

SEGMENTASI DAERAH PRIORITAS PENANGANAN *STUNTING* DI JAWA BARAT MENGGUNAKAN METODE *K-MEANS CLUSTERING*

Aliffia Kurniawati Komaria, Bayu Priyatna, Tukino, Agustia Hananto

Sistem Informasi, Universitas Buana Perjuangan Karawang
Jalan Ronggo Waluyo Sirnabaya, Puseurjaya, Telukjambe Timur, Indonesia
Si22.aliffiakomaria@mhs.ubpkarawang.ac.id

ABSTRAK

Stunting pada balita masih menjadi permasalahan kesehatan masyarakat di Provinsi Jawa Barat yang ditandai dengan adanya perbedaan tingkat prevalensi antarwilayah. Permasalahan utama terletak pada belum optimalnya pemetaan wilayah berbasis analisis tren kuantitatif, sehingga penentuan prioritas intervensi belum sepenuhnya tepat sasaran. Penelitian ini bertujuan untuk menganalisis pola perubahan prevalensi *stunting* serta mengelompokkan kabupaten/kota berdasarkan kesamaan tren selama periode 2014–2024. Metode yang digunakan adalah *K-Means Clustering* dengan tahapan pre-processing data, normalisasi menggunakan metode Min-Max, penentuan jumlah kluster menggunakan *Elbow Method*, serta evaluasi menggunakan *Silhouette Score*. Hasil penelitian menunjukkan bahwa wilayah dapat dikelompokkan menjadi tiga kluster, yaitu tinggi, sedang, dan rendah. Kluster tinggi memiliki jumlah kasus yang relatif besar dan fluktuatif, kluster sedang menunjukkan tren penurunan yang konsisten, sedangkan kluster rendah memiliki jumlah kasus yang relatif kecil dan stabil. Hasil ini diharapkan dapat mendukung penentuan kebijakan penanganan *stunting* yang lebih terarah dan berbasis data.

Kata kunci : *stunting, K-Means Clustering, data mining, Jawa Barat, tren data, clustering*

1. PENDAHULUAN

Stunting pada balita masih menjadi salah satu permasalahan kesehatan yang krusial di Provinsi Jawa Barat. Kondisi ini tidak hanya berdampak pada pertumbuhan fisik, tetapi juga berpengaruh terhadap perkembangan kognitif dan kualitas sumber daya manusia di masa depan [1].

Perbedaan tingkat prevalensi *stunting* antar kabupaten/kota menunjukkan adanya ketimpangan kondisi yang perlu dianalisis secara lebih mendalam. Ketersediaan data terbuka dari pemerintah daerah sebenarnya memberikan peluang untuk melakukan analisis berbasis data guna memahami pola dan dinamika *stunting* secara lebih komprehensif [2].

Namun demikian, pemanfaatan data tersebut belum optimal karena sebagian besar analisis masih bersifat deskriptif dan belum mampu mengidentifikasi kesamaan karakteristik antarwilayah secara kuantitatif [3].

Kondisi ini menyebabkan penentuan wilayah prioritas penanganan *stunting* menjadi kurang terarah dan berpotensi tidak tepat sasaran. Oleh karena itu, diperlukan pendekatan analitis yang lebih sistematis untuk mengolah data yang tersedia sehingga dapat menghasilkan informasi yang lebih bermakna dan mendukung pengambilan keputusan berbasis data [4].

Seiring dengan perkembangan teknologi, Dalam dunia pengolahan data, teknik *K-Means Clustering* hadir sebagai solusi untuk memilah-milah kumpulan informasi. Prinsip kerjanya sederhana: algoritma ini bakal mencari titik temu atau kesamaan karakteristik tertentu, lalu menyatukan data-data tersebut ke dalam kelompok yang sejenis [5].

Penelitian ini bertujuan untuk menganalisis tren prevalensi *stunting* pada balita di Provinsi Jawa Barat serta mengelompokkan kabupaten/kota berdasarkan

pola perubahan selama periode 2014–2024. Pendekatan ini diharapkan mampu menghasilkan segmentasi wilayah yang lebih terstruktur, sehingga dapat mendukung penentuan prioritas penanganan *stunting* secara lebih efektif dan berbasis data [6].

2. TINJAUAN PUSTAKA

Sebagai dasar kerangka berpikir, tinjauan ini memuat landasan teori dan kajian terkait yang mendukung analisis prevalensi *stunting* serta penerapan metode data mining. Pembahasan difokuskan pada konsep *stunting* sebagai objek penelitian, metode *clustering* sebagai pendekatan analisis, serta teknik evaluasi yang digunakan untuk menilai kualitas hasil pengelompokan [7].

Selain itu, tinjauan pustaka juga mencakup pemanfaatan data terbuka sebagai sumber informasi dalam analisis kesehatan masyarakat. Penggunaan data sekunder yang terstruktur memungkinkan dilakukan pengolahan dan analisis secara lebih sistematis untuk mengidentifikasi pola serta dinamika perubahan tahunannya [8].

