

ANALISIS KLASTERISASI WILAYAH PENYANDANG DISABILITAS DI PROVINSI JAWA BARAT

Maulidina Azizah, Nining Rahaningsih, Raditya Danar Dana

STMIK IKMI Cirebon

Jalan Perjuangan No.10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Jawa Barat 45135

azizahmaulidina18@gmail.com

ABSTRAK

Penelitian ini menggunakan algoritma K-Means untuk mengklasterisasi beberapa wilayah di Provinsi Jawa Barat berdasarkan jumlah penduduk penyandang disabilitas. Data yang digunakan berasal dari Dinas Kependudukan dan Pencatatan Sipil Provinsi Jawa Barat, yang mencakup berbagai jenis disabilitas seperti cacat fisik, netra atau buta, rungu atau berbicara, mental atau jiwa, fisik, atau mental. Selain itu, penelitian ini menggunakan program RapidMiner untuk melakukan analisis data yang efektif. Hasil klasterisasi menghasilkan tiga klaster yang dianggap penting. Klaster 0 fokus pada Kabupaten Ciamis, Klaster 1 fokus pada Kota Bandung, dan Klaster 2 fokus pada Kabupaten Pangandaran. Klasterisasi ini dianggap ideal dengan nilai *Davies-Bouldin Index* (DBI) sebesar 0.214. Diharapkan hasil penelitian ini akan berkontribusi pada pembuatan kebijakan yang lebih sesuai dan efektif untuk masyarakat penyandang disabilitas di Provinsi Jawa Barat.

Kata kunci : *Klastering, K-Means, Disabilitas*

1. PENDAHULUAN

Penelitian ini dilakukan karena perkembangan pesat dibidang teknologi informasi yang menyebabkan pengaruh yang besar pada berbagai aspek kehidupan, salah satunya adalah kesehatan. Maka, diperlukan penelitian untuk mengklasterisasi wilayah di Provinsi Jawa Barat berdasarkan jumlah penduduk penyandang disabilitas menggunakan *K-Means*.

Tantangan dalam mengklasterisasi wilayah berdasarkan jumlah penduduk penyandang disabilitas menggunakan algoritma *K-Means* adalah kompleks dan berdampak sosial. Era digital membawa peluang besar untuk menggali data berharga. Namun, keterbatasan data, teknologi, sumber daya manusia dan pemahaman tentang disabilitas perlu diatasi. Upaya ini sangat penting karena akan memberikan dasar yang kuat untuk perencanaan kebijakan yang lebih inklusif bagi masyarakat penyandang disabilitas di Provinsi Jawa Barat. Selain itu, penggunaan algoritma *K-Means* dalam konteks ini dapat mengisi kesenjangan pengetahuan dan memperkuat peran informatika dalam pembangunan sosial.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Sebelumnya, penelitian telah dilakukan terkait dengan analisis klaster pada data penderita disabilitas di wilayah lain menggunakan algoritma *K-Means*. Meskipun belum terdapat studi yang secara khusus menangani klasterisasi wilayah berdasarkan jumlah penduduk penyandang disabilitas di Provinsi Jawa Barat menggunakan algoritma *K-Means*. Oleh karena itu, terdapat kesenjangan penelitian yang dapat diisi melalui penelitian ini untuk memberikan pemahaman yang lebih mendalam tentang distribusi penduduk penyandang disabilitas di wilayah tersebut.

Studi oleh Dwina, yang implementasi algoritma *K-Means* untuk mengelompokkan disabilitas

berdasarkan jenisnya di Kabupaten Sidoarjo. Hasil analisis menunjukkan empat klaster optimal dengan nilai koefisien siluet tertinggi sebesar 0,33[1].

Studi oleh Wildani Putri, yang menerapkan algoritma *K-Means* untuk pengelompokan data penyandang disabilitas di Kabupaten Rokan Hilir. Hasil dari pengelompokan data tersebut menghasilkan tiga klaster, yaitu cluster_0 yang merupakan kelompok kecamatan dengan disabilitas rendah, cluster_2 yang merupakan kelompok kecamatan dengan disabilitas sedang, dan cluster_1 yang merupakan kelompok kecamatan dengan disabilitas terbanyak. Dari 18 kecamatan yang ada di Kabupaten Rokan Hilir, terdapat satu kecamatan yang termasuk dalam cluster tingkat tinggi untuk wilayah penyandang disabilitas terbanyak, yaitu kecamatan Bangko. Selain itu, terdapat tiga kecamatan yang termasuk dalam cluster tingkat sedang, yaitu kecamatan Rimba Melintang, Rantau Kopar, dan Pujud, serta 14 kecamatan lainnya yang termasuk dalam cluster tingkat rendah[2].

