

ANALISIS DATASET STATUS GIZI PADA BALITA MENGGUNAKAN ALGORITMA K-MEANS CLUSTERING

Fina Raudotul Janah, Rudi Kurniawan, Tati Suprapti

Teknik Informatika, STMIK IKMI Cirebon

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

finaraudotuljanah21@gmail.com

ABSTRAK

Status gizi balita merupakan satu permasalahan yang harus diperhatikan, terutama pada asupan gizi serta nutrisi yang diperoleh dari makanan sehari-hari. Sebagai bagian dari siklus pertumbuhan dan perkembangan buah-buahan, mereka membutuhkan asupan gizi yang lebih baik, karena balita paling mudah menderita kelainan gizi. Pemenuhan kebutuhan gizi merupakan faktor penting untuk mencapai hasil tumbuh kembang pada balita dan anak-anak. Puskesmas Karangsembung merupakan fasilitas pelayanan kesehatan yang menyelenggarakan upaya kesehatan masyarakat dan upaya kesehatan perseorangan tingkat pertama. Pihak puskesmas belum melakukan proses data mining untuk meng-*clustering* dataset balita. Sehingga belum mengetahui kelompok status gizi balita dan belum mengetahui nilai evaluasi DBI. Tujuan dari Penelitian ini untuk mengelompokkan status gizi balita menggunakan algoritma K-means untuk mendapatkan nilai *Cluster* terbaik dengan evaluasi nilai DBI dengan memanfaatkan tools Rapidminer. Hasil dari penelitian didapatkan jumlah *Cluster* optimal dengan K=5. Untuk *Cluster* 0: 41 Balita dengan rata-rata berjenis kelamin Laki-laki, *Cluster* 1: 83 Balita dengan rata-rata berjenis kelamin Laki-laki, *Cluster* 2: 82 Balita dengan rata-rata berjenis kelamin Perempuan, *Cluster* 3: 70 Balita dengan rata-rata berjenis kelamin Perempuan, *Cluster* 4: 56 Balita dengan rata-rata berjenis kelamin Laki-laki dan nilai DBI yang optimal sebesar 0.132 dimana nilai tersebut mendekati 0 yang berarti klaster yang di evaluasi menghasilkan klaster yang baik.

Kata kunci : K-Means, *Clustering*, Status Gizi Balita, *Davies Bouldin Index* (DBI)

1. PENDAHULUAN

Anak-anak berusia dibawah lima tahun atau yang dikenal dengan sebutan Balita, memiliki resiko tinggi terhadap permasalahan Kesehatan dan gizi. Permasalahan gizi pada balita memiliki karakteristik yang berbeda dari masalah gizi pada orang dewasa, masa pertumbuhan dan perkembangan balita merupakan periode kritis dalam membentuk dasar kesehatan dan kualitas hidup di masa mendatang. Gizi menjadi kebutuhan pokok bagi tubuh manusia, terutama saat masa tumbuh kembang pada balita dan anak-anak. Balita mengalami siklus pertumbuhan dan perkembangan yang membutuhkan asupan gizi yang lebih baik dari kelompok unsur lain, sehingga balita paling mudah menderita kelainan gizi [1]. Salah satu perbedaannya adalah bahwa masalah gizi pada balita sering kali sulit terdeteksi oleh masyarakat, bahkan oleh keluarga mereka sendiri. Akibatnya, jika beberapa anak di suatu desa mengalami masalah gizi, kemungkinan mereka tidak segera mendapatkan perhatian yang dibutuhkan karena gejala mereka tidak selalu terlihat dengan jelas [2].

Pada penelitian ini terkait tentang pengelompokan Status gizi balita pada Puskesmas Karangsembung, saat ini kurangnya pengetahuan tentang pentingnya gizi seimbang di kalangan Orang Tua dapat mempengaruhi usaha dalam mencegah kasus gizi pada Balita [3]. Oleh karena itu perlu dilaksanakan sosialisasi tentang asupan gizi yang baik untuk pertumbuhan balita dan anak. Hal ini perlu diperhatikan, karena hanya dengan mengetahui perkembangan balita berdasarkan fisik saja tentu

tidak cukup untuk mengetahui status gizi dari balita tersebut. Semestinya orang tua perlu memperhatikan lebih jauh mengenai perkembangan balita. Perlu dilakukan analisis terhadap data Status gizi pada balita menggunakan algoritma K-Means *Clustering*. Melakukan analisis data dengan menggunakan Teknik K-Means *Clustering* menggunakan tools RapidMiner [4]. Analisis pengelompokan menggunakan data mining yang bertujuan untuk membentuk kelompok gizi balita dengan *cluster* terbaik. Perlu dilakukan evaluasi untuk melihat tingkat performa dari hasil pemodelan yang dilakukan menggunakan *Davies Bouldin Index* (DBI) [5].

Beberapa penelitian yang dilakukan oleh peneliti sebelumnya mengenai algoritma K-Means *Clustering* Berdasarkan studi-studi sebelumnya, terbukti bahwa algoritma K-Means efisien untuk mengelompokkan data berlaku untuk berbagai jenis data [6].