Dengan mengintegrasikan konsep *stunting*, teknik *data mining*, dan evaluasi *clustering*, diharapkan penelitian ini memiliki landasan analitis yang kuat sehingga hasil yang diperoleh dapat dipertanggungjawabkan secara ilmiah [9].

Beberapa penelitian terdahulu menunjukkan bahwa metode *K-Means Clustering* efektif digunakan dalam analisis *stunting*. Penelitian oleh Amalia dan Arianto menerapkan metode *K-Means* dengan evaluasi *Silhouette Score* dan berhasil mengidentifikasi tiga kluster dengan karakteristik risiko *stunting* yang berbeda [10].

Sementara itu, Fadilah et al. menggunakan metode *K-Means* dengan pendekatan *Elbow Method*

dan menghasilkan dua kluster utama, yaitu wilayah dengan risiko tinggi dan rendah [11].

Penelitian lain oleh Syalom et al. juga menggunakan *K-Means* dengan evaluasi *Silhouette Score* dan menunjukkan bahwa wilayah dengan tingkat kerawanan *stunting* berkaitan dengan rendahnya cakupan layanan kesehatan esensial. Berdasarkan penelitian-penelitian tersebut, dapat disimpulkan bahwa metode *clustering*, khususnya *K-Means*, memiliki kemampuan yang baik dalam mengelompokkan wilayah berdasarkan karakteristik *stunting*, sehingga relevan untuk digunakan dalam penelitian ini [12].

2.1. Stunting

Stunting dan malnutrisi merupakan dua kondisi yang saling berkaitan. Periode 1000 hari pertama kehidupan, termasuk masa janin, adalah fase krusial di mana gangguan nutrisi kronis bisa memicu terjadinya *stunting* pada balita [4].

Kondisi ini umumnya belum tampak saat bayi dilahirkan, tetapi mulai teridentifikasi ketika anak memasuki usia dua tahun. Dampak *stunting* meliputi keterlambatan perkembangan motorik dan bahasa, yang berdampak pada degradasi kapasitas intelektual, eskalasi risiko morbiditas, serta keterbatasan daya saing individu saat memasuki fase usia produktif [13].

Efektivitas penyerapan gizi yang terganggu akibat infeksi serta rendahnya konsumsi makanan merupakan determinan utama yang memicu *stunting* secara langsung. Namun, kondisi ini juga sangat bergantung pada faktor-faktor eksternal seperti tingkat ketersediaan bahan pangan di lingkungan keluarga dan kualitas pola asuh dalam memberikan perawatan kesehatan serta asupan nutrisi bagi anak [14].

2.2. K-Means

Proses pengelompokan data ke dalam berbagai sub-set dilakukan melalui mekanisme algoritma *K-Means* atau kluster melalui metode pengelompokan data nonhierarkis [15].

Secara dasar, tugas utama algoritma *K-Means* adalah mengelompokkan data yang memiliki karakteristik dan fitur umum tertentu ke dalam kategori data yang terpisah dan berbeda. Pertama, pusat kluster (*centroids*) acak untuk setiap kluster data yang akan dibentuk dipilih dari data yang tersedia. Proses perhitungan diulang untuk menghitung posisi terbaik dari pusat kluster tersebut [16].

$$J = \sum_{j=1}^k \sum_{i \in S_j} \|x_i - \mu_j\|^2 \quad (1)$$

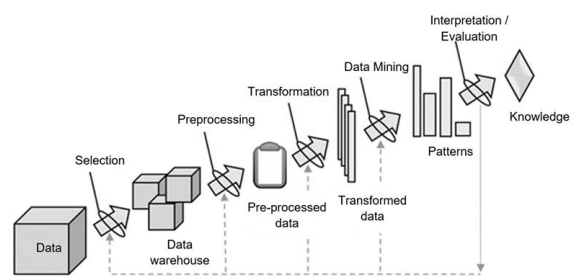
Berbeda dengan metode lain, setiap titik data dalam *K-Means* harus termasuk dalam suatu kluster pada satu tahap proses, tetapi dapat termasuk dalam kluster lain pada tahap-tahap berikutnya [17].

Secara esensial, penerapan *K-Means* dalam teknik pengelompokan bergantung tergantung pada variabel data yang diolah serta target hasil yang ingin dicapai. Korelasi dan kemiripan karakteristik intrinsik

pada atribut data akan mendikte bagaimana algoritma *K-Means* mengonfigurasi pengelompokan data secara sistematis [18].

2.3. Data Mining

Transformasi kumpulan data masif menjadi wawasan yang aplikatif dilakukan melalui proses *data mining*, dengan fokus utama pada penemuan pola-pola unik dan hubungan korelasi yang tidak terlihat secara manual. *Data mining* pada dasarnya adalah proses menganalisis dan mengeksplorasi data. *Data mining* membantu mengidentifikasi pola dalam data, pengelompokan data, dan membuat prediksi untuk mendukung pengambilan keputusan dengan menggunakan berbagai teknik. Teknik-teknik yang terlibat mencakup tahap-tahap seperti pemodelan, evaluasi, persiapan data, dan interpretasi hasil [19].