Studi oleh Kristiyo Indriadi Wardoyo, menerapkan algoritma *K-Means* untuk mengelompokkan siswa penyandang disabilitas berdasarkan tingkat tunagrahita. Berdasarkan hasil penelitian, Dari perhitungan nilai jarak Euclidean dengan menggunakan metode pengambilan jarak terdekat, ditemukan bahwa 47,9% individu dalam cluster 1 (tingkat tunagrahita ringan), 37% dalam cluster 2 (tingkat tunagrahita sedang), dan 15,1% dalam cluster 3 (tingkat tunagrahita berat)[3].

Meskipun studi-studi sebelumnya telah dilakukan, namun terdapat peluang untuk penelitian lebih lanjut. Beberapa keterbatasan dalam studi sebelumnya adalah tidak adanya aplikasi yang mendukung dalam mempermudah pengelompokan data penyandang disabilitas berdasarkan disabilitas di seluruh Indonesia. Maka, diharapkan penelitian ini

mampu memberikan kontribusi dalam mengatasi keterbatasan tersebut.

Tujuan utama penelitian ini adalah mengklasterisasi wilayah di Provinsi Jawa Barat berdasarkan jumlah penduduk penyandang disabilitas menggunakan algoritma *K-Means*. Dengan kata lain, penelitian ini bertujuan untuk mengelompokkan wilayah-wilayah tersebut menjadi klaster berdasarkan kemiripan dalam jumlah penduduk penyandang disabilitas, sehingga dapat memberikan pemahaman yang lebih baik mengenai distribusi dan karakteristik penyandang disabilitas di setiap klaster.

Algoritma klastering *K-Means* akan digunakan dalam penelitian ini. *K-Means* merupakan teknik klasterisasi yang populer dalam analisis data dan akan diterapkan untuk mengelompokkan wilayah-wilayah di Provinsi Jawa Barat berdasarkan jumlah penduduk penyandang disabilitas. Algoritma *K-Means* mengelompokkan data berdasarkan total pencarian iterative centroid dan menetapkan bahwa area setiap klaster data adalah jarak minimal dari centroid. Kelebihan algoritma *K-Means* adalah kesederhanaannya dan kemampuan untuk mengelompokkan data besar dan outlier dengan sangat cepat[1].

Hasil dari penelitian klasterisasi menggunakan *K-Means*, klaster 0 terdiri dari Kabupaten Bandung, Kabupaten Bogor, Kabupaten Ciamis, Kabupaten Cianjur, Kabupaten Cirebon, Kabupaten Garut, Kabupaten Kuningan, Kabupaten Sukabumi, dan Kabupaten Tasikmalaya dengan Kabupaten Ciamis sebagai fokus utama yang mendominasi dalam berbagai jenis disabilitas. Klaster 1 melibatkan Kota Bandung, Kota Banjar, Kota Bekasi, Kota Bogor, Kota Cimahi, Kota Cirebon, Kota Depok, Kota Sukabumi, dan Kota Tasikmalaya. Kota Bandung menjadi pusat perhatian dengan jumlah penyandang disabilitas tertinggi dalam beberapa kategori, mendominasi Klaster 1. Klaster 2 terdiri dari Kabupaten Bandung Barat, Kabupaten Bekasi, Kabupaten Indramayu, Kabupaten Karawang, Kabupaten Majalengka, Kabupaten Pangandaran, Kabupaten Purwakarta, Kabupaten Subang, dan Kabupaten Sumedang. Meskipun jumlah penyandang disabilitas lebih rendah, Kabupaten Pangandaran menonjol sebagai yang tertinggi dalam berbagai jenis disabilitas, sehingga mendominasi Klaster 2.

2.2. Disabilitas

Goffman berpendapat bahwa individu dengan disabilitas dianggap memiliki keterbatasan yang membuat mereka kesulitan berkomunikasi dengan orang lain. Lingkungan sering melihat mereka sebagai individu yang tidak mampu melakukan aktivitas tanpa menimbulkan masalah. Akibatnya, karena adanya keterbatasan dan stigma negatif dari orang lain, mereka berusaha untuk tidak tergantung pada individu lain dan berupaya meyakinkan diri agar tidak bergantung sepenuhnya pada orang lain[4].

2.3. Data Mining

Data mining adalah proses eksplorasi pola yang secara otomatis menganalisis dan mengekstrak pengetahuan dari kumpulan data untuk memecahkan masalah. Oleh karena itu, *Knowledge Discovery in Database* (KDD) adalah istilah umum untuk data mining. [5].