Penelitian sebelumnya oleh Chandra, M. D., Irawan, E., Saragih, I. S dalam penelitiannya yang berjudul Penerapan Algoritma K-Means Dalam Mengelompokkan Balita yang Mengalami Gizi Buruk menurut Provinsi dari hasil penelitian ini melalui pengelompokan, teridentifikasi 2 kelompok utama, yaitu kelompok tinggi dan kelompok rendah. Kelompok tinggi mencakup 15 Provinsi sementara kelompok rendah melibatkan 19 Provinsi.

Dan pada penelitian sebelumnya oleh Saputro dan Suci hermayanti dalam jurnal nya yang berjudul Penerapan Klaterisasi Menggunakan K-Means untuk

menentukan Tingkat Kesehatan Bayi dan Balita di Kabupaten Bengkulu Utara berdasarkan pengujian Algoritma K-Means diterapkan untuk mengelompokkan setiap kecamatan berdasarkan indikator seperti angka kematian bayi, angka kematian balita, angka kesakitan, dan status gizi. Hasil dari pengolahan indikator ini menghasilkan tiga kelompok, yang mencakup tingkat Kesehatan tinggi, sedang, dan rendah [7].

Tujuan dari analisis ini untuk mengelompokkan status gizi balita menggunakan algoritma K-Means untuk mendapatkan nilai terbaik dengan menggunakan *Davis Bouldin Index* (DBI) dilengkapi dengan alat Rapidminer. DBI digunakan sebagai index dalam analisis *Cluster* [8]. Memperbaiki atau meningkatkan *Cluster* dalam Algoritma *Clustering* dengan melakukan validasi *Davies Bouldin index* (DBI), yang membantu mengidentifikasi kualitas *Cluster* terbaik pada dataset yang digunakan [9]. Teknik pengumpulan data yang digunakan dalam penelitian ini yaitu *clustering* [10]. *Clustering* merupakan salah satu teknik data mining yang masuk dalam kategori tanpa pengawasan (*unsupervised*), yang artinya metode ini digunakan tanpa pelatihan, panduan dan tidak memerlukan *output* target [11].

Penelitian ini menggunakan teknik *Clustering* dengan metode Algoritma K-Means yang merupakan metode untuk membagi sejumlah objek ke dalam kelompok berdasarkan karakteristik yang sama (jumlah bilangan bulat positif). Menunjukkan bahwa sebuah *Cluster* memiliki massa, yang merupakan representasi mean dari *Cluster* tersebut [12]. Data yang dimanfaatkan dalam penelitian ini yaitu data mengenai status gizi, yang terdiri dari berbagai atribut. Atribut-atribut itu meliputi No, Nama, Jenis kelamin, Usia (Bulan), Tinggi Badan dan Berat Badan. Data mining merupakan serangkaian langkah untuk mengeksplorasi pengetahuan yang sebelumnya tidak terdeteksi melalui metode manual [13].

Dengan demikian Hasil analisis diharapkan dapat memberikan pengetahuan tentang algoritma K-Means, selain itu hasil analisis ini dapat memberikan kontribusi dalam meningkatkan pemahaman dan Tindakan yang lebih efektif mengenai status gizi balita, sehingga upaya dalam meningkatkan kesejahteraan anak-anak dapat dilakukan secara tepat sasaran. Pada penelitian ini diharapkan memberikan pemahaman bagi masyarakat terutama Orang Tua, tentang status gizi dan pemberian gizi yang cukup untuk memenuhi perkembangan anak.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data Mining merupakan proses analitik yang memanfaatkan Teknik-teknik statistika, matematika dan kecerdasan buatan untuk mengeksplorasi dan menggali informasi yang berharga atau pola yang tersembunyi dari kumpulan data yang besar dan kompleks [14]. Tujuan dari data mining untuk mengenali relasi, kecenderungan atau pola yang

mungkin tak terlihat secara langsung, sehingga memberikan wawasan serta pemahaman yang lebih mendalam mengenai data tersebut.

2.2. Algoritma K-Means Clustering

Pemberian nutrisi yang memadai untuk memenuhi pertumbuhan anak. Algoritma ini relatif mudah digunakan dalam proses pengelompokan, di mana dataset dibagi menjadi beberapa kelompok atau klaster. [15]. Tujuan utamanya adalah mengurangi Perbedaan di antara data dalam satu klaster, sambil memaksimalkan perbedaan dengan klaster lainnya [16]. Dalam melakukan pengelompokan data, metode K-means melibatkan beberapa langkah, dengan kata lain [15]:

1. Tentukan nilai K

2. Inisialisasi *Cluster* K ke *cluster* pusat

Pusat klaster K dapat diinisialisasi menggunakan berbagai metode, tetapi yang umumnya digunakan adalah metode acak. Pusat-pusat klaster diberikan nilai awal secara acak dengan menggunakan angka acak [17]. Pada fase iterasi, rumus yang diterapkan adalah sebagai berikut:

$$V_{ij} = \frac{1}{N_i} \sum_{k=0}^{N_i} X_{kj}$$

Keterangan:

V_{ij} = Centroid rata-rata *Cluster* ke-I untuk variable ke-j

N_i = Jumlah anggota *Cluster* ke-ii, k = Indeks dari variable

X_{kj} = Nilai data ke-k variable ke-j untuk *Cluster* tersebut

3. Letakkan semua data ke dalam cluster yang paling dekat. Kedekatan dua objek ditentukan oleh jarak di antara keduanya, dan hal serupa berlaku untuk kedekatan data ke klaster.

