilakukan melalui teknik clustering menggunakan algoritma *K-Means*, yang termasuk dalam metode *unsupervised learning* dan efektif untuk pengelompokan data [8], [9]. Teknik *data mining* digunakan untuk menelaah kumpulan data berukuran besar guna mengidentifikasi pola, keterkaitan, serta informasi yang memiliki nilai analitis dengan memanfaatkan pendekatan dan algoritma tertentu. Proses ini menjadi bagian dari rangkaian Knowledge Discovery in Databases (KDD), yaitu suatu tahapan sistematis yang bertujuan memperoleh pengetahuan baru berdasarkan hasil analisis data yang tersedia [10], [11]. Hasil analisis ini memberikan pemetaan yang lebih sistematis mengenai karakteristik pengangguran berdasarkan jenjang pendidikan. Informasi tersebut dapat dimanfaatkan sebagai referensi dalam perumusan kebijakan, khususnya dalam pengembangan program peningkatan keterampilan dan perluasan lapangan kerja yang lebih sesuai dengan kondisi masing-masing wilayah, sehingga berkontribusi terhadap penguatan ekonomi daerah [12].

Penelitian ini menggunakan perangkat lunak *RapidMiner* dengan menerapkan algoritma *K-Means* sebagai metode analisis. *RapidMiner* dimanfaatkan sebagai platform analitik untuk mengakses, mengolah, dan menganalisis data secara sistematis [13]. Penelitian difokuskan pada kabupaten/kota dengan kondisi lapangan pekerjaan yang terbatas sehingga berdampak pada tingkat pengangguran. Penelitian ini memanfaatkan data sekunder yang diperoleh melalui Portal Open Data Jabar dengan cakupan periode 2011–2024.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Menurut Andryadi dan Nisa (2023), algoritma *K-Means* dapat dimanfaatkan untuk melakukan pengelompokan data tingkat pengangguran pada wilayah kabupaten/kota di Provinsi Jawa Barat. Penelitian tersebut menggunakan data pengangguran terbuka dari Portal Open Data Jabar periode 2007–2022 yang dianalisis melalui tahapan Knowledge Discovery in Databases (KDD). Hasil penelitian menunjukkan bahwa wilayah yang dianalisis dapat dikelompokkan ke dalam dua kluster utama, yaitu kelompok daerah dengan tingkat pengangguran tinggi dan kelompok daerah dengan tingkat pengangguran rendah, sehingga hasil pengelompokan tersebut dapat memberikan gambaran dalam mendukung pengambilan kebijakan terkait penanganan pengangguran [14]

2.2. Pengangguran

Permasalahan pengangguran memberikan dampak yang cukup besar terhadap stabilitas ekonomi suatu daerah. Dengan demikian, dibutuhkan pendekatan analitis yang lebih menyeluruh guna menguraikan determinan yang memengaruhi kondisi tersebut. Tingkat pendidikan menjadi salah satu variabel yang berhubungan langsung dengan peluang seseorang dalam memperoleh pekerjaan. Namun, mengidentifikasi pola keterkaitan antara pendidikan dan pengangguran pada data berukuran besar memerlukan pendekatan analitis, karena pengolahan secara manual berisiko menimbulkan kekeliruan dalam penafsiran akibat kompleksitas struktur data [15]

2.3. Data Mining

Data mining merupakan proses penggalian dan analisis data dalam jumlah besar untuk memperoleh informasi yang bernilai serta mengidentifikasi pola yang tersembunyi di dalamnya. Salah satu teknik utama dalam data mining adalah *clustering*, yaitu metode pengelompokan data tanpa menggunakan label atau contoh kelompok sebelumnya, sehingga data dikelompokkan berdasarkan kemiripan karakteristiknya [16], [17]. *Data mining* merupakan bagian dari proses yang dikenal sebagai Knowledge Discovery in Databases (KDD) [18]

