

ANALISIS CLUSTERING HARGA PANGAN DI PASAR TRADISIONAL WILAYAH KALIMANTAN MENGGUNAKAN ALGORITMA K-MEANS DAN FUZZY C-MEANS

Emmanuel Garcia Sumargo, Teny Handhayani

Teknik Informatika, Universitas Tarumanagara
Jln. Letjen S. Parman No. 1, Jakarta, 11440, Indonesia
tenyh@fti.untar.ac.id

ABSTRAK

Kalimantan, sebagai salah satu pulau vital di Indonesia, memegang peran penting dalam ketahanan pangan nasional. Namun, stabilitas harga pangan di wilayah ini menghadapi tantangan signifikan berupa fluktuasi harga yang tinggi dan disparitas antar wilayah, yang berpotensi mengganggu stabilitas ekonomi dan sosial. Pemahaman mendalam mengenai pola kewilayahan harga menjadi krusial untuk mitigasi risiko tersebut. Oleh karena itu, penelitian ini bertujuan untuk mengkluster kota-kota di Kalimantan berdasarkan karakteristik data harga pangan menggunakan algoritma K-Means dan Fuzzy C-Means. Metode penelitian melibatkan analisis data *time series* harian dari 10 kota/kabupaten dari 1 Januari 2018 hingga 31 Desember 2023, yang mencakup lima komoditas utama: beras, telur ayam, daging ayam, bawang putih, dan bawang merah. Hasil evaluasi menggunakan *Silhouette Score* dan *Davies-Bouldin Index* menunjukkan algoritma K-Means lebih unggul dengan membentuk dua *cluster* optimal. *Cluster* pertama (Kabupaten Sintang, Kota Palangkaraya, Kota Tarakan) menunjukkan karakteristik harga yang lebih tinggi dengan fluktuasi lebih besar, dan secara geografis terkonsentrasi di wilayah tengah dan utara. *Cluster* kedua, yang mencakup tujuh kota lainnya, cenderung memiliki harga yang lebih rendah dan stabil. Penelitian ini diharapkan membantu perencanaan kebijakan yang lebih baik dalam mengelola harga pangan di Kalimantan.

Kata kunci : *Clustering, K-Means, Fuzzy C-Means, Kalimantan, Harga Pangan*

1. PENDAHULUAN

Kalimantan merupakan salah satu pulau terbesar yang berada di Indonesia, dengan luas sekitar 736.000 kilometer persegi dan terdiri dari lima provinsi: Kalimantan Barat, Kalimantan Tengah, Kalimantan Selatan, Kalimantan Timur, dan Kalimantan Utara [1]. Pulau ini berbatasan dengan Malaysia dan Brunei Darussalam. Pada wilayah perbatasan seringkali terdapat daerah-daerah yang kurang maju, terutama dalam pembangunan sektor pertanian yang relatif rendah dibandingkan dengan daerah lain. Sektor pertanian memegang peran penting dalam ketahanan pangan suatu daerah, namun di kawasan perbatasan Kalimantan, pembangunan sektor ini belum optimal. Peningkatan pembangunan sektor pertanian di wilayah perbatasan ini sangat diperlukan untuk meningkatkan ketahanan pangan dan kesejahteraan masyarakat setempat [2]. Pangan adalah kebutuhan dasar manusia yang memiliki dampak signifikan terhadap stabilitas ekonomi dan sosial suatu negara. Stabilitas ekonomi dan sosial sangat bergantung pada ketersediaan pangan yang cukup, aman, bermutu, bergizi, dan terjangkau [3]. Di Indonesia, beberapa bahan pangan pokok meliputi beras, telur ayam, daging ayam, daging sapi, bawang putih, bawang merah, dan cabai [4] [5]. Namun, kenaikan harga pangan yang terjadi setiap tahun dapat mengakibatkan krisis pangan yang berdampak luas [6].