Gambar 1. KDD Methodology Process

Salah satu cara untuk mengakses informasi dari basis data adalah melalui *data mining*. *Data mining* umumnya melibatkan teknik seperti *machine learning* dan kecerdasan buatan untuk mengekstrak data dari basis data yang sangat besar [20].

Banyak industri seperti bisnis, penelitian, kesehatan, keuangan, dan beberapa industri lainnya memanfaatkan penambangan data untuk membantu organisasi memanfaatkan informasi yang terkandung dalam data mereka dengan lebih baik [21].

2.4. Normalisasi Data

Upaya meminimalisir bias akibat perbedaan skala pengukuran pada dataset dilakukan melalui proses normalisasi data. Sebagai bagian integral dari tahapan *data preprocessing*, teknik ini memastikan bahwa setiap fitur memiliki kontribusi yang proporsional dalam pemodelan statistik tanpa terdistorsi oleh perbedaan satuan atau rentang nilai. Normalisasi data dapat dijelaskan sebagai mengubah skala data agar lebih mudah diproses. Selain itu, normalisasi data memastikan tidak memerlukan ruang penyimpanan tambahan maupun perhitungan CPU yang lebih besar [22].

Min-Max Normalization adalah metode yang digunakan dalam penormalan data. Metode penormalan skala desimal, yang merupakan transformasi data yang melibatkan penormalan untuk membuat nilai-nilai memiliki pentingnya yang sama dengan menyelaraskan semua nilai data ke skala yang diinginkan masing-masing berdasarkan posisi

desimalnya, dalam kerangka prosedur penormalan data untuk mereduksi disparitas magnitudo antarvariabel[23].

$$nv = f(v) = \frac{v - \min(v)}{\max(v) - \min(v)} \quad (2)$$

Tabel 1. Penelitian Terdahulu

Judul	Metode	Hasil
Implementasi Algoritma <i>K-Means Clustering</i> dalam Klasterisasi Kabupaten/Kota Provinsi Jawa Barat berdasarkan Faktor Pemicu <i>Stunting</i> pada Balita (Amalia & Arianto, 2024)	<i>K-Means</i> + <i>Silhouette</i>	Identifikasi tiga kluster dengan karakteristik risiko berbeda; nilai <i>silhouette</i> menunjukkan kluster cukup baik.
Pengelompokan Kabupaten/Kota di Indonesia Berdasarkan Faktor Penyebab <i>Stunting</i> pada Balita menggunakan Algoritma <i>K-Means</i> (Fadilah et al., 2022)	<i>K-Means</i> + <i>Elbow</i>	Menghasilkan dua kluster, wilayah risiko tinggi dan risiko rendah.
<i>Clustering</i> Daerah Rawan <i>Stunting</i> di Indonesia menggunakan Algoritma <i>K-Means</i> (Syalom et al., 2025)	<i>K-Means</i> + <i>Silhouette</i>	Mengungkap kluster rawan <i>stunting</i> yang berkaitan dengan rendahnya cakupan layanan esensial.

3. METODE PENELITIAN

Melalui pemanfaatan data sekunder dari portal Open Data Jabar (2014-2024), penelitian ini menggunakan instrumen *data mining* dalam kerangka penelitian kuantitatif. Fokus utamanya adalah membedah tren prevalensi *stunting* balita di tingkat provinsi dengan tingkat akurasi yang tinggi. Pendekatan berbasis data ini dipilih untuk menyingkap keterkaitan antar-variabel wilayah di Jawa Barat, sehingga pola perubahan dan kesamaan karakteristik daerah dapat dipetakan secara representatif.

Identifikasi karakteristik wilayah dan penetapan prioritas penanggulangan *stunting* dilakukan melalui analisis hasil *clustering* yang dihasilkan oleh algoritma *K-Means*. Prosedur ini diawali dengan prapemrosesan dan sinkronisasi rentang data, yang kemudian dilanjutkan dengan pemodelan menggunakan *Elbow Method* dan *Silhouette Score* untuk menjamin validitas jumlah kelompok. Tahapan yang terintegrasi ini memastikan bahwa setiap kluster yang terbentuk merepresentasikan kondisi faktual di lapangan secara sistematis.

3.1. Objek Penelitian

Pemanfaatan data sekunder dari portal Open Data Jabar dalam penelitian ini diarahkan untuk membedah tren prevalensi *stunting* di tingkat kabupaten/kota se-Jawa Barat. Pendekatan ini memungkinkan dilakukannya pemetaan yang komprehensif terhadap capaian indikator kesehatan balita di tiap wilayah administratif. Data dipilih karena bersifat terbuka, terstruktur, dan menyediakan informasi historis yang relevan untuk mengkaji pola perubahan prevalensi selama periode 2014-2024. Pendekatan ini memungkinkan dilakukan analisis spasial-temporal guna memahami dinamika *stunting* secara lebih komprehensif antarwilayah.