2.4. Klastering

Klastering atau klasterisasi merupakan Teknik klasifikasi data yang digunakan untuk menemukan struktur klaster dalam kumpulan data.[6].

2.5. K-Means

Algoritma klastering *K-Means* memecah data ke dalam berbagai kelompok, di mana data dengan kesamaan tinggi dikelompokkan bersama, sedangkan data yang memiliki perbedaan karakteristik ditempatkan dalam kelompok yang berbeda[7].

2.6. DBI

Davies-Bouldin Index adalah parameter evaluasi yang digunakan dalam analisis data untuk mengukur kualitas hasil pengelompokan. Fungsinya adalah untuk mengukur seberapa efektif algoritma pengelompokan dalam memisahkan berbagai kelompok data dan mendekati pusat masing-masing klaster[8].

2.7. RapidMiner

Dr. Markus Hofmann dari Institut Teknologi Blanchardstown dan Ralf Klinkenberg dari rapid-i.com membuat perangkat lunak yang disebut *RapidMiner*. Aplikasi ini dibuat dengan antarmuka pengguna grafis (GUI) yang ramah pengguna. *RapidMiner*, yang dibangun menggunakan bahasa pemrograman Java dan dirilis di bawah lisensi publik GNU, dapat digunakan di banyak sistem operasi, menjadikannya alat yang populer.[9].

3. METODE PENELITIAN

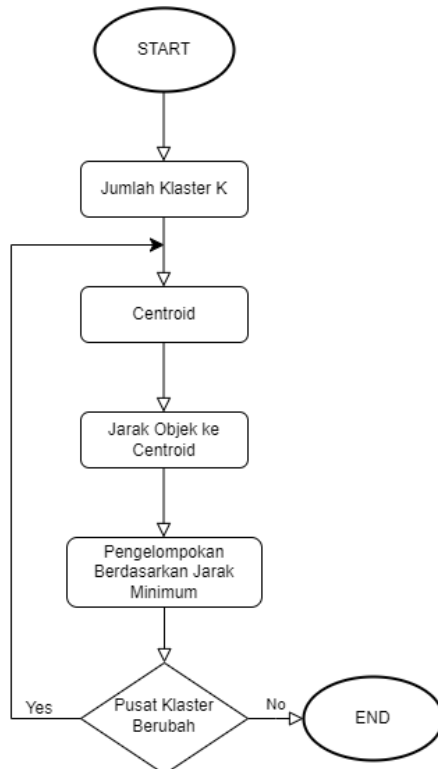
3.1. Flowchart Metode Klasterisasi

Gambar flowchart metode klasterisasi, yang digunakan algoritma *K-Means*, dijelaskan pada gambar 1.

- Menentukan Jumlah Cluster**
Langkah pertama adalah mengetahui berapa banyak klaster yang akan digunakan untuk mengelompokkan data.
- Inisialisasi Centroid**
Untuk setiap klaster, centroid diaktifkan di posisi awalnya.
- Pembagian Data**
Berdasarkan kedekatan centroid, data dibagi ke dalam berbagai klaster. Penempatan data dalam klaster yang tepat dilakukan dengan menghitung jarak geometris.
- Iterasi dan Konvergensi**
Sebelum konvergensi terjadi, iterasi dilakukan untuk memastikan bahwa data ditempatkan dengan benar dalam klaster yang telah terbentuk.

Pembaruan posisi centroid dilakukan dalam proses iteratif ini berdasarkan rata-rata data di setiap kluster setelah pembagian.

Langkah 3 dan 4 diulang hingga posisi centroid tidak berubah atau iterasi mencapai batas yang ditentukan. Ini memastikan konvergensi dan penentuan posisi centroid yang ideal.



Gambar 1. Flowchart Metode Klusterisasi

3.2. Metode Penelitian

Untuk mendapatkan pengetahuan dari database yang tersedia, pendekatan *Knowledge Discovery in Database (KDD)* digunakan[3]. Berikut tahapan dari KDD :

a. Data Selection

Tahap ini bertujuan untuk memahami atribut-atribut dan kebutuhan pengguna terhadap hasil analisis data.

b. Data Preprocessing

Tahap ini bertujuan untuk membersihkan data dari *noise*, *missing value*, dan *outlier*.

c. Data Transformation

Melakukan transformasi data seperti normalisasi, *discretization*, dan *feature selection*.

d. Data Mining

Menerapkan algoritma *K-Means* untuk mengelompokkan wilayah di Provinsi Jawa Barat berdasarkan jumlah penduduk penyandang disabilitas.

e. Interpretation

Menginterpretasi hasil klustering untuk memberikan informasi berupa kluster data penyandang disabilitas yang dapat dimanfaatkan oleh pemerintah terkait dalam menentukan

kebijakan yang berhubungan dengan penyandang disabilitas

3.3. Sumber Data

Dataset ini diambil dari Open Data Jabar, dataset ini berisi jumlah penyandang disabilitas di Provinsi Jawa Barat berdasarkan kategori disabilitas dari tahun 2013 hingga 2022. Data ini dibuat oleh Dinas Kependudukan dan Pencatatan Sipil dan didistribusikan selama satu tahun.