Krisis pangan merupakan suatu masalah global yang serius, hal ini tidak dapat diabaikan karena sangat berdampak pada kesehatan masyarakat [7]. Indonesia juga mengalami kesulitan dalam menghadapi krisis pangan, ditambah masalah krisis ekonomi yang terus

berlanjut membuat hal ini semakin berat [8][9]. Krisis ekonomi menyebabkan terjadinya inflasi, dimana inflasi akan menyebabkan fluktuasi harga pangan [9]. Analisis data harga pangan dari tahun-tahun sebelumnya dapat membantu dalam melakukan prediksi dan mengambil langkah pencegahan terhadap krisis pangan di masa depan [10]

Salah satu metode analisis yang efektif adalah *clustering*, yang mengelompokkan data ke dalam *cluster-cluster* di mana setiap anggota dalam *cluster* tertentu memiliki karakteristik yang mirip [11] [12] [13] [14] [13] [15]. Tujuan penelitian ini yaitu melakukan pengelompokan kota - kota di Kalimantan berdasarkan harga pangan yang ada di pasar tradisional. Pengelompokan dilakukan menggunakan algoritma *K-Means* dan *Fuzzy C-Means*. Dengan mengidentifikasi pola dan tren harga pangan serta mengelompokkan wilayah-wilayah berdasarkan karakteristik harga pangan masing-masing, diharapkan hasil dari analisis ini dapat memberikan kontribusi yang berarti dalam memahami dinamika harga pangan dan memperbaiki krisis pangan di wilayah Kalimantan [16].

2. TINJAUAN PUSTAKA

Penelitian terdahulu menunjukkan bahwa algoritma *clustering* telah efektif digunakan untuk analisis data regional di Indonesia. Algoritma K-Means terbukti andal dalam memetakan disparitas pada level provinsi, baik untuk analisis harga komoditas [17] maupun volume produksi pangan [18]. Di sisi lain, Fuzzy C-Means (FCM) menawarkan pendekatan yang lebih bernuansa untuk data sosio-

ekonomi yang kompleks. Deviyana [19] telah berhasil menerapkan FCM untuk pengelompokan indikator kesejahteraan di Kalimantan, sementara validitasnya untuk data pertanian dikonfirmasi dalam studi klasterisasi hasil panen di Bali [20].

Meskipun demikian, sintesis literatur mengidentifikasi beberapa celah penelitian yang signifikan. Pertama, analisis harga pangan yang ada cenderung terbatas pada skala makro (provinsi) dan komoditas tunggal, sehingga kurang memiliki granularitas spasial. Kedua, aplikasi FCM untuk menganalisis dinamika harga pangan multi-komoditas di Kalimantan belum dieksplorasi secara mendalam. Terakhir, belum ditemukan studi di Indonesia yang secara langsung membandingkan kinerja metode hard clustering (K-Means) dan soft clustering (FCM) pada dataset harga pangan yang identik. Oleh karena itu, penelitian ini bertujuan mengisi celah tersebut dengan melakukan analisis komparatif pada level kota/kabupaten di Kalimantan menggunakan data multi-komoditas.

2.1. K-Means

K-Means adalah algoritma *clustering* yang membagi data ke dalam K jumlah *cluster* dengan menentukan *centroid*, yaitu titik pusat yang mewakili setiap *cluster*. Proses dimulai dengan pemilihan *centroid* secara acak [21]. Kemudian, data-data akan dikelompokkan ke dalam *cluster* berdasarkan kedekatan jaraknya dengan *centroid* tersebut menggunakan rumus perhitungan jarak antar dua titik yaitu *euclidean distance*. Rumus *euclidean distance* pada persamaan (1) [22]:

$$distance(x, y) = \sqrt{\sum_{i=0}^n (X_i - Y_i)^2} \quad (1)$$

Dimana :

- $distance(x, y)$ adalah jarak antar titik data x dengan *centroid* y .
- X_i adalah nilai fitur ke- i dalam titik data x .
- Y_i adalah nilai dari fitur ke- i dalam *centroid* y .
- n adalah jumlah fitur dalam setiap titik data dan *centroid*.