4. Hitung Jarak antara setiap data dan pusat cluster. ada menggunakan rumus jarak euclidean sebagai berikut:

Keterangan:

$$D(i,j) = \sqrt{(X_{1i} - X_{1j})^2 + \dots + (X_{ki} - X_{kj})^2}$$

$D(i,j)$ = Jarak data I ke pusat j

X_{ki} = Data ke-i dari atribut data ke-k X_{kj} = Titik tengah ke-j dari atribut ke- k

Setelah itu, hitung ulang jarak dari pusat cluster dengan mempertimbangkan keanggotaan cluster saat ini. Pusat cluster adalah nilai rata-rata dari semua data yang ada di cluster tertentu. Evaluasi *clustering* dilakukan untuk mengukur sejauh mana kualitas pengelompokan yang akan diperoleh [18]. Metode yang diterapkan dalam menilai *cluster* ini memanfaatkan *Davies Bouldin Index* (DBI) [19], DBI didefinisikan sebagai perbandingan antara jarak rata-rata di dalam setiap *cluster* dan di antara *cluster* dengan *cluster* lain terdekatnya [20]. DBI merupakan salah satu metode evaluasi internal yang menilai cluster dalam suatu metode pengelompokan dengan memperhatikan nilai kohesi dan separasi [19].

$$DBI = \frac{1}{k} \sum_{i=1}^k \max_{i \neq j} (R_{ij})$$

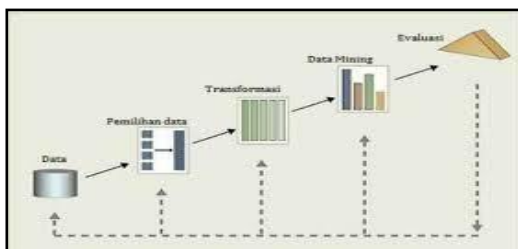
Dari rumus tersebut, k adalah jumlah kelompok yang diterapkan. Nilai yang dihasilkan yang lebih rendah (tidak negatif ≤ 0), semakin optimal penerapan K-means dalam menghasilkan cluster [21].

2.3. RapidMiner

Perangkat RapidMiner bersifat terbuka dan dapat digunakan oleh siapa saja. Berbagai Teknik, termasuk deskriptif dan prediksi, digunakan oleh RapidMiner sebagai solusi untuk menganalisa data pengolahan. Untuk menjalankan RapidMiner ini, menggunakan Bahasa Java [22].

3. METODE PENELITIAN

pada penelitian ini adalah menggunakan proses *Knowledge Discovery in Databases (KDD)*. Teknik analisis ini akan dilakukan menggunakan tools RapidMiner, berikut tahapan untuk analisis data dapat dilihat pada Gambar 1 di bawah ini.



Gambar 1. tahapan proses kdd , sumber: [14]

1. Data Selection

Pada tahap ini dataset balita yang terkumpul dilakukan pembersihan dan seleksi pada atribut data yang akan digunakan dalam analisis seperti penghapusan atribut-atribut data yang tidak relevan untuk analisis dan memilih atribut untuk di jadikan sebagai id.

2. Data Preprocessing

Kemudian dilakukan Data *Preprocessing* dengan tujuan untuk membersihkan data seperti data *null* dan data duplikat pada informasi yang telah dipilih dalam RapidMiner.

3. Data Transformation

Data yang dipilih dan dipulihkan kemudian di ubah atau di transformasikan untuk memungkinkan analisis yang baik seperti mengubah bentuk data, tipe data dan lainnya dalam RapidMiner.

4. Data Mining

Pada titik ini, data akan diproses dengan menggunakan operator K-means pada RapidMiner. Terdapat beberapa parameter yang perlu digunakan salah satunya *Davies Bouldin Index (DBI)*.

5. Evaluasi

Pada tahap evaluasi digunakan untuk mengukur keefektifan dan validitas model atau pola yang ditemukan. Evaluasi dilakukan dengan cara mencari nilai DBI yang menghasilkan nilai terbaik pada pengelompokan gizi balita.