2.4. Algoritma K-means Clustering

Algoritma *K-Means Clustering* merupakan teknik pengelompokan data yang banyak digunakan untuk memproses dataset dalam jumlah besar secara efisien, meskipun metode ini cenderung sensitif terhadap keberadaan data pencilan (*outlier*). Pengelompokan ini dilakukan dengan mempertimbangkan tingkat kemiripan antar data sehingga objek yang memiliki karakteristik serupa berada dalam satu kluster yang sama. Tujuan utama metode ini adalah membentuk kelompok dengan tingkat kesamaan yang tinggi di dalam kluster dan

perbedaan yang jelas antar kluster. Proses pengelompokan dilakukan tanpa memerlukan informasi atau label awal mengenai kategori data [19], [20]

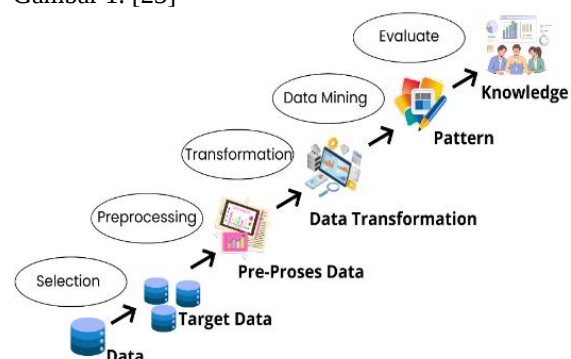
2.5. RapidMiner

Penelitian ini menggunakan perangkat lunak *RapidMiner* sebagai alat bantu dalam proses pengolahan dan analisis data. *RapidMiner* merupakan platform analisis dan penambangan data berbasis *Graphical User Interface* (GUI) yang memudahkan pengguna dalam melakukan tahapan data mining secara sistematis. Perangkat lunak ini bersifat *open source*, dibangun menggunakan bahasa pemrograman Java, serta mendukung berbagai teknik pemrosesan data, termasuk proses input, transformasi, dan analisis data [21], [22]

3. METODE PENELITIAN

3.1. Metode Penelitian

Pelaksanaan penelitian dilakukan secara bertahap mulai dari pengumpulan data hingga proses evaluasi akhir, yang meliputi seleksi data, pra-pemrosesan, transformasi, serta penerapan metode *data mining*. Rangkaian tahapan tersebut divisualisasikan dalam Gambar 1. [23]



Gambar 1. Metode Penelitian [23]

3.2. Pengumpulan Data

pendidikan pada kabupaten/kota di Provinsi Jawa Barat. Data yang digunakan merupakan data sekunder yang bersumber dari Portal Open Data Jabar sebagai platform resmi penyedia data pemerintah daerah. Proses pengumpulan dilakukan untuk memastikan ketersediaan dataset yang sesuai, akurat, dan relevan dengan kebutuhan analisis dalam penelitian ini.

3.3. Data Selection

Tahapan *data selection* dilakukan untuk menyaring variabel yang sesuai dengan fokus penelitian. Dataset yang digunakan berasal dari sumber sekunder, yaitu Portal Open Data Jabar, yang menyediakan informasi mengenai tingkat pengangguran berdasarkan jenjang pendidikan pada kabupaten/kota di Provinsi Jawa Barat untuk periode 2011–2024. Variabel yang digunakan dalam analisis mencakup nama wilayah, tahun pengamatan, kategori pendidikan, serta jumlah pengangguran.

3.4. Data Preprocessing

Tahap *data preprocessing* dilakukan untuk membersihkan dan menyiapkan dataset sebelum proses analisis. Pada tahap ini dilakukan pengecekan terhadap nilai yang tidak konsisten, penanganan data kosong (*missing values*), serta penyesuaian format data agar sesuai untuk proses komputasi. Proses ini bertujuan untuk memastikan bahwa data dalam kondisi bersih dan siap digunakan pada tahap selanjutnya.

3.5. Data Transformasi

Tahap *data transformasi* dilakukan untuk menyesuaikan format dan struktur dataset agar dapat diproses oleh algoritma *K-Means*. Pada proses ini, data yang awalnya berbentuk kategori tingkat pendidikan diubah menjadi representasi numerik dan disusun dalam bentuk tabel yang menampilkan jumlah pengangguran pada setiap jenjang pendidikan. Transformasi ini memungkinkan data dianalisis secara kuantitatif dalam proses *clustering*.