Prosedur Algoritma *K-Means*:

- Tentukan banyak *cluster* sejumlah k .
- Inisialisasi *centroid* dengan memilih k titik secara acak sebagai titik pusat awal (*centroid*) dari setiap *cluster*.
- Setiap titik data dalam *dataset* akan dimasukkan ke *centroid* terdekat (*cluster* terdekat) berdasarkan metrik jarak, seperti *euclidean distance*.
- Menentukan *centroid* baru dengan menghitung ulang posisi *centroid* sebagai rata-rata dari semua titik data yang termasuk dalam *cluster* tersebut.
- Jika posisi *centroid* sudah stabil atau tidak ada perubahan anggota *cluster*, algoritma selesai. Jika tidak, kembali ke langkah 3.

2.2. Fuzzy C-Means

Fuzzy C-Means (FCM) adalah algoritma *clustering* yang memungkinkan setiap data untuk memiliki derajat keanggotaan terhadap setiap *cluster*. Derajat keanggotaan ini mencerminkan sejauh mana data tersebut berhubungan dengan setiap *cluster* dan nilainya berkisar antara 0 hingga 1. Dalam FCM, setiap data tidak hanya menjadi anggota satu *cluster*, tetapi dapat menjadi anggota beberapa *cluster* dengan nilai keanggotaan yang berbeda. Jika jumlah *cluster* yang telah ditentukan adalah c , maka setiap data akan memiliki c nilai keanggotaan [23].

Algoritma *Fuzzy C-Means* mempartisi sebuah himpunan $X = \{X_1, X_2, \dots, X_n\}$ menjadi c *cluster* berdasarkan kriteria tertentu. Setiap elemen X_i adalah sebuah vektor dengan d dimensi. *Cluster* diwakili oleh pusat-pusatnya $V = \{V_1, V_2, \dots, V_c\}$. Matriks U adalah representatif untuk keanggotaan setiap elemen dalam setiap *cluster*. Matriks U memiliki beberapa karakteristik sebagai berikut [24]:

- $U(i, k)$ adalah nilai keanggotaan elemen X_i dalam *cluster* dengan pusat V_k .
- $0 \leq U(i, k) \leq 1$ untuk semua i dan k , dan jumlah nilai keanggotaan untuk setiap elemen di semua *cluster* adalah $\sum_{k=1}^c U(i, k) = 1$
- Semakin besar $U(i, k)$, semakin tinggi tingkat kemungkinan bahwa X_i termasuk dalam *cluster* k .

Fungsi objektif $J(U, V)$:

$$J(U, V) = \sum_{i=1}^N \sum_{k=1}^c U(i, k)^m D(i, k)^2 \quad (2)$$

Dimana :

- $D(i, k)^2 = \|X_i - V_k\|^2$ adalah jarak kuadrat antara elemen X_i dan pusat *cluster* V_k .
- m adalah koefisien fuzzifikasi.

Langkah-langkah algoritma FCM:

- Inisialisasi
 - Inisialisasi pusat *cluster* V .
 - Tetapkan $l = 0$ (penghitung iterasi), $\epsilon > 0$ (kriteria penghentian), dan $m > 1$ (koefisien fuzzifikasi).
- Perbarui Matriks Keanggotaan
 - Untuk setiap elemen X_i dan setiap *cluster* k , perbarui U menggunakan:
$$U(i, k) = \left(\sum_{j=1}^c \left(\frac{D(i, k)}{D(i, j)} \right)^{\frac{2}{m-1}} \right)^{-1} \quad (3)$$
- Perbarui Pusat *Cluster* V

$$V(k) = \frac{\sum_{i=1}^N U(i, k)^m X_i}{\sum_{i=1}^N U(i, k)^m} \quad (4)$$
- Pemeriksaan Konvergensi
 - Jika $\|V^{(l)} - V^{(l+1)}\| < \epsilon$, hentikan algoritma.
 - Jika tidak, tetapkan $l = l + 1$ dan kembali ke langkah 2.
- Akhir
 - Algoritma berhasil dan mengembalikan *cluster-cluster* yang sudah dipartisi.