4. HASIL DAN PEMBAHASAN

4.1 Data Selection

Tahapan pertama, dalam proses pengelompokan data, proses seleksi data digunakan untuk memilih atau memilih data yang dikumpulkan. Dataset yang digunakan Balita pada Bulan Juli 2023 yang diperoleh dari Departemen Kesehatan (Puskesmas). Dataset yang diperoleh dalam bentuk *Microsoft Excel*, Dataset terdiri dari 332 dengan 6 variabel. Variabel tersebut diantaranya No, Nama, Jenis Kelamin, Usia (Bulan), Berat Badan, Tinggi Badan terlihat Pada Tabel 1, dibawah ini:

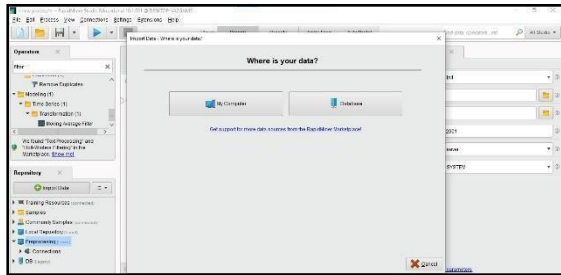
Tabel 1. Dataset

No.	Nama	JK	Usia (Bulan)	Berat Badan	Tinggi Badan
1.	Meisya Aira A	P	58	14,44	103
2.	Ahmad Zibrikun	L	58	19,87	108,1
3.	Zapar	L	56	13,6	98,3
4.	Vidia ArtikaS	P	56	14,44	108
...		...			
...		...			
332.	Akmal Firdaus	L	7	6,98	69

Proses pengambilan data dalam penelitian ini dilakukan sleksi dengan memanfaatkan fungsi *Edit List* pada Parameter *Read Excel* di RapidMiner. Sebelum melakukan proses sleksi dataset yang sudah disiapkan di import terlebih dahulu kedalam RapidMiner. Setelah itu lakukan lembar proses yang berisi implementasi algoritma dan melakukan uji coba dengan RapidMiner menggunakan algoritma K-Means. Berikut Langkah-langkah proses dapat dilihat di bawah ini

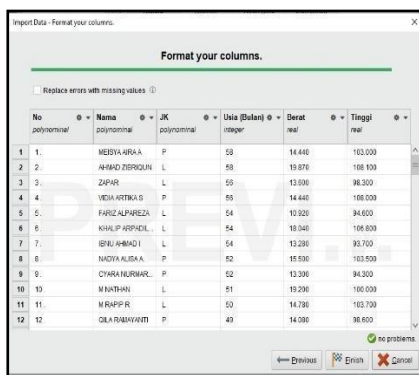
4.2 Import Data

Sebelum penggunaan Algoritma K-means Langkah pertama data dalam format *.xsl perlu dilakukan import data kedalam RapidMiner dengan mengklik "Add Data", kemudian memilih data pada lokasi di mana data yang akan digunakan akan disimpan. Karena kumpulan data yang akan digunakan berbentuk *.xsl bukan berbentuk Jika menggunakan SQL, Dalam proses import data, pilih lokasi dataset. dapat dilihat pada Gambar 1 dibawah ini.



Gambar 2. import Data

Selanjutnya, Seleksi data yang akan di Mengimpor ke dalam RapidMiner. Jika data yang dipilih tidak menunjukkan adanya kendala sehingga timbul “no problem” sepanjang proses import data. Kemudian, lanjutkan hingga selesai untuk memastikan bahwa data berhasil. diimpor kedalam RapidMiner.

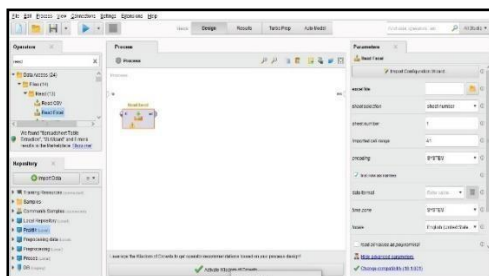


Gambar 3. proses import data

Dalam proses import data, dataset harus diberi identitas atau identifikasi. Hal ini dilakukan dengan mengubah atribut “Nama Balita” melalui klik logo pengaturan, kemudian mengganti tipe peran dengan mengklik logo pengaturan, kemudian mengganti tipe peran dengan mengklik “change role type” kemudian pilih opsi ID lalu klik ok.

4.3 Membuat Lembar Proses

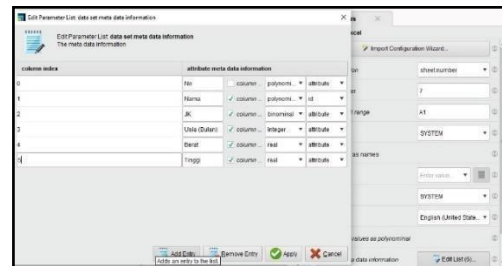
Setelah dataset Balita diimpor kemudian pilih operators “Read Excel” kemudian Drag dan Drop dataset yang telah diimpor kedalam operator Read Excel pada halaman Proses. Implementasi ini dapat dilihat pada Gambar 4 di bawah dalam RapidMiner.



Gambar 4. read excel data yang telah diimpor

4.4 Select Atribut

Seleksi Atribut dilakukan pada Operators Read Excel dengan melakukan Edit List. Pada dataset awal terdapat 6 atribut, diantaranya atribut No, Nama Balita, Jenis Kelamin, Usia (Bulan), Berat Badan dan Tinggi Badan. Kemudian seleksi atribut mana saja yang akan digunakan dan yang tidak akan digunakan. Proses seleksi data dilihat pada Gambar 5 di bawah ini.



Gambar 5. parameter list

Berikut Data yang telah diseleksi dengan proses select atribut dilihat pada Gambar 6 di bawah ini.