3.6. Data Mining

Pada tahap ini, proses analisis dilakukan dengan menggunakan algoritma *K-Means* untuk membentuk kelompok data berdasarkan kesamaan pola tingkat pengangguran dan jenjang pendidikan. Jumlah kluster ditentukan terlebih dahulu melalui proses evaluasi guna memperoleh konfigurasi yang paling sesuai. Hasil pengolahan menghasilkan beberapa kelompok wilayah dengan karakteristik pengangguran yang berbeda satu sama lain.

3.7. Evaluasi

Evaluasi dilakukan untuk menentukan jumlah kluster optimal dan mengukur kualitas hasil klusterisasi. Penelitian ini menggunakan *Elbow Method* untuk melihat perubahan nilai inertia serta *Silhouette Score* untuk mengukur tingkat pemisahan antar kluster. Berdasarkan hasil evaluasi, jumlah kluster sebanyak empat ($k = 4$) dinilai cukup optimal

dan konsisten untuk digunakan dalam proses klusterisasi.

3.8. Alat dan Bahan Penelitian

Dalam pelaksanaan penelitian ini digunakan perangkat lunak *RapidMiner* sebagai alat bantu analisis *clustering*. Dataset yang digunakan merupakan data sekunder tingkat pengangguran dan pendidikan kabupaten/kota di Provinsi Jawa Barat yang bersumber dari Portal Open Data Jabar. Proses komputasi dilakukan pada perangkat komputer dengan spesifikasi yang mendukung analisis data, meliputi prosesor Intel Core/AMD A8, RAM 4 GB, media penyimpanan 500 GB, serta sistem operasi Windows [25], [26].

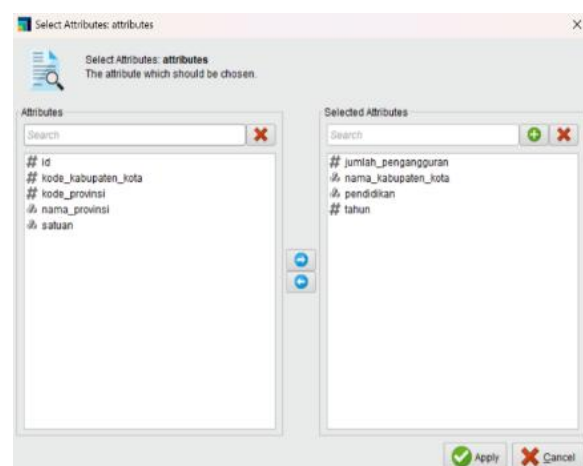
3.9. Metode Analisis Data

Penelitian ini menerapkan algoritma *K-Means* sebagai metode analisis untuk melakukan pengelompokan data. Pendekatan ini digunakan guna mengelompokkan kabupaten/kota di Provinsi Jawa Barat berdasarkan karakteristik tingkat pengangguran dan jenjang pendidikan. Proses pengelompokan dilakukan dengan menetapkan empat kluster sebagai jumlah kelompok yang digunakan, kemudian menghitung jarak setiap data terhadap pusat kluster (*centroid*) secara iteratif hingga terbentuk hasil pengelompokan yang konvergen.

4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Pada penelitian ini, proses *data selection* dilakukan menggunakan operator *Select Attributes* pada perangkat lunak *RapidMiner*. Melalui tahap ini, atribut-atribut yang bersifat administratif dan tidak berpengaruh terhadap proses klusterisasi dihapus dari dataset.



Gambar 5. Proses *Data Selection* menggunakan *Rapidminer*

Proses *data selection* dilakukan menggunakan operator *Select Attributes* pada *RapidMiner* untuk memilih atribut yang relevan dengan penelitian. Tahap ini menghasilkan dataset yang terfokus dan siap

digunakan pada tahap *data preprocessing* dan analisis selanjutnya.