2.3. Silhouette Score

Silhouette score adalah sebuah metrik evaluasi yang digunakan untuk menilai seberapa baik sebuah titik data cocok dengan *cluster* yang telah ditentukan oleh algoritma *clustering*. Metrik ini menghitung koefisien *silhouette* untuk setiap titik data dengan mempertimbangkan jarak rata-rata intra-*cluster* (a) dan jarak rata-rata terdekat ke *cluster* lain (b). Koefisien *silhouette* untuk sebuah titik data didefinisikan sebagai $(b - a) / \max(a, b)$. Interpretasi dari nilai *silhouette score* adalah sebagai berikut [25]:

- Nilai dekat +1 menunjukkan bahwa titik data tersebut berada dalam *cluster* yang tepat.
- Nilai dekat 0 menandakan bahwa titik data mungkin lebih baik ditempatkan di *cluster* lain.
- Nilai dekat -1 menunjukkan bahwa titik data mungkin berada di *cluster* yang salah.

Rumus perhitungan *Silhouette Score* dapat dilihat pada persamaan (5) [26]:

$$S(i) = \frac{b(i) - a(i)}{\max(a(i), b(i))} \quad (5)$$

Dimana:

- $a(i)$: rata-rata jarak antara titik data i dengan semua titik data lain yang berada dalam *cluster* yang sama dengan i .
- $b(i)$: rata-rata jarak antara titik data i dengan semua titik data di *cluster* lain yang memiliki jarak terdekat dengan i .

2.4. Davies Bouldin Index

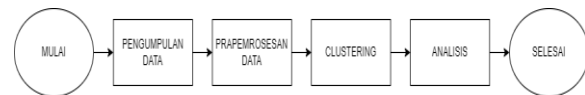
Indeks Davies-Bouldin (DBI) adalah metrik yang digunakan untuk menilai kinerja *clustering*. Ini menggabungkan dua aspek utama dalam mengevaluasi pengelompokan: meminimalkan varians intra-*cluster* dan memaksimalkan jarak antar-*cluster* [27][28]. DBI memiliki nilai minimum 0 dan semakin kecil nilainya maka semakin baik kinerja *clustering* [29]. DBI didefinisikan sebagai rata-rata dari nilai yang dihitung untuk setiap *cluster* dalam solusi pengelompokan. Perhitungan DBI dapat dilihat pada persamaan (6) [25]:

$$DB(k) = \frac{1}{k} \sum_{i=1}^k \max_{j \neq i} \left\{ \frac{S_i + S_j}{d(C_i, C_j)} \right\} \quad (6)$$

Di sini, k menyatakan jumlah *cluster*, S_i merupakan rata-rata jarak antara *centroid cluster* i dan semua anggotanya, dan $d(C_i, C_j)$ menyatakan jarak antara *centroid cluster* i dan j . DBI menunjukkan korelasi positif untuk kompakitas intra-*cluster* dan korelasi negatif untuk pemisahan inter-*cluster*. Metrik ini digunakan untuk mengevaluasi seberapa baik algoritma pengelompokan mengorganisir data ke dalam *cluster* yang secara internal koheren dan terpisah dengan baik [28].

3. METODE PENELITIAN

Penelitian ini berlangsung melalui serangkaian tahapan yang terstruktur dan dapat dilihat pada Gambar 1. Pada langkah pertama, dilakukan pengumpulan data. Dataset merupakan data harga pangan di wilayah Kalimantan. Langkah kedua adalah prapemrosesan data, yang mencakup pembersihan data, penanganan nilai yang hilang, dan normalisasi data untuk memastikan kualitas data yang optimal untuk analisis lebih lanjut. Langkah ketiga adalah proses *clustering*, di mana algoritma K-Means dan Fuzzy C-Means diterapkan untuk mengelompokkan data harga pangan ke dalam beberapa *cluster* berdasarkan kesamaan karakteristiknya. Setelah proses *clustering* selesai, langkah keempat adalah analisis hasil *clustering*. Pada tahap ini, hasil dari kedua algoritma dibandingkan dan dianalisis untuk mendapatkan pemahaman yang lebih mendalam tentang pola harga pangan di wilayah tersebut [29].