Objek penelitian mencakup data jumlah balita *stunting* per wilayah yang dianalisis untuk mengidentifikasi pola dan karakteristik berdasarkan tren tahunan. Proses penelitian meliputi pengumpulan

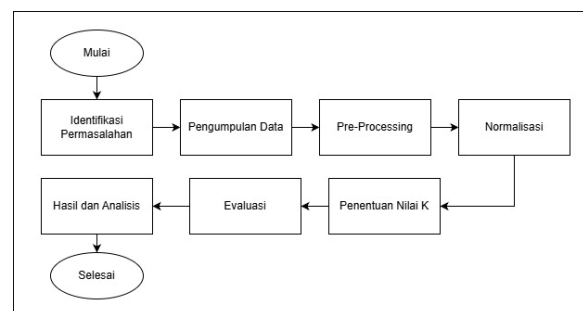
2.5. Penelitian Terdahulu

Untuk memperkuat landasan teoretis dan menemukan kelemahan dalam penelitian ini, penelitian sebelumnya digunakan sebagai acuan.

data, *pre-processing*, normalisasi, pemodelan menggunakan *K-Means Clustering*, serta evaluasi dan interpretasi kluster. Seluruh tahapan dilaksanakan menggunakan Python melalui platform Google Colab dalam rentang waktu Desember hingga Maret 2026.

3.2. Prosedur Penelitian

Bagian ini menjelaskan desain alur penelitian yang dilaksanakan secara sistematis, dimulai dari tahap identifikasi permasalahan, pengumpulan data, *pre-processing*, dan normalisasi data. Selanjutnya dilakukan penentuan nilai K, evaluasi model, serta analisis hasil pengelompokan hingga diperoleh kesimpulan penelitian. Gambar 2 menyajikan diagram alur prosedur penelitian, sedangkan uraian rinci setiap tahapan dijelaskan pada subbagian berikutnya.



Gambar 2. Prosedur Penelitian

3.3. Identifikasi Permasalahan

Tahap ini dilakukan untuk merumuskan permasalahan penelitian serta menetapkan tujuan analisis yang akan dicapai. Peneliti mengkaji kondisi objek penelitian dan menentukan kebutuhan pengelompokan data sebagai pendekatan penyelesaian masalah. Hasil identifikasi digunakan sebagai dasar dalam menyusun tahapan pengolahan data selanjutnya.

3.4. Pengambilan Data

Data penelitian diperoleh dari portal resmi Open Data Jabar pada dataset prevalensi *stunting* balita

tingkat kabupaten/kota. Data yang digunakan mencakup periode 2014–2024 dengan cakupan 27 kabupaten/kota. Dataset diunduh dalam format CSV/Excel melalui fitur unduh pada portal, kemudian digabungkan menjadi satu file induk. Total data yang terkumpul terdiri dari (jumlah baris $\times$ jumlah variabel) yang selanjutnya digunakan sebagai data awal penelitian.

3.5. Pre-processing Data

Data awal diperiksa untuk mengidentifikasi duplikasi, nilai kosong, dan ketidakkonsistenan format. *Record* duplikat dihapus, nilai kosong ditangani melalui imputasi atau penghapusan baris, dan tipe data diseragamkan. Selanjutnya dilakukan transformasi struktur data sehingga setiap baris merepresentasikan satu kabupaten/kota dengan atribut prevalensi per tahun. Hasil tahap ini menghasilkan dataset bersih yang siap dianalisis.

3.6. Normalisasi Data

Dataset yang telah dibersihkan selanjutnya dinormalisasi untuk menyamakan rentang nilai antar elemen numerik dalam dataset untuk meniadakan ketimpangan rentang nilai selama ekstraksi kluster. Penyetaraan ini dilakukan dengan menerapkan metode normalisasi *Min-Max*, *Z-Score*, atau *Decimal Scaling* sesuai karakteristik data. Proses ini menghasilkan data dengan skala yang lebih seragam sehingga perhitungan jarak pada algoritma *K-Means* menjadi lebih akurat dan stabil.

3.7. Pemodelan K-Means Clustering

Implementasi algoritma *K-Means* dilakukan dengan memanfaatkan hasil normalisasi data sebagai fondasi analitis untuk menghasilkan pengelompokan yang akurat dan objektif. Pengujian dilakukan pada beberapa nilai jumlah kluster (k) untuk menemukan konfigurasi pengelompokan yang paling representatif. Proses pemodelan dijalankan menggunakan perangkat lunak analisis data, dan setiap kabupaten/kota diberikan label kluster sesuai kedekatan karakteristik tren prevalensi *stunting*.

3.8. Evaluasi Model

Guna menjamin reliabilitas hasil pengelompokan, dilakukan uji validitas menggunakan koefisien *Silhouette Score*. Metrik ini berfungsi untuk menganalisis kepadatan data dalam satu kluster dibandingkan dengan jarak antar-kluster lainnya. Nilai indeks yang paling mendekati angka satu diidentifikasi sebagai konfigurasi model yang paling akurat dalam merepresentasikan distribusi data secara proporsional.