Row No.	Nama	No	JK	Usia (Bulan)	Berat	Tinggi
1	MEITYA AURA	1	P	58	14.440	103
2	ARHARD ZIERLI	2	L	58	19.870	106.100
3	ZAPAR	3	L	56	13.600	96.300
4	VEDA ARIKA	4	P	56	14.440	100
5	FARIZ ALFAREZ	5	L	54	19.920	94.600
6	KHALIP ARPHIDI	6	L	54	18.040	106.800
7	IBNU ARHADI	7	L	54	13.280	93.700
8	NADYA ALISA	8	P	52	15.500	103.500
9	CHARA KUR	9	P	52	13.300	94.300
10	MATHIAS	10	L	51	19.200	100

Gambar 6. data yang telah diseleksi

4.5 Pre-processing Data

Tahap Pre-processing Data adalah proses pre-processing data ketika ada variable yang tidak dapat dihitung (*Missing Value*). Pada tahap ini, tidak ada data yang ditemukan tentang ketidak lengkapan dalam set data, atau mungkin disebut sebagai "No Missing Value". seperti yang ditunjukkan oleh hasil statistik RapidMiner pada Gambar 7 di bawah ini.

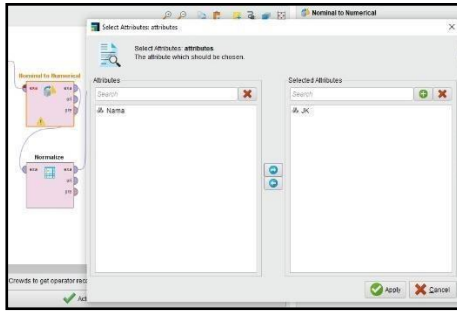
Name	Type	Missing	Statistics	Filter (14 attributes)	Sort on Statistics
Nama	Nominal	0	ZUSKARMEV (1)	FAUZI ALFAREZ (2)	FAUZI ALFAREZ
No	Nominal	0	9 (1)	100 (1)	100 (1)
JK	Nominal	0	P (152)	L (106)	L (106)
Usia (Bulan)	Integer	0	51	58	58
Berat	Real	0	5.290	19.870	19.870

Gambar 7. pengecekan missing value

4.6 Transformation Data

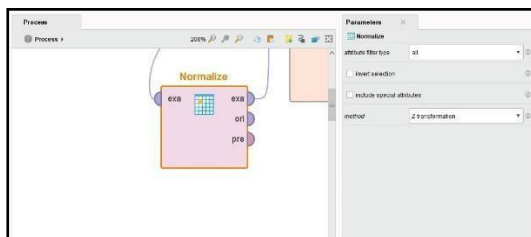
Pada tahap Transformation Data, Operators Nominal to Numerical digunakan data dan Operators Normalize digunakan untuk memungkinkan analisis dan pengolahan dataset yang telah melewati Langkah- langkah sebelumnya dengan menerapkan metode K- Means Clustering. Operators Normalize pada dataset Balita digunakan untuk menormalisasi sebelum Algoritma K-Means menggunakan nilai DBI

untuk pengelompokan. Gambar 8 menunjukkan proses transfer data di bawah ini.



Gambar 8. parameter operator nominal to numerical

Berikut *Operators Normalize* pada dataset Balita digunakan untuk menormalisasi nilai DBI dapat dilihat pada Gambar 9 di bawah ini.



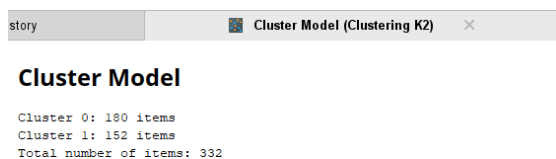
Gambar 9. parameter operator normalize

4.7 Data Mining

Pada fase eksplorasi data (Data Mining) pada penelitian ini, dilakukan kegiatan klasterisasi dengan memanfaatkan Algoritma K-Means untuk menemukan nilai Klaster terbaik dalam pengelompokan Status Gizi Balita dengan *tools* RapidMiner. Dalam Langkah-langkah tersebut, digunakan *operators cluster distance performance* untuk melakukan evaluasi menggunakan metode nilai DBI. Dengan melakukan percobaan iterasi sebanyak 7 kali dengan menentukan nilai K mulai dari K2, K3, K4, K5, K6 dan K7. Berikut adalah hasil dari pengujian iterasi tersebut di bawah ini:

1. Hasil iterasi dengan nilai K2

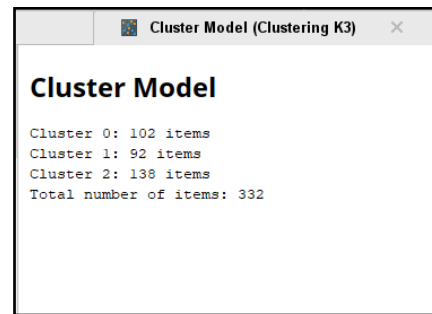
Berikut merupakan hasil iterasi K2 mendapatkan nilai *Avg. within centroid distance* 0.597 dan didapatkan *Cluster model* pada Gambar 10 di bawah ini:



Gambar 10. cluster model K2

2. Hasil iterasi dengan nilai K3

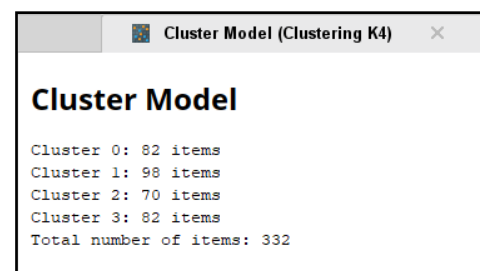
Berikut merupakan hasil iterasi K3 mendapatkan nilai *Avg. within centroid distance* 0.375 dan didapatkan *Cluster model* pada Gambar 11 di bawah ini:



Gambar 11. cluster model K3

3. Hasil iterasi dengan nilai K4

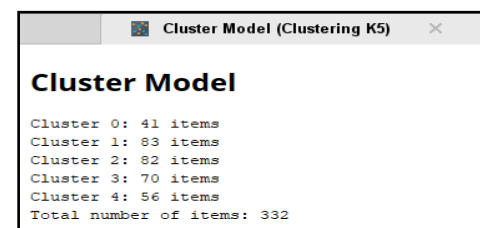
Berikut merupakan hasil iterasi K4 mendapatkan nilai *Avg. within centroid distance* 0.218 dan didapatkan *Cluster model* pada Gambar 12 berikut:



Gambar 12. cluster model K4

4. Hasil iterasi dengan nilai K5

Berikut merupakan hasil iterasi K5 mendapatkan nilai *Avg. within centroid distance* 0.162 dan didapatkan *Cluster model* pada Gambar 13 di bawah ini:

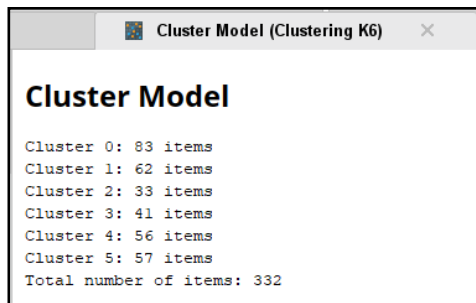


Gambar 13. cluster model K5

Hasil dari evaluasi diperoleh nilai DBI yang mendekati 0 didapatkan pada nilai K5 yaitu sebesar 0.132.

5. Hasil iterasi dengan nilai K6

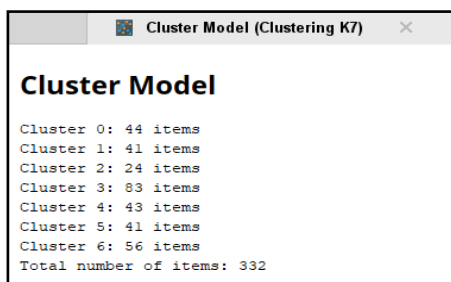
Berikut merupakan hasil iterasi K6 mendapatkan nilai *Avg. within centroid distance* 0.125 dan didapatkan *Cluster model* pada Gambar 14 di bawah ini:



Gambar 14. cluster model K6

6. Hasil iterasi dengan nilai K7

Berikut merupakan hasil iterasi K7 mendapatkan nilai *Avg. within centroid distance* 0.106 dan didapatkan *Cluster* model pada Gambar 15 di bawah ini:



Gambar 15. cluster model K7

4.8 Evaluasi

Pada tahap evaluasi dilakukan pengujian pada nilai DBI yang mendekati 0. Berikut merupakan hasil nilai DBI dari 7 kali percobaan iterasi pada tabel 1, berikut:

Tabel 1. evaluasi hasil nilai dbi

K	<i>Avg. within centroid distance</i>	DBI
2	0.597	0.210
3	0.375	0.172
4	0.218	0.136
5	0.162	0.132
6	0.125	0.137
7	0.106	0.144

4.9 Pembahasan

- Penerapan Algoritma K-Means Clustering
Sebelum melakukan penerapan algoritma K-means. data sebanyak 332 baris, dilakukan analisis menggunakan metode KDD yaitu tahap selection, preprocessing, dan data mining, pada tahap data mining dilakukan penerapan algoritma K-means *Clustering* menggunakan parameter DBI, sehingga dapat mengetahui nilai (K) terbaik pada evaluasi DBI dengan dilakukan percobaan iterasi sebanyak 7 kali.
- Hasil Analisis Pengelompokan Dataset Balita dengan Algoritma K-Means
Dari hasil percobaan iterasi 7 kali dengan nilai K2, K3, K4, K5, K6 dan K7 diperoleh nilai DBI

terendah yaitu pada nilai K5 dengan menghasilkan *Cluster* 0: 41 Balita, dengan rata-rata balita berjenis kelamin Laki-laki, tinggi badan rata-rata 72,724 CM dan berat badan rata-rata 8,533 KG; *Cluster* 1: 83 balita laki-laki dengan rata-rata tinggi badan 87,571 CM dan rata-rata berat badan 11,49 KG; *Cluster* 2: 82 balita perempuan dengan rata-rata tinggi badan 96,941 CM dan rata-rata berat badan 13,69 KG; dan *Cluster* 3: 70 balita perempuan dengan rata-rata tinggi badan 77,084 CM dan rata-rata berat badan 85,33KG.