4.2. Data Preprocessing

Tahap *data preprocessing* dilakukan untuk membersihkan dan menyiapkan dataset hasil *data selection* agar siap dianalisis. Pada tahap ini dilakukan penanganan nilai yang tidak konsisten serta pengecekan kelengkapan data, sehingga seluruh atribut berada dalam kondisi bersih dan siap digunakan pada proses klusterisasi.

RowNo	nama_kabupaten	pendidikan	jumlah_peserta	tahun
1	KABUPATEN...	SD KE BAWAH	80529	2011
2	KABUPATEN...	SLTP	36779	2011
3	KABUPATEN...	SLTA	9730	2011
4	KABUPATEN...	SD KE BAWAH	23242	2011
5	KABUPATEN...	SLTP	26191	2011
6	KABUPATEN...	SLTA	35451	2011
7	KABUPATEN...	SD KE BAWAH	59647	2011
8	KABUPATEN...	SLTP	27178	2011
9	KABUPATEN...	SLTA	19875	2011
10	KABUPATEN...	SD KE BAWAH	50490	2011
11	KABUPATEN...	SLTP	61548	2011
12	KABUPATEN...	SLTA	43167	2011
13	KABUPATEN...	SD KE BAWAH	23675	2011

Gambar 6. Hasil *Data Preprocessing* Menggunakan *RapidMiner*

4.3. Data Transformasi

Pada tahap ini, dataset yang telah melalui pra-pemrosesan disusun ulang sehingga strukturnya mendukung pelaksanaan analisis kluster. Pada tahap ini, data diubah dari bentuk baris menjadi kolom berdasarkan kategori jenjang pendidikan sehingga setiap atribut merepresentasikan jumlah pengangguran pada masing-masing jenjang pendidikan. Selain itu, nilai kosong yang muncul akibat proses transformasi ditangani dengan menggantinya menggunakan nilai nol. Dataset hasil data transformasi telah berada dalam bentuk numerik dan terstruktur, sehingga siap digunakan pada proses penentuan jumlah kluster dan klusterisasi menggunakan algoritma *K-Means*.

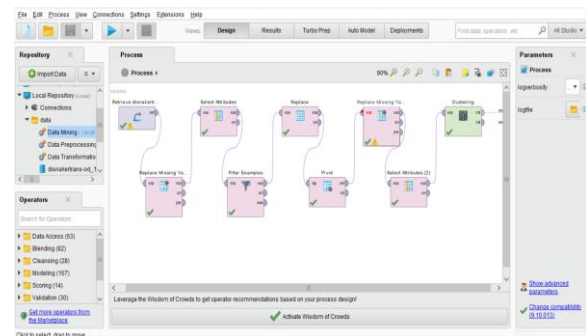
RowNo	nama_kabupaten	tahun	sd_ke_bawah	sltp	slta	sd_ke_bawah	sltp	slta	sd_ke_bawah	sltp	slta	sd_ke_bawah
1	KABUPATEN...	2011	0	0	0	80529	36779	9730	0	0	0	0
2	KABUPATEN...	2012	2650	60468	0	0	0	16209	0	0	0	57873
3	KABUPATEN...	2013	3408	61777	0	0	0	21284	0	0	0	42651
4	KABUPATEN...	2014	9119	46395	0	0	0	76136	0	0	0	49540
5	KABUPATEN...	2015	5986	23953	0	0	0	12019	0	0	0	39703
6	KABUPATEN...	2017	32193	34128	0	0	0	13468	0	0	0	47728
7	KABUPATEN...	2018	20942	42098	0	0	0	13672	0	0	0	43485
8	KABUPATEN...	2019	21991	0	0	0	0	129723	0	0	0	39554
9	KABUPATEN...	2020	39977	67122	0	0	0	119629	75798	72817	0	0
10	KABUPATEN...	2021	17435	0	0	0	0	160864	70395	62464	0	0
11	KABUPATEN...	2022	11189	0	0	0	0	68724	61156	61183	0	0
12	KABUPATEN...	2023	16913	0	0	0	0	67288	65788	39642	0	0
13	KABUPATEN...	2024	11115	0	0	0	0	71668	61616	51476	0	0