3.9. Interpretasi Kluster

Setiap kluster yang terbentuk dianalisis dengan menghitung rata-rata dan pola tren prevalensi *stunting* pada masing-masing kelompok. Analisis ini digunakan untuk mengidentifikasi karakteristik umum, seperti wilayah dengan tren tinggi, sedang, atau

rendah, sehingga diperoleh gambaran kondisi tiap kelompok secara lebih jelas dan informatif.

3.10. Kesimpulan

Kesimpulan disusun berdasarkan hasil evaluasi dan interpretasi kluster untuk menjawab rumusan masalah penelitian. Temuan utama dirangkum secara sistematis serta dilengkapi implikasi praktis yang dapat dimanfaatkan sebagai dasar pertimbangan dalam perencanaan kebijakan berbasis data.

4. HASIL DAN PEMBAHASAN

Interpretasi terhadap hasil pengelompokan data dan dinamika prevalensi yang ditemukan melalui pendekatan *K-Means Clustering* akan dijabarkan secara sistematis dalam bagian pembahasan ini. Adapun hasil yang ditampilkan meliputi tahap pengolahan data, penentuan jumlah kluster, proses *clustering*, visualisasi, serta interpretasi hasil untuk mendukung penentuan wilayah prioritas penanganan *stunting* di Provinsi Jawa Barat.

4.1. Dataset

Pemanfaatan data jumlah balita *stunting* di Jawa Barat pada skala kabupaten/kota ini merujuk pada basis data Open Data Jabar periode pengamatan terkait. Data mencakup periode tahun 2014 hingga 2024 dan memiliki struktur time series yang memungkinkan analisis tren secara longitudinal. Setelah dilakukan proses *pre-processing*, data diubah ke dalam bentuk matriks dengan setiap baris merepresentasikan wilayah dan setiap kolom merepresentasikan tahun pengamatan.

nama_kabupaten_kota	KABUPATEN BANDUNG	KABUPATEN BANDUNG BARAT	KABUPATEN BEKASI	KABUPATEN BOGOR	KABUPATEN CIAMIS	KABUPATEN CIANJUR
tahun						
2014	46412.0	20233.0	23473.0	40314.0	2929.0	26687.0
2015	16357.0	14809.0	20699.0	33721.0	2236.0	26582.0
2016	20391.0	10734.0	10843.0	32035.0	4092.0	26666.0
2017	12272.0	10487.0	3861.0	23797.0	5831.0	27625.0
2018	21131.0	7578.0	6045.0	7642.0	4847.0	12476.0

Gambar 3. Dataset

4.2. Normalisasi Data

nama_kabupaten_kota	KABUPATEN BANDUNG	KABUPATEN BANDUNG BARAT	KABUPATEN BEKASI	KABUPATEN BOGOR	KABUPATEN CIAMIS
tahun					
2014	1.000000	0.435943	0.505753	0.868612	0.063109
2015	0.470819	0.423643	0.603145	1.000000	0.040472
2016	0.622940	0.310223	0.313753	1.000000	0.095139
2017	0.423513	0.356488	0.107690	0.856263	0.181661
2018	0.822380	0.278607	0.217100	0.281175	0.169034

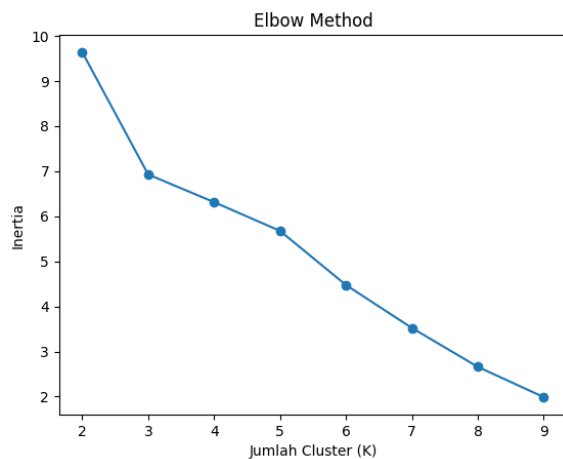
Gambar 4. Normalisasi Data

Guna menjamin stabilitas hasil pengelompokan, teknik normalisasi diterapkan untuk menyinkronkan interval nilai pada setiap variabel dataset. Ketimpangan skala antartahun ditekan sedemikian

rupa agar tidak memonopoli komputasi jarak dalam algoritma *K-Means*. Penyetaraan rentang nilai ini memastikan setiap atribut berperan setara dalam identifikasi pola, sehingga luaran klastering dapat mencerminkan kondisi lapangan secara akurat.

4.3. Penentuan Jumlah Klaster

Melalui pendekatan *Elbow Method*, titik optimal jumlah klaster diidentifikasi berdasarkan penurunan drastis nilai *inertia* yang kemudian melandai secara konsisten. Validitas dari temuan tersebut diuji kembali menggunakan instrumen *Silhouette Score* untuk memastikan reliabilitas pemisahan antarkelompok. Berdasarkan hasil integrasi kedua metode tersebut, ditetapkan bahwa $k=3$ merupakan jumlah klaster paling ideal karena mampu menyeimbangkan aspek akurasi statistik dengan kepentingan interpretasi sosiogeografis wilayah.