3.4. Teknik Pengumpulan Data

Penelitian ini menggunakan data sekunder dari Dinas Kependudukan dan Pencatatan Sipil Provinsi Jawa Barat. Data ini menunjukkan jumlah penduduk penyandang disabilitas di Provinsi Jawa Barat dari tahun 2013 hingga 2022, berdasarkan kategori disabilitas.

3.5. Teknik Analisis Data

Penelitian ini menggunakan analisis kuantitatif. Menggunakan data sekunder, algoritma *K-Means* dapat digunakan untuk mengklasifikasikan data wilayah berdasarkan jumlah penduduk penyandang disabilitas di Provinsi Jawa Barat. Algoritma *K-Means*, yang pertama kali diterbitkan oleh Stuart Lloyd pada tahun 1984 dan banyak digunakan sebagai algoritma klustering, adalah metode analisis data yang dapat digunakan setelah data dikumpulkan. Algoritma *K-Means* sangat cepat, dapat disesuaikan, dan digunakan secara luas[10]. Algoritma *K-Means* digunakan untuk mengelompokkan data menjadi beberapa kluster berdasarkan kesamaan jumlah penyandang disabilitas di setiap wilayah. Hasil pengelompokan kemudian diinterpretasikan dan disimpulkan.

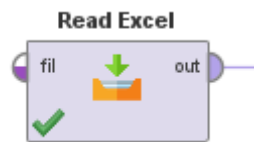
4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Row No.	id	nama_kabu...	kategori_dis...	jumlah_pen...
1	1	KABUPATEN ...	CACAT FISIK	3629
2	2	KABUPATEN ...	CACAT METR...	1389
3	3	KABUPATEN ...	CACAT RUN...	1782
4	4	KABUPATEN ...	CACAT MENT...	2437
5	5	KABUPATEN ...	CACAT FISIK	482
6	6	KABUPATEN ...	CACAT LAI...	1714
7	7	KABUPATEN ...	CACAT FISIK	1686
8	8	KABUPATEN ...	CACAT METR...	1009
9	9	KABUPATEN ...	CACAT RUN...	1122
10	10	KABUPATEN ...	CACAT MENT...	1559
11	11	KABUPATEN ...	CACAT FISIK	178
12	12	KABUPATEN ...	CACAT LAI...	841
13	13	KABUPATEN ...	CACAT FISIK	2565
14	14	KABUPATEN ...	CACAT METR...	967
15	15	KABUPATEN ...	CACAT RUN...	1108

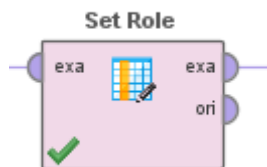
Gambar 2. Hasil data selection

Gambar 2 menunjukkan *output* data pemilihan dengan *RapidMiner*. Tahapan yang digunakan untuk proses pemilihan data yaitu *Read Excel*, *Set Role*, *Select Attributes*.



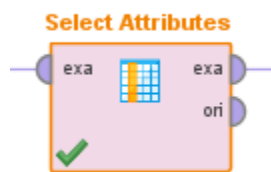
Gambar 3. Operator read excel

Tahap ini melibatkan membaca data dari file Excel ke dalam *RapidMiner*. Data yang dibaca dapat berupa spreadsheet atau file Excel



Gambar 4. Operator Set Role

Set Role sangat penting untuk mengidentifikasi atribut yang akan dijadikan variabel sasaran atau variabel prediktor. Hal ini memastikan bahwa analisis yang dilakukan konsisten dengan tujuan penelitian.



Gambar 5. Operator Select Attributes

Pemilihan Atribut (*Attribute Selection*) Langkah ini dilakukan untuk memfilter atribut yang paling relevan untuk analisis kluster. Atribut yang digunakan yaitu jumlah penduduk, kategori disabilitas dan nama kabupaten.