Gambar 7. Hasil *Data Transformasi* Menggunakan *RapidMiner*

4.4. Data Mining

Proses analisis utama dilakukan melalui penerapan metode *K-Means* guna mengidentifikasi pengelompokan wilayah di Provinsi Jawa Barat berdasarkan variasi angka pengangguran dan tingkat Pendidikan. Proses ini dilakukan untuk

mengidentifikasi pola serta kemiripan karakteristik antarwilayah dengan memanfaatkan data yang sebelumnya telah melalui *data selection*, *data Preprocessing*, dan *data transformasi*.



Gambar 8. Proses *Data Mining* Menggunakan *Rapidminer*

Berdasarkan hasil penerapan algoritma *K-Means*, dataset yang dianalisis tersegmentasi menjadi empat kluster dengan karakteristik yang berbeda.. Berdasarkan hasil klusterisasi, kluster 0 memiliki jumlah anggota terbanyak dengan 1.477 data, diikuti oleh kluster 1 dengan 59 data, kluster 3 dengan 22 data, dan kluster 2 dengan jumlah anggota paling sedikit yaitu 8 data. Secara keseluruhan, jumlah data yang dianalisis pada proses klusterisasi ini adalah sebanyak 1.566 data. Perbedaan jumlah anggota pada setiap kluster menunjukkan adanya variasi karakteristik data antar kluster yang terbentuk, yang selanjutnya dianalisis pada tahap interpretasi hasil.

Cluster Model

```
Cluster 0: 1477 items
Cluster 1: 59 items
Cluster 2: 8 items
Cluster 3: 22 items
Total number of items: 1566
```

Gambar 9. Hasil *Clustering K-Means*

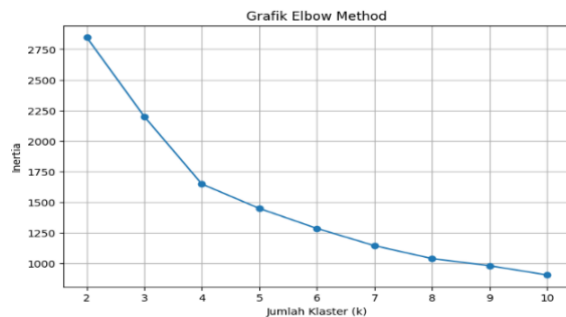
4.5. Evaluasi dan Penentuan Jumlah Kluster

Penentuan jumlah kluster dilakukan sebagai langkah untuk memastikan bahwa pembagian data menggunakan algoritma *K-Means* telah menghasilkan struktur kelompok yang paling tepat. Dalam penelitian ini, proses penilaian tersebut dilakukan dengan memanfaatkan dua metode evaluasi, yaitu *Elbow Method* dan *Silhouette Score*.

4.6. Elbow Method

Penentuan jumlah kluster dilakukan dengan memanfaatkan metode *Elbow*, yaitu dengan mengamati pola penurunan nilai inertia pada setiap variasi jumlah kluster. Inertia menunjukkan seberapa jauh jarak data terhadap pusat klasternya. Semakin rendah nilai inertia yang diperoleh, semakin kecil penyebaran data dalam kluster tersebut, sehingga

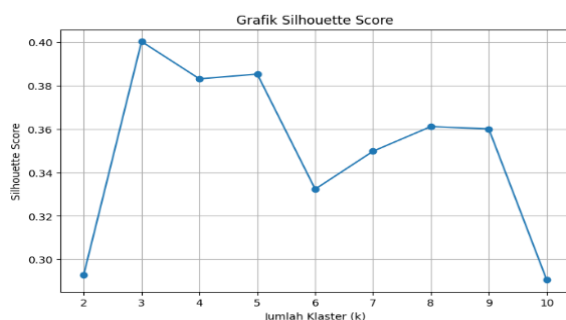
struktur pengelompokan menjadi lebih terpusat dan seragam.