Gambar 5. Penentuan Nilai K

4.4. Hasil Clustering

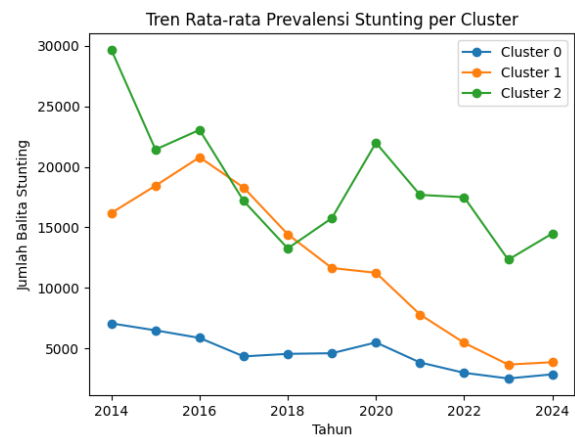
Tiga kelompok wilayah dengan pola perkembangan *stunting* yang homogen telah diidentifikasi melalui mekanisme *K-Means Clustering*. Penentuan klaster ini didasarkan pada kemiripan atribut tren selama rentang waktu pengamatan, di mana setiap klaster menyajikan profil unik yang membedakannya dari kelompok lain. Struktur pengelompokan yang dihasilkan menjadi instrumen krusial dalam memandu analisis lanjutan guna memahami urgensi penanganan di setiap klaster secara objektif.

nama_kabupaten_kota	tahun	Cluster	Kategori
KABUPATEN BANDUNG	2		Tinggi
KABUPATEN BANDUNG BARAT	1		Sedang
KABUPATEN BEKASI	0		Rendah
KABUPATEN BOGOR	2		Tinggi
KABUPATEN CIAMIS	0		Rendah

Gambar 6. Hasil Clustering

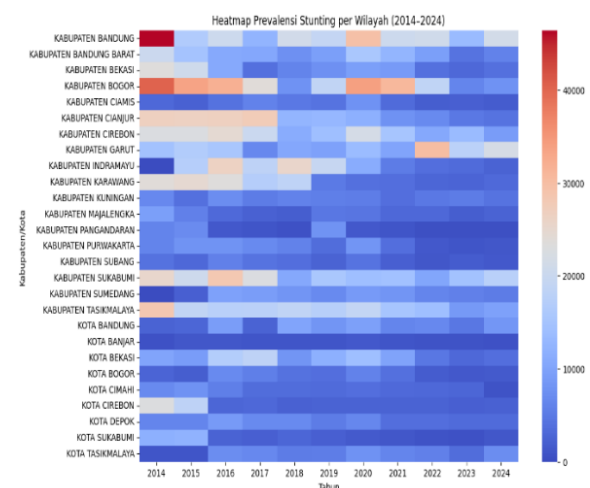
4.5. Tren Stunting per Klaster

Visualisasi tren *stunting* per klaster digunakan untuk menggambarkan pola perubahan rata-rata jumlah balita *stunting* pada masing-masing kelompok wilayah selama periode 2014–2024. Grafik yang dihasilkan menunjukkan adanya perbedaan yang cukup jelas antar klaster, baik dari segi tingkat prevalensi maupun kecenderungan perubahannya dari waktu ke waktu. Setiap klaster memiliki karakteristik tren yang berbeda, di mana terdapat kelompok wilayah dengan tingkat prevalensi yang relatif tinggi dan cenderung fluktuatif, sementara kelompok lainnya menunjukkan pola yang lebih stabil atau mengalami penurunan secara bertahap. Guna mempertegas signifikansi hasil pengelompokan, visualisasi ini digunakan untuk memetakan fluktuasi *stunting* secara lebih sistematis, sehingga karakteristik unik pada setiap klaster dapat teridentifikasi dengan lebih presisi.



Gambar 7. Tren Stunting per Klaster

4.6. Heatmap Distribusi Stunting



Gambar 8. Heatmap Distribusi Stunting

Guna meninjau dinamika sebaran *stunting* secara sistematis, *heatmap* digunakan dengan memanfaatkan transisi intensitas warna sebagai parameter tingkat nilai data. Pendekatan grafis ini memungkinkan deteksi dini terhadap wilayah dengan anomali jumlah