- Hasil Evaluasi Nilai Davies Bouldin Index (DBI)
Hasil evaluasi DBI diperoleh nilai DBI yang mendekati 0 yaitu pada K5 dengan nilai sebesar 0.132 sehingga K5 dinyatakan Kelompok terbaik. Hal ini sejalan dengan penelitian oleh [23] Dalam penelitiannya yang menyatakan semakin rendah nilai *Devies Bouldin Indeks* (DBI), semakin baik kualitas kluster yang dihasilkan.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil dan pembahasan dapat diambil kesimpulan bahwa algoritma K-Means dapat diterapkan dalam mengelompokan dataset balita menghasilkan 5 *Cluster* dengan *Cluster* 0: 41 Balita, dengan rata-rata balita berjenis kelamin Laki-laki, *Cluster* 1: 83 balita laki-laki dengan rata-rata tinggi badan 87,571 CM dan rata-rata berat badan 11,49 KG; *Cluster* 2: 82 balita perempuan dengan rata-rata tinggi badan 96,941 CM dan rata-rata berat badan 13,69 KG; dan *Cluster* 3: 70 balita perempuan dengan rata-rata tinggi badan 72,724 CM dan rata-rata berat badan 8,533 KG. dengan *Cluster* 4-56 balita, dengan rata-rata perempuan 77,084 CM dan rata-rata berat badan 9,246 KG. Laki-laki memiliki rata-rata tinggi badan 100,97 CM dan rata-rata berat badan 9,246 KG. 15,12 dan nilai DBI terbaik adalah 0.132, hampir 0, menunjukkan bahwa klaster yang dievaluasi menghasilkan klaster yang baik. Penelitian ini dapat di kembangkan dengan algoritma lainnya seperti K-Medoid, X-Means dan lainnya untuk mengetahui perbandingan hasil yang diperoleh dari masing-masing algoritma.

DAFTAR PUSTAKA

- W. Malinda, Z. Hayati, N. Ramadhanty, and Y. F. Putri, "PERKEMBANGAN ANAK Kesehatan Diri Dan Lingkungan : Pentingnya Gizi Bagi Perkembangan Anak perkembangan yang fundamental bagi kehidupannya kelak . Pada tahapan usia kehidupannya kelak . untuk meningkatkan sumber daya manusia secara sistematis dan memaksimal," *J. Multidisipliner Bharasumba*, vol. 01, no. 02, pp. 269–278, 2022.
- N. Rizki Octaviyani and R. Mayasari, "Implementasi Algoritma K-Means Clustering Status Gizi Balita," *J. Ilm. Wahana Pendidik.*, vol. 8, no. 13, pp. 370–381, 2022,