Gambar 10. Grafik *Elbow Method*

4.7. Silhouette Score

Silhouette Score dimanfaatkan untuk mengukur seberapa baik hasil klusterisasi yang diperoleh dengan menilai tingkat kedekatan setiap data terhadap kelompoknya. Nilai yang dihasilkan berada pada kisaran -1 sampai 1 . Semakin mendekati angka 1 , semakin menunjukkan bahwa data tersebut memiliki kemiripan yang kuat dengan kluster tempatnya tergabung dan memiliki jarak yang cukup jelas dibandingkan dengan kluster lain.



Gambar 11. Grafik *Silhouette Score*

Berdasarkan hasil *clustering* dan evaluasi yang telah dilakukan, masing-masing kluster menunjukkan karakteristik tingkat pengangguran yang berbeda. Perbedaan tersebut mencerminkan variasi kondisi ketenagakerjaan antar wilayah di Provinsi Jawa Barat berdasarkan faktor pendidikan. Hasil ini menunjukkan bahwa algoritma *K-Means* mampu mengidentifikasi pola pengelompokan wilayah secara sistematis.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil analisis terhadap data tingkat pengangguran kabupaten/kota di Provinsi Jawa Barat periode 2011–2024, penerapan algoritma *K-Means* menghasilkan empat kluster dari total 1.566 data. Distribusi data menunjukkan bahwa Kluster 0 terdiri dari 1.477 data, Kluster 1 sebanyak 59 data, Kluster 2 sebanyak 8 data, dan Kluster 3 sebanyak 22 data, yang mencerminkan perbedaan karakteristik tingkat pengangguran berdasarkan jenjang pendidikan antarwilayah. Hasil evaluasi menggunakan *Elbow Method* dan *Silhouette Score* menunjukkan bahwa

konfigurasi empat kluster telah memberikan struktur pengelompokan yang cukup representatif. Temuan ini menunjukkan bahwa pendekatan clustering mampu mengidentifikasi variasi kondisi ketenagakerjaan secara lebih tersegmentasi. Oleh karena itu, pemerintah daerah dapat memanfaatkan hasil pengelompokan ini sebagai dasar dalam menyusun kebijakan yang lebih tepat sasaran, seperti penguatan program pelatihan kerja sesuai kebutuhan wilayah, peningkatan keterampilan berbasis pendidikan, serta perluasan kesempatan kerja pada kluster dengan tingkat pengangguran yang relatif lebih tinggi.

DAFTAR PUSTAKA

- [1] S. Muharni and S. Andriyanto, "Penerapan Metode K-Means Clustering Pada Data Tingkat Pengangguran Terbuka Tahun 2016-2018 Dan 2019-2021," *Jurnal Informatika*, vol. 22, Jun. 2022.
- [2] F. A. Tanjung, A. P. Windarto, and M. Fauzan, "Penerapan Metode K-Means Pada Pengelompokan Pengangguran Di Indonesia," *Jurnal Riset Sistem Informasi Dan Teknik Informatika (JURASIK)*, vol. 6, pp. 61–74, Feb. 2021, [Online]. Available: <https://tunasbangsa.ac.id/ejurnal/index.php/jurasik>
- [3] Andi Diah Kuswanto, Azumardi Nabil Fadhila, Paulus Tri Setiawan, Muhammad Kevin Setiawan, and Dody Renal Syahputra, "Penerapan K-Means Clustering Untuk Menentukan Jumlah Pengangguran Berdasarkan Umur (Studi Kasus Di Badan Statistik Provinsi DKI Jakarta 2020-2022)," *Repeater : Publikasi Teknik Informatika dan Jaringan*, vol. 2, no. 3, pp. 135–146, Jul. 2024, doi: 10.62951/repeater.v2i3.116.
- [4] J. Halif, D. Wahiddin, I. Sanjaya, and S. Faisal, "Model Regresi Linear Berganda untuk Prediksi Tingkat Pengangguran di Provinsi Jawa Barat," *Jurnal Algoritma*, vol. 22, no. 1, pp. 324–335, May 2025, doi: 10.33364/algoritma/v.22-1.2312.
- [5] M. B. Gultom, P. A. Simbolon, and N. S. Nainggolan, "Prediksi Tingkat Pengangguran Berdasarkan Pendidikan Menggunakan Regresi Linear (Studi Kasus : Kota Medan)," *SNISTIK : Seminar Nasional Inovasi Sains Teknologi Informasi Komputer*, vol. 1, no. Mei, pp. 3025–8715, May 2024, [Online]. Available: <https://www.kaggle.com/>
- [6] N. P. Dharshinni, G. Singh, A. Simamora, J. A. T. Naibaho, and R. D. L. Tobing, "Penerapan Algoritma K-Means Pada Data Pengangguran Di Jawa Barat," *Data Sciences Indonesia (DSI)*, vol. 3, no. 1, pp. 23–34, Aug. 2023, doi: 10.47709/dsi.v3i1.2767.
- [7] D. Agustina, Y. I. Mukti, and S. Muntari, "Integrasi Particle Swam Optimization Menggunakan K-Means untuk Klusterisasi Pengangguran di Kota Pagar Alam," *Jurnal*