kasus tinggi maupun rendah pada tiap entitas kabupaten/kota selama periode pengamatan. Visualisasi ini memudahkan dalam mengidentifikasi perbedaan tingkat prevalensi antarwilayah dan antar waktu secara simultan dalam satu tampilan. Hasil *heatmap* menunjukkan adanya variasi distribusi yang cukup signifikan, di mana beberapa wilayah terlihat memiliki intensitas warna yang lebih tinggi secara konsisten, sementara wilayah lainnya berada pada tingkat yang lebih rendah. Pola tersebut mengindikasikan bahwa distribusi *stunting* di Provinsi Jawa Barat tidak merata dan memiliki dinamika yang berbeda-beda pada setiap wilayah sepanjang periode penelitian.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil analisis komputasi yang telah dilakukan, dapat disimpulkan bahwa segmentasi wilayah melalui *K-Means* mampu menyingkap pola distribusi *stunting* di Jawa Barat yang sebelumnya tidak teramati secara manual. Integrasi berbagai tahapan analitis telah menghasilkan klasifikasi tiga kluster yang mencerminkan derajat kerentanan wilayah yang berbeda sepanjang periode 2014-2024. Visualisasi yang dihasilkan, baik melalui grafik tren maupun *heatmap*, memberikan pemahaman visual yang mendalam mengenai disparitas kesehatan antarwilayah serta dinamika perubahan tahunannya. Struktur data yang terorganisir ini diharapkan dapat menjadi rujukan fundamental dalam perumusan skala prioritas penanganan *stunting* yang lebih efisien. Sebagai langkah lanjutan, penelitian ini merekomendasikan adanya integrasi data makro ekonomi dan infrastruktur kesehatan sebagai variabel pendukung, serta penggunaan teknik validasi silang dengan metode *unsupervised learning* lainnya guna memperkuat validitas ilmiah dan kedalaman informasi dalam studi berikutnya.

DAFTAR PUSTAKA

- [1] S. Anwar *Et Al.*, “Systematic Review Faktor Risiko, Penyebab Dan Dampak Stunting Pada Anak (Systematic Review Risk Factors, Causes And Impact Of Stunting In Children),” *Jurnal Ilmu Kesehatan*, Vol. 11, No. 1, 2022.
- [2] H. Aulia Qodrina And R. Kurnia Sinuraya, “Faktor Langsung Dan Tidak Langsung Penyebab Stunting Di Wilayah Asia: Sebuah Review,” *Jurnal Penelitian Kesehatan Suara Forikes*, Vol. 12, No. 4, 2021, Doi: 10.33846/Sf12401.
- [3] S. Annisa Rizki Manaf, A. Fitrianto, R. Amelia, And I. Pertanian Bogor Jl Raya Dramaga Kampus Ipb Dramaga Bogor, “Faktor-Faktor Yang Mempengaruhi Permasalahan Stunting Di Jawa Barat Menggunakan Regresi Logistik Biner,” *J Statistika*, Vol. 15, No. 2, 2022. [Online]. Available: Www.Unipasby.Ac.Idtelp./Fax.
- [4] Taufik Hidayat, Mohamad Jajuli, And Susilawati, “Clustering Daerah Rawan Stunting Di Jawa Barat Menggunakan Algoritma K-Means,” *Infotech: Jurnal Informatika & Teknologi*, Vol. 4, No. 2, Pp. 137–146, Dec. 2023, Doi: 10.37373/Infotech.V4i2.642.
- [5] A. F. Afra, A. L. Hananto, A. Hananto, And B. Priyatna, “Klasterisasi Supplier Berdasarkan Kinerja Menggunakan Algoritma K-Means,” *Jurnal Teknologi Dan Sistem Informasi Bisnis*, Vol. 7, No. 2, Pp. 334–341, Apr. 2025, Doi: 10.47233/Jteksis.V7i2.1935.
- [6] A. Riznawati, D. Yudhistira, M. Rahmaniati, T. Sipahutar, And T. Eryando, “Autokorelasi Spasial Prevalensi Stunting Di Jawa Barat Tahun 2021,” *Jurnal Biostatistik, Kependudukan, Dan Informatika Kesehatan*, Vol. 3, No. 1, Nov. 2022, Doi: 10.7454/Bikfokes.V3i1.1035.
- [7] A. Adityaningrum *Et Al.*, “Faktor Penyebab Stunting Di Indonesia: Analisis Data Sekunder Data Ssgi Tahun 2021 Factors Causing Stunting In Indonesia: 2021 Ssgi Secondary Data Analysis,” *Jambura Journal Of Epidemiology*, Vol. 3, No. 1, Pp. 1–10, 2021, Doi: 10.56796/Jje.V2i1.21542.
- [8] S. Suraya, M. Sholeh, And U. Lestari, “Evaluation Of Data Clustering Accuracy Using K-Means Algorithm,” *International Journal Of Multidisciplinary Approach Research And Science*, Vol. 2, No. 01, Pp. 385–396, Dec. 2023, Doi: 10.59653/Ijmars.V2i01.504.
- [9] H. Munawaroh, N. Khoirun Nada, And A. Hasjiandito, “Peranan Orang Tua Dalam Pemenuhan Gizi Seimbang Sebagai Upaya Pencegahan Stunting Pada Anak Usia 4-5 Tahun,” *Sentra Cendekia*, Vol. 3, No. 2, 2022. [Online]. Available: [Http://E-Journal.Ivet.Ac.Id/Index.Php/Sc](http://E-Journal.Ivet.Ac.Id/Index.Php/Sc)
- [10] M. F. Amalia And D. B. Arianto, “Implementasi Algoritma K-Means Clustering Dalam Klasterisasi Kabupaten/Kota Provinsi Jawa Barat Berdasarkan Faktor Pemicu Stunting Pada Balita,” *Simkom*, Vol. 9, No. 1, Pp. 36–46, Jan. 2024, Doi: 10.51717/Simkom.V9i1.356.
- [11] A. Fadilah, M. Nurfaizy P, S. Lumbanbatu, And S. Defiyanti, “Pengelompokan Kabupaten/Kota Di Indonesia Berdasarkan Faktor Penyebab Stunting Pada Balita Menggunakan Algoritma K-Means,” *Jurnal Informatika Dan Komputer*, Vol. 6, No. 2, Pp. 223–230, 2022.
- [12] C. Syalom, U. Khaira, And D. Arsa, “Clustering Daerah Rawan Stunting Di Indonesia Menggunakan Algoritma K-Means,” *Jurnal Teknologi Informasi*, Vol. 6, No. 1, 2025, Doi: 10.46576/Djtechno.
- [13] N. A. Latifah, F. Fajrini, N. Romdhona, D. Herdiansyah, Erniyasih, And Suherman, “Systematic Literature Review: Stunting Pada Balita Di Indonesia Dan Faktor Yang Mempengaruhinya,” *Jurnal Kedokteran Dan*