4.2. Data Preprocessing

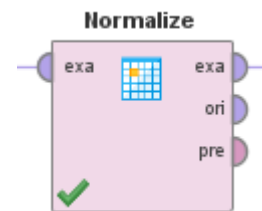
Name	Type	Missing	Statistics	Filter (6 attributes)	Search for Attributes
id	Integer	0	Min 1 Max 162 Average 81.500		
cluster	Nominal	0	Local cluster_2 (54) Global cluster_0 (54)		
nama_kabupaten_kota	Numeric	0	Min 0 Max 26		
kategori_disabilitas	Numeric	0	Min 0 Max 5 Average 2.500		
tahun	Numeric	0	Min 0 Max 1 Average 0.006		

Gambar 6. Hasil missing value

Dalam tahap Data Preprocessing, dilakukan pemeriksaan terhadap keberadaan nilai yang hilang (*missing value*) pada dataset. Pada dataset jumlah penduduk penyandang disabilitas di Provinsi Jawa Barat, tidak ditemukan adanya *missing value* seperti pada gambar 2.

4.3. Data Transformation

Proses transformasi data melalui Normalize dan Nominal to Numerical pada *RapidMiner* dilakukan untuk menyempurnakan persiapan data sebelum analisis klusterisasi.



Gambar 7. Operator normalize

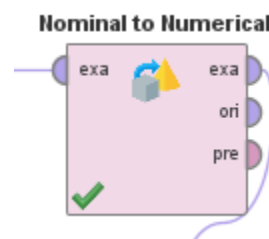
Pada tahap normalisasi menggunakan subset dengan atribut "jumlah penduduk" dan menerapkan metode range transformation, dilakukan standarisasi skala untuk mencegah dominasi variabel berbasis skala besar terhadap variabel lain dalam analisis klusterisasi. Pendekatan ini penting untuk memastikan setiap variabel numerik memiliki rentang yang seragam, sehingga bobot yang setara diberikan dalam pembentukan kelompok oleh algoritma *K-Means*. Penggunaan subset pada atribut "jumlah penduduk" menjadikan proses normalisasi lebih terfokus, hanya memengaruhi variabel yang dianggap relevan untuk analisis klusterisasi. Pemilihan metode *range transformation* dipertimbangkan guna memberikan skala antar variabel yang konsisten, menjaga integritas distribusi data, dan meningkatkan keakuratan dalam konteks analisis klusterisasi.

Berikut adalah hasil dari normalisasi

Row No.	id	jumlah_pes...	nama_kab...	kategori_dis...	tahun
1	1	0.038	KABUPATEN ...	CACAT FISIK ...	2013
2	2	0.012	KABUPATEN ...	CACAT NETR ...	2013
3	3	0.014	KABUPATEN ...	CACAT RUNT...	2013
4	4	0.013	KABUPATEN ...	CACAT MENI...	2013
5	5	0.006	KABUPATEN ...	CACAT FISIK ...	2013
6	6	0.006	KABUPATEN ...	CACAT LAH...	2013
7	7	0.018	KABUPATEN ...	CACAT FISIK ...	2013
8	8	0.011	KABUPATEN ...	CACAT NETR...	2013
9	9	0.009	KABUPATEN ...	CACAT NETR...	2013
10	10	0.007	KABUPATEN ...	CACAT MENI...	2013
11	11	0.002	KABUPATEN ...	CACAT FISIK ...	2013
12	12	0.008	KABUPATEN ...	CACAT LAH...	2013
13	13	0.024	KABUPATEN ...	CACAT FISIK ...	2013
14	14	0.010	KABUPATEN ...	CACAT NETR...	2013
15	15	0.010	KABUPATEN ...	CACAT RUNT...	2013

Gambar 8. Hasil normalisasi

Setelah normalisasi kemudian dilanjutkan dengan *nominal to numerical*.



Gambar 9. Operator nominal to numerical

Dalam langkah *nominal to numerical*, subset yang dipilih melibatkan atribut "kategori disabilitas"

dan "nama kabupaten". Langkah ini bertujuan mengkonversi variabel kategori disabilitas dari tipe nominal menjadi bentuk numerik menggunakan coding type unique integers. Pendekatan ini dilakukan agar informasi kategori disabilitas dapat diolah oleh algoritma klasterisasi. Penggunaan coding type unique integers memberikan representasi numerik unik untuk setiap label kategori disabilitas, memungkinkan pengukuran yang lebih obyektif terhadap perbedaan kategori disabilitas dalam analisis klasterisasi. Proses ini secara positif mendukung kemampuan algoritma *K-Means* dalam memahami pola dan hubungan antar data, khususnya dalam membentuk kelompok dengan karakteristik serupa.