- Pustaka AI (Pusat Akses Kajian Teknologi Artificial Intelligence)*, vol. 3, no. 1, pp. 34–41, Apr. 2023, doi: 10.55382/jurnalpustakaai.v3i1.543.
- [8] M. Rachma and C. Budy Santoso, “Klasterisasi Data Pengangguran Di Pulau Jawa Menggunakan Algoritma K-Means Dalam Penanggulangan Pengangguran Tahun 2020-2023,” *Idealis: Indonesia Journal Information System*, vol. 8, no. 2, pp. 202–209, Jul. 2025.
- [9] A. Bima Saputra, U. Pradema Sanjaya, and I. Aristia Sa’ida, “Algoritma K-Means Untuk Mengelompokkan Tingkat Pengangguran Terbuka (Tpt) Menurut Provinsi Di Indonesia,” *JURNAL ILMIAH INFORMATIKA*, vol. 15, no. 10, Aug. 2024.
- [10] R. Sari, A. Irma Purnamasari, and A. Rinaldi Dikanda, “Implementasi Data Mining Dalam Menentukan Pola Penjualan Vitamin Blackmores,” *Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 2, Apr. 2023, [Online]. Available: <https://www.kaggle.com/datasets/kiranbudati/mobile->
- [11] S. Rohaniyah and A. Irma Purnamasari, “Clustering Data Pencari Kerja Menurut Tingkat Pendidikan Menggunakan Algoritma K-Means,” *Jurnal Minfo Polgan*, vol. 12, no. 1, pp. 1–14, Mar. 2023, doi: 10.33395/jmp.v12i1.12217.
- [12] T. Yuniarsih and K. Jinan, “Pengelompokan Data Pengangguran Terbuka Menggunakan Algoritma K-Means Berdasarkan Provinsi Kalimantan Barat,” *COBIT*, vol. 01, no. 01, pp. 25–31, May 2024, doi: 10.
- [13] A. Thoriq Basalamah and R. Setyadi, “Penerapan Algoritma K-Means Clustering Pada Tingkat Penyelesaian Pendidikan Di Provinsi Indonesia,” *Jurnal Informatika dan Teknologi Komputer (J-ICOM)*, vol. 04, pp. 114–121, Feb. 2023, [Online]. Available: <https://ejurnalunsam.id/index.php/jicom/>
- [14] A. Ansen Andryadi, “Penerapan Metode Algoritma K-Means Untuk Data Pengangguran Di Jawa Barat,” *Jurnal Wahana Informatika (JWI)*, vol. 2, no. 2, pp. 277–283, 2023.
- [15] F. T. Ismar, M. Reza, M. F. Afif, M. R. Saputra, and A. Ammar, “Implementasi Algoritma C4.5 untuk Klasifikasi Status Tingkat Pengangguran di Indonesia Berdasarkan Jenjang Pendidikan,” *Jurnal Pengabdian Masyarakat dan Riset Pendidikan*, vol. 4, no. 3, pp. 17639–17644, Jan. 2026, doi: 10.31004/jerkin.v4i3.5076.
- [16] A. L. Hananto, A. Hananto, B. Huda, Y. Rahman, E. Novalia, and B. Priyatna, “Determination of Training Participants in Community Work Training Centers Using the Naïve Bayes Classifier Algorithm,” *JOIV: International Journal on Informatics Visualization*, vol. 8, 2024, [Online]. Available: www.joiv.org/index.php/joiv