- Kesehatan, Vol. 20, No. 1, 2024. [Online]. Available: <https://jurnal.umj.ac.id/index.php/jkk>
- [14] N. Rusliani, W. R. Hidayani, And H. Sulistyoningih, "Literature Review: Faktor-Faktor Yang Berhubungan Dengan Kejadian Stunting Pada Balita," *Buletin Ilmu Kebidanan Dan Keperawatan*, Vol. 1, No. 01, Pp. 32–40, Aug. 2022, Doi: 10.56741/Bikk.V1i01.39.
- [15] N. F. Rahman *Et Al.*, "Penerapan Algoritma K-Means Untuk Klasterisasi Film Pada Platform Amazon Prime Video," *Jati (Jurnal Mahasiswa Teknik Informatika)*, Vol. 9 No. 5, 2025.
- [16] F. Salsabila, T. Ridwan, And H. H, "Analisa Volume Penyebaran Sampah Di Karawang Menggunakan Algoritma K-Means Clustering," *Jurnal Informatika Dan Teknik Elektro Terapan*, Vol. 12, No. 2, Apr. 2024, Doi: 10.23960/Jitet.V12i2.4226.
- [17] I. Ferdiansyah, B. Huda, A. Hananto, And U. Buana Perjuangan Karawang, "Analisis Clustering Menggunakan Metode K-Means Pada Kemiskinan Di Jawa Timur Tahun 2020," *Innovative: Journal Of Social Science Research*, Vol. 4, No. 3, 2024.
- [18] H. Jurnal, Ruf Ubaidillah, And Z. Fatah, "Implementasi Rapidminer Pada Klasterisasi Gempa Bumi Di Indonesia Berdasarkan Kedalaman Menggunakan K-Means," *Jurnal Ilmiah Multidisiplin Ilmu*, Vol 1, No. 6, 2024, Doi: 10.69714/W0m9zv32.
- [19] N. Hendrastuty, "Penerapan Data Mining Menggunakan Algoritma K-Means Clustering Dalam Evaluasi Hasil Pembelajaran Siswa," *Jurnal Ilmiah Informatika Dan Ilmu Komputer (Jima-Ilkom)*, Vol. 3, No. 1, Pp. 46–56, Mar. 2024, Doi: 10.58602/Jima-Ilkom.V3i1.26.
- [20] I. Julia, B. Priyatna, And S. Shofia Hilabi, "Penerapan Metode K-Means Clustering Untuk Menentukan Jumlah Penjualan Terlaris Pada Cv. Equipment & Tools," *Jutisi: Jurnal Ilmiah Teknik Informatika Dan Sistem Informasi*, Vol. 13, No. 1, 2024.
- [21] Dodi Nofri Yoliandi, "Data Mining Dalam Analisis Tingkat Penjualan Barang Elektronik Menggunakan Algoritma K-Means," *Insearch (Information System Research) Journal*, Vol 3, No. 1, 2023.
- [22] J. Homepage, I. Permana, And F. Nur Salisah, "Pengaruh Normalisasi Data Terhadap Performa Hasil Klasifikasi Algoritma Backpropagation," *Ijirse (Indonesian Journal Of Informatic Research And Software Engineering)*, Vol.2, No.1, 2022.
- [23] M. R. Kusnaldi, T. Gulo, And S. Aripin, "Penerapan Normalisasi Data Dalam Mengelompokkan Data Mahasiswa Dengan Menggunakan Metode K-Means Untuk Menentukan Prioritas Bantuan Uang Kuliah Tunggal," *Journal Of Computer System And Informatics (Josyc)*, Vol. 3, No. 4, Pp. 330–338, Sep. 2022, Doi: 10.47065/Josyc.V3i4.2112.