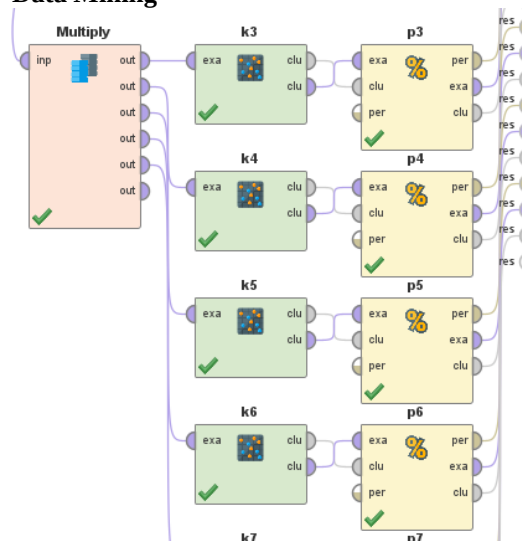
Tujuan dari *nominal to numerical* adalah untuk mengubah atribut kategorikal menjadi bentuk numerik sehingga dapat dimasukkan dalam algoritma klastering. Selama proses ini, metode bilangan bulat unik dipilih untuk mewakili setiap kategori. Gambar 10 merupakan hasil dari *nominal to numerical*.

Row No.	id	nama_kabu...	kategori_dis...	jumlah_pen...
1	1	0	0	0.052
2	2	0	1	0.019
3	3	0	2	0.025
4	4	0	3	0.035
5	5	0	4	0.006
6	6	0	5	0.024
7	7	1	0	0.024
8	8	1	1	0.014
9	9	1	2	0.016
10	10	1	3	0.022
11	11	1	4	0.002
12	12	1	5	0.011
13	13	2	0	0.037
14	14	2	1	0.013
15	15	2	2	0.015

ExampleSet (162 examples, 1 special attribute, 3 regular attributes)

Gambar 10. Hasil *nominal to numerical*

4.4. Data Mining



Gambar 11. Data mining

Dalam implementasi data mining untuk klasterisasi wilayah berdasarkan jumlah penyandang disabilitas di Provinsi Jawa Barat, penggunaan parameter khusus dalam algoritma *K-Means* menjadi langkah penting. Ada beberapa parameter, seperti "*Max runs 10*", "*Type of measure: numerical measure*", "*Numerical measure: Euclidean Distance*", dan "*Step optimization maksimum 100*". Pengaturan ini mencerminkan pendekatan yang matang untuk memaksimalkan kinerja algoritma dan mengekstrak pola kluster yang optimal. Pemilihan *Measure type*, terutama *Euclidean Distance*, menjadi faktor kunci dalam menentukan kedekatan antar-poin dalam ruang numerik, memungkinkan pembentukan kluster yang responsif terhadap karakteristik numerik jumlah penyandang disabilitas di wilayah tersebut. Penggunaan "*Max runs*" dan "*Max optimization steps*" memberikan batasan efektif untuk mencegah konvergensi prematur dan memastikan pencarian solusi klasterisasi yang optimal. Hasil klasterisasi dengan parameter ini berhasil membentuk tiga kluster dengan karakteristik serupa, dibuktikan oleh evaluasi performa menggunakan *Davies-Bouldin Index (DBI)* dengan nilai terendah sebesar 0,214, yang mengonfirmasi kehomogenan yang baik antar kluster. Oleh karena itu, penggunaan parameter khusus ini tidak hanya meningkatkan efektivitas algoritma, tetapi juga memberikan dasar yang kuat untuk interpretasi lebih mendalam terkait distribusi penyandang disabilitas di wilayah Provinsi Jawa Barat.

Proses Data Mining pada *RapidMiner* menghasilkan *Cluster Model* yang menunjukkan hasil kluster yang optimal untuk dataset penduduk penyandang disabilitas di Provinsi Jawa Barat. Rincian hasil kluster adalah sebagai berikut:

Cluster Model

```
Cluster 0: 54 items
Cluster 1: 54 items
Cluster 2: 54 items
Total number of items: 162
```

Gambar 12. Hasil cluster model

4.5. Interpretation

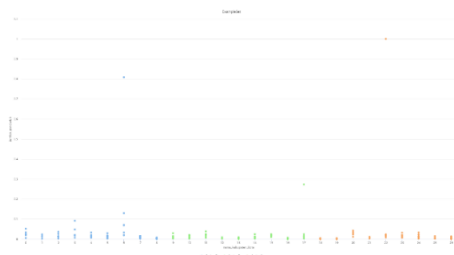
Tabel 1. Nilai DBI Tiap Kluster

Kluster	Performance	DBI
K3	P3	0.214
K4	P4	0.239
K5	P5	0.270
K6	P6	0.301
K7	P7	0.323