- [Online]. Available: <https://doi.org/10.5281/zenodo.6962588>
- [3] Sri Rahayu Savitri *et al.*, “Pencegahan Kasus Stunting Melalui Penyuluhan Remaja Dan Pmt (Pemberian Makanan Tambahan) Di Desa Purbosono,” *J-ABDI J. Pengabd. Kpd. Masy.*, vol. 2, no. 7, pp. 5521–5528, 2022, doi: 10.53625/jabdi.v2i7.3990.
 - [4] R. Ordila, R. Wahyuni, Y. Irawan, and M. Yulia Sari, “PENERAPAN DATA MINING UNTUK PENGELOMPOKAN DATA REKAM MEDIS PASIEN BERDASARKAN JENIS PENYAKIT DENGAN ALGORITMA CLUSTERING (Studi Kasus : Poli Klinik PT.Inecda),” *J. Ilmu Komput.*, vol. 9, no. 2, pp. 148–153, 2020, doi: 10.33060/jik/2020/vol9.iss2.181.
 - [5] I. Soliani and S. Juanita, “Grouping the Prevalence of Disease Cases By Age in Bandung City Hospitals Using K-Means,” *J.Tek. Inform.*, vol. 3, no. 6, pp. 1647–1654, 2022, doi: 10.20884/1.jutif.2022.3.6.430.
 - [6] N. Fadhila Aini and S. Sulastri, “Perbandingan Algoritma K-Means dan K-Median Clustering Terhadap Nilai Ujian Nasional SMP di Jawa Tengah,” *Smart Comp Jurnalnya Orang Pint. Komput.*, vol. 11, no.4, 2022, doi: 10.30591/smartcomp.v11i4.3915.
 - [7] D. T. Saputro and W. P. Sucihermayanti, “Penerapan Klasterisasi Menggunakan K-Means untuk Menentukan Tingkat Kesehatan Bayi dan Balita di Kabupaten Bengkulu Utara,” *J. Buana Inform.*, vol. 12, no. 2, pp. 146–155, 2021, doi: 10.24002/jbi.v12i2.4861.
 - [8] A. Supriyadi, A. Triayudi, and I. D. Sholihati, “Perbandingan Algoritma K-Means Dengan K-Medoids Pada Pengelompokan Armada Kendaraan Truk Berdasarkan Produktivitas,” *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.)*, vol. 6, no. 2, pp. 229–240, 2021, doi: 10.29100/jupi.v6i2.2008.
 - [9] F. Tempola, M. Muhammad, and A. Mubarak, “Penggunaan Internet Dikalangan Siswa SD di Kota Ternate: Suatu Survey, Penerapan Algoritma Clustering dan Validasi DBI,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 7, no. 6, pp. 1153–1160, 2020, doi: 10.25126/jtiik.2020722370.
 - [10] T. Hardiani, “Analisis Clustering Kasus Covid 19 di Indonesia Menggunakan Algoritma K-Means,” *J. Nas. Pendidik. Tek.Inform.*, vol. 11, no. 2, pp. 156–165, 2022, doi: 10.23887/janapati.v11i2.45376.
 - [11] R. Hablum, A. Khairan, and R. Rosihan, “Clustering Hasil Tangkap Ikan Di Pelabuhan Perikanan Nusantara (Ppn) Ternate Menggunakan Algoritma K-Means,” *JIKO (Jurnal Inform. dan Komputer)*, vol. 2, no. 1, pp. 26–33, 2019, doi: 10.33387/jiko.v2i1.1053.
 - [12] L. Studi, K. Pada, U. D. Toko, and B. Yd, “Penerapan Algoritma K-Means Untuk Menentukan Bahan Bangunan”.
 - [13] M. R. Alhapizi, M. Nasir, and I. Effendy, “Penerapan Data Mining Menggunakan Algoritma K-Means Clustering Untuk Menentukan Strategi Promosi Mahasiswa Baru Universitas Bina Darma Palembang,” *J. Softw. Eng. Ampera*, vol. 1, no. 1, pp. 1–14, 2020, doi: 10.51519/journalsea.v1i1.10.
 - [14] Agung Nugraha, Odi Nurdian, and Gifthera Dwilestari, “Penerapan Data Mining Metode K-Means Clustering Untuk Analisa Penjualan Pada Toko Yana Sport,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 6, no. 2, pp. 1–7, 2022.
 - [15] ary bagus jiwandono, “Analisa Pengelompokan Kelas BPJS Kesehatan Dengan Menggunakan Metode K-Means,” 2021.
 - [16] R. Muliono and Z. Sembiring, “Data Mining Clustering Menggunakan Algoritma K-Means Untuk Klasterisasi Tingkat Tridarma Pengajaran Dosen,” *CESS (Journal Comput. Eng. Syst. Sci.)*, vol. 4, no. 2, pp. 2502–714, 2019.
 - [17] D. Sartika and J. Jumadi, “Seminar Nasional Teknologi Komputer & Sains (SAINTEKS) Clustering Penilaian Kinerja Dosen Menggunakan Algoritma K-Means (Studi Kasus: Universitas Dehasen Bengkulu),” pp. 703–709, 2019, [Online]. Available: <https://seminar-id.com/seminas-sainteks2019.html>
 - [18] M. Mughnyanti, S. Efendi, and M. Zarlis, “Analysis of determining centroid clustering x-means algorithm with davies-bouldin index evaluation,” *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 725, no. 1, 2020, doi: 10.1088/1757-899X/725/1/012128.
 - [19] Z. Nabila, A. Rahman Isnain, and Z. Abidin, “Analisis Data Mining Untuk Clustering Kasus Covid-19 Di Provinsi Lampung Dengan Algoritma K-Means,” *J. Teknol. dan Sist. Inf.*, vol. 2, no. 2, p. 100, 2021, [Online]. Available: <http://jim.teknokrat.ac.id/index.php/JTSI>
 - [20] Y. Sopyan, A. D. Lesmana, and C. Juliane, “Analisis Algoritma K-Means dan Davies Bouldin Index dalam Mencari Cluster Terbaik Kasus Perceraian di Kabupaten Kuningan,” *Build. Informatics, Technol. Sci.*, vol. 4, no. 3, pp. 1464–1470, 2022, doi: 10.47065/bits.v4i3.2697.
 - [21] Amalia Yunia Rahmawati, *Pengenalan Data Mining*, no. July. 2020.
 - [22] Y. R. Sari, A. Sudewa, D. A. Lestari, and T. Jaya, “Penerapan Algoritma K-Means Untuk Clustering Data Kemiskinan Provinsi Banten Menggunakan Rapidminer,” *CESS (Journal Comput. Eng. Syst. Sci.)*, vol. 5, no. 2, p. 192, 2020, doi: 10.24114/cess.v5i2.18519.
 - [23] R. Adhitama, A. Burhanuddin, and R. Ananda,

“Penentuan Jumlah Cluster Ideal Smk Di Jawa Tengah Dengan Metode X- Means Clustering Dan K-Means Clustering Determining Vocational Ideal ClusterNumber in Central Java With X-Means Clustering and K-Means Clustering Methods,” *J. Inform. dan Komputer) Akreditasi KEMENRISTEKDIKTI*, vol. 3, no. 1, pp. 1–5, 2020, doi: 10.33387/jiko.