- [17] D. Setiadi, B. Irawan, and A. Bahtiar, “Penerapan Algoritma K-Means Clustering Pada Pembesaran Jenis Ikan Unggulan Di Provinsi Jawa Barat,” *Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 6, Dec. 2023.
- [18] S. Kusuma Arum, R. Astuti, and F. Muhammad Basysyar, “Penerapan Algoritma K-Means Pada Dataset Pengangguran Terbuka Berdasarkan Pendidikan Di Provinsi Jawa Barat,” *Jurnal Mahasiswa Teknik Informatika*, vol. 8, no. 2, Apr. 2024.
- [19] A. Alif, R. Mulyana, A. Ratna Juwita, A. M. Siregar, and A. Fauzi, “Penerapan Algoritma K-means Clustering dan Hierarchical Clustering dalam Mengelompokkan Data Pengangguran di Karawang,” *Jurnal Algoritma*, vol. 21, 2024, doi: 10.33364/algoritma/v.21-2.2155.
- [20] E. Dwiguna and A. Bahtiar, “Penerapan Data Mining Untuk Menentukan Penerima Bantuan Blt Menggunakan Metode Clustering K-Means Pada Desa Pamulihan,” *Jurnal Mahasiswa Teknik Informatika*, vol. 8, no. 2, Apr. 2024.
- [21] D. Sutris Martua Simanjuntak, I. Gunawan, and I. Purnama Sari, “Penerapan Algoritma K-Medoids Untuk Pengelompokkan Pengangguran Umur 25 tahun Keatas Di Sumatera Utara,” *Jurnal Krisnadana*, vol. 3, Jan. 2023, [Online]. Available: <https://ejournal.catuspata.com/index.php/jkdn/index>
- [22] R. Sebastian, A. P. Kusuma, and W. D. Puspitasari, “Penggunaan Metode K-Nearest Neighbor Untuk Mendiagnosa Kerusakan Laptop Dengan Teknik Data Mining,” *Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 2, Apr. 2023.
- [23] R. Lorentiana Wijayanti, R. Kurniawan, R. Herdiana, and H. Susana, “Komparasi Algoritma Apriori Dan Fp-Growth Untuk Memberikan Strategi Diskon,” *Jurnal Mahasiswa Teknik Informatika*, vol. 8, no. 1, p. 1096, Feb. 2024.
- [24] E. Novalia, J. Na’, G. W. Nurcahyo, and A. Voutama, “Website Implementation with the Monte Carlo Method as a Media for Predicting Sales of Cashier Applications,” *SYSTEMATICS*, vol. 2, no. 3, pp. 118–131, 2020.
- [25] B. Huda, A. S. Amin, F. Nurapriani, and A. Damuri, “Aplikasi Monitoring Perkembangan Edukasi Anak Usia Dini Berbasis Web,” *Jurnal Informatika Utama*, vol. 1, no. 1, pp. 1–10, Jun. 2023, doi: 10.55903/jitu.v1i1.70
- [26] T. Paryono, E. Novalia, B. Y. Astario, and A. Hananto, “Pelatihan Penggunaan Paint Untuk Melatih Motorik Siswa Anak Sekolah Dasar,” *VIDHEAS: Jurnal Nasional Abdimas Multidisiplin*, vol. 2, no. 2, Dec. 2024, [Online]. Available: <https://vinicho.id/index.php/vidheas>