Hasil evaluasi menunjukkan bahwa jumlah kluster optimal untuk dataset ini adalah K3 dengan nilai DBI 0.214. Pemilihan ini didasarkan pada nilai terendah karena semakin mendekati nol, semakin baik performa klasterisasi [11]. Perhitungan *Davies-*

Bouldin Index (DBI) dilakukan secara otomatis pada RapidMiner sebagai bagian dari evaluasi klusterisasi. Untuk menjelaskan cara RapidMiner menghitung DBI dan mengapa nilai terendahnya adalah 0.214, berikut adalah penjelasan umumnya:

- Konfigurasi Algoritma K-Means**
Langkah pertama melibatkan pengaturan algoritma K-Means dengan menentukan jumlah kluster yang diinginkan, dalam hal ini tiga kluster (K3).
- Proses Algoritma K-Means**
Algoritma ini mengelompokkan data ke dalam kluster berdasarkan seberapa dekat data berada di ruang fitur.
- Perhitungan DBI**
Nilai Davies-Bouldin Index (DBI) secara otomatis dihitung oleh RapidMiner setelah proses klusterisasi selesai. Ini melibatkan menghitung jarak antar-kluster dan intra-kluster serta menghitung rasio yang sesuai dengan rumus DBI.
- Penilaian Kualitas Klusterisasi**
Nilai DBI yang dihasilkan menunjukkan seberapa baik klusterisasi dilakukan. Nilai terendah, yaitu 0.214, menunjukkan bahwa klusterisasi tiga kluster memberikan hasil terbaik.



Gambar 13. Visualisasi K3

Dari gambar visualisasi menjelaskan bahwa kluster 0 terdiri dari Kabupaten Bandung, Kabupaten Bogor, Kabupaten Ciamis, Kabupaten Cianjur, Kabupaten Cirebon, Kabupaten Garut, Kabupaten Kuningan, Kabupaten Sukabumi, dan Kabupaten Tasikmalaya. Jenis disabilitas yang dominan meliputi cacat fisik, netra/buta, rungu/wicara, mental/jiwa, fisik dan mental, serta jenis cacat lainnya. Kabupaten Ciamis menjadi fokus utama dengan jumlah penyandang disabilitas tertinggi, khususnya pada kategori cacat fisik (55,499), cacat fisik dan mental (1,395), mental/jiwa (4,808), netra/buta (8,965), rungu/wicara (4,906), dan jenis cacat lainnya (3,327). Dengan demikian, Kabupaten Ciamis mendominasi Kluster 0.

Kluster 1 mencakup Kota Bandung, Kota Banjar, Kota Bekasi, Kota Bogor, Kota Cimahi, Kota Cirebon, Kota Depok, Kota Sukabumi, dan Kota Tasikmalaya. Jenis disabilitas yang signifikan termasuk cacat fisik (2,978), cacat fisik dan mental (884), mental/jiwa (2,105), netra/buta (1,701), rungu/wicara (2,116), dan jenis cacat lainnya (68,584). Kota Bandung menjadi pusat perhatian dengan jumlah penyandang disabilitas

tertinggi dalam beberapa kategori, sehingga mendominasi Kluster 1.

Kluster 2 terdiri dari Kabupaten Bandung Barat, Kabupaten Bekasi, Kabupaten Indramayu, Kabupaten Karawang, Kabupaten Majalengka, Kabupaten Pangandaran, Kabupaten Purwakarta, Kabupaten Subang, dan Kabupaten Sumedang. Meskipun jumlah penyandang disabilitas lebih rendah (62,683), beberapa kabupaten menonjol, seperti Kabupaten Pangandaran yang menjadi yang tertinggi pada cacat fisik (18,777), netra/buta (1,767), mental/jiwa (2,105), rungu/wicara (1,410), dan cacat fisik dan mental (795). Sehingga, Kabupaten Pangandaran menjadi dominan dalam Kluster 2. Identifikasi wilayah dengan status terbanyak dapat menjadi dasar bagi pihak berwenang untuk fokus pada upaya pencegahan dan rehabilitasi.

5. KESIMPULAN DAN SARAN

Hasil penelitian yang menggunakan algoritma K-Means untuk klusterisasi wilayah berdasarkan jumlah penduduk penyandang disabilitas di Provinsi Jawa Barat menunjukkan bahwa pembentukan tiga kluster berhasil menggambarkan distribusi penyandang disabilitas dengan jelas. Kabupaten Ciamis adalah rumah bagi Kluster 0, yang menonjol dalam berbagai jenis disabilitas. Kota Bandung berada di Kluster 1 dengan jumlah penyandang disabilitas tertinggi dalam berbagai kategori. Kabupaten Pangandaran berada di Kluster 2, dengan jumlah penyandang disabilitas yang lebih rendah. Evaluasi Davies-Bouldin Index (DBI) dengan nilai rendah sebesar 0,214 menunjukkan hasil klusterisasi yang cukup baik dengan batasan yang jelas dan homogenitas internal yang memadai; penentuan jumlah kluster ideal (K3) didukung oleh nilai rendah DBI, yang menunjukkan pembentukan kluster yang baik dengan batasan dan homogenitas yang memadai.

Disarankan untuk melakukan penelitian lebih lanjut pada masing-masing kluster, seperti Kabupaten Ciamis, Kota Bandung, dan Kabupaten Pangandaran. Ini akan membantu kita memahami variabel yang memengaruhi distribusi penyandang disabilitas di setiap daerah. Untuk mendapatkan pemahaman yang lebih luas, disarankan untuk mempertimbangkan variabel tambahan seperti tingkat pendidikan, kebijakan kesehatan, dan aksesibilitas infrastruktur. Diharapkan dengan menerapkan rekomendasi ini, kebijakan dan program dukungan di Provinsi Jawa Barat akan lebih tepat sasaran dan efektif dalam meningkatkan kesejahteraan penyandang disabilitas.

DAFTAR PUSTAKA

- [1] F. I. Dwiana, W. D. Utami, A. Hamid, and S. Asih, "Implementasi *K-Means* pada Klusterisasi Jenis Disabilitas," *J. Math. Comput. Stat.*, vol. 6, no. 1, pp. 92–101, 2023, doi: 10.35580/jmathcos.v6i1.41719.
- [2] W. Putri and M. Afdal, "Application Of The *K-Means* Algorithm For Data Grouping Persons With Disabilities In Rokan Hilir District Penerapan Algoritma *K-Means* Untuk

- Pengelompokan Data Penyandang Disabilitas Di Kabupaten Rokan Hilir,” vol. 3, no. 1, pp. 30–38, 2023.
- [3] K. I. Wardoyo and J. Fredricka, “KLAUSTERISASI SISWA PENYANDANG TUNAGRAHITA MENGGUNAKAN ALGORITMA K-MEANS,” vol. 19, no. 1, pp. 1–10, 2023.
- [4] et al. Kurniadi, Y U., “Penyandang Disabilitas di Indoneisa,” *Nusant. J. Ilmu Pengetah. Sos.*, vol. 7, no. 2, pp. 408–420, 2020.
- [5] L. N. Ningsi and H. S. Tambunan, “Implementasi Data Mining Kluster Pada Rumah Tangga Yang Memiliki Akses Hunian Layak Berdasarkan Provinsi,” vol. 1, no. 5, pp. 228–234, 2021.
- [6] K. P. Sinaga and M. Yang, “Unsupervised K-Means Clustering Algorithm,” vol. 8, 2020, doi: 10.1109/ACCESS.2020.2988796.
- [7] S. Nelson Butarbutar¹, Agus Perdana Windarto², Dedi Hartama³, “KOMPARASI KINERJA ALGORITMA FUZZY C-MEANS DAN K-MEANS DALAM PENGELOMPOKAN DATA SISWA BERDASARKAN PRESTASI NILAI AKADEMIK SISWA (Studi Kasus : SMP Negeri 2 Pematangsiantar),” vol. 1, no. 2012, 2016.
- [8] I. T. Umagapi and B. Umaternate, “Uji Kinerja K-Means Clustering Menggunakan Davies-Bouldin Index Pada Pengelompokan Data Prestasi Siswa,” pp. 303–308.
- [9] B. G. Sudarsono, M. I. Leo, A. Santoso, and F. Hendrawan, “Analisis Data Mining Data Netflix Menggunakan Aplikasi Rapid Miner,” *JBASE - J. Bus. Audit Inf. Syst.*, vol. 4, no. 1, pp. 13–21, 2021, doi: 10.30813/jbase.v4i1.2729.
- [10] D. Marlina, N. Lina, A. Fernando, and A. Ramadhan, “Implementasi Algoritma K-Medoids dan K-Means untuk Pengelompokan Wilayah Sebaran Cacat pada Anak,” *J. CoreIT J. Has. Penelit. Ilmu Komput. dan Teknol. Inf.*, vol. 4, no. 2, p. 64, 2018, doi: 10.24014/coreit.v4i2.4498.
- [11] Y. A. Wijaya, D. A. Kurniady, E. Setyanto, W. S. Tarihoran, D. Rusmana, and R. Rahim, “Davies Bouldin Index Algorithm for Optimizing Clustering Case Studies Mapping School Facilities,” *TEM J.*, vol. 10, no. 3, pp. 1099–1103, 2021, doi: 10.18421/TEM103-13.