Gambar 1. Alur penelitian

3.1. Preprocessing Data

Data *preprocessing* adalah langkah penting yang dilakukan sebelum pengolahan data, dengan tujuan mengubah data mentah menjadi format yang ideal agar dapat diolah dengan lebih efisien pada tahap selanjutnya [30]. Pada tahap ini, data akan melalui proses pembersihan, integrasi, dan transformasi. Dalam proses pembersihan data, dilakukan penanganan nilai yang hilang, pada dataset yang digunakan untuk penelitian ini terdapat baris dengan nilai “-” pada kolomnya yang artinya nilai tidak diketahui, data-data tersebut diubah nilainya menjadi null, kemudian digunakan metode *forward fill* dan *backward fill* untuk mengisi nilai null tersebut. Bentuk data awal adalah data time series harian dari setiap kota yang akan di *clustering*. Data-data ini perlu di gabungkan dengan melakukan sinkronisasi data dari berbagai kota ke dalam satu matriks untuk setiap sumber data.

Selanjutnya, dilakukan proses integrasi data dari berbagai kota di wilayah Kalimantan, mencakup 10 kota yang berbeda. Setiap kota memiliki data harga pangan harian yang mencakup beberapa komoditas seperti beras, telur ayam, daging ayam, bawang merah, dan bawang putih. Data dari setiap kota kemudian disinkronkan agar memiliki *timestamp* yang sama sehingga dapat dianalisis secara bersamaan.

Setelah itu, dilakukan transformasi data dengan mengatur ulang semua fitur menjadi sebuah rangkaian. Matriks dari setiap kota yang berisi data *time series* harian diatur ulang menjadi rangkaian panjang. Semua urutan data dari berbagai kota kemudian diatur dalam sebuah matriks baru di mana baris-barisnya merepresentasikan kota-kota dan kolom-kolomnya adalah data harga pangan yang telah dinormalisasi.

Matriks baru ini kemudian digunakan sebagai input untuk proses *clustering* [16][29].

3.2. Clustering

Dalam bidang data mining, analisis clustering adalah salah satu topik penelitian utama [31]. *Clustering* merupakan metode untuk menemukan struktur *cluster* dalam dataset, di mana data dalam satu *cluster* memiliki karakteristik yang mirip, sementara data dalam *cluster* yang berbeda memiliki karakteristik yang berbeda secara signifikan [32]. Metode *clustering* berbeda dengan metode data mining lainnya karena dapat mengklasifikasikan data tanpa memerlukan pengetahuan awal tentang dataset. Oleh karena itu, *clustering* termasuk dalam pendekatan *unsupervised learning* dalam *machine learning* [31].

Machine learning memiliki dua pendekatan utama: *supervised learning* dan *unsupervised learning*. Dalam *supervised learning*, algoritma belajar dari data yang telah dilabeli, sedangkan dalam *unsupervised learning*, algoritma tidak memerlukan data berlabel dan dapat menemukan pola serta hubungan dalam data tersebut secara mandiri. Clustering, sebagai bagian dari *unsupervised learning*, memanfaatkan struktur dasar distribusi data untuk mengelompokkan data dengan karakteristik serupa. Ini menjadikannya sangat cocok untuk mengidentifikasi pola dan hubungan dalam data harga pangan dari berbagai kota di wilayah Kalimantan, di mana label atau kategori tidak tersedia. Dalam penelitian ini, algoritma yang akan digunakan untuk melakukan *clustering* adalah *K-Means* dan *Fuzzy C-Means* dan *Silhouette Score* dan *Davies-Bouldin Index* untuk mengukur kinerja *clustering* [33].

4. HASIL DAN PEMBAHASAN

Pada bab ini, disajikan hasil eksperimen yang dilakukan menggunakan *dataset* harga pangan dari 10 kota/kabupaten di wilayah Kalimantan. *Dataset* ini mencakup periode dari Januari 2018 hingga Desember 2023 dan merupakan data *time series* harian yang hanya mencakup hari Senin hingga Jumat, sementara data untuk hari Sabtu dan Minggu tidak tersedia. Data dikumpulkan melalui *website* Pusat Informasi Harga Pangan Strategis Nasional [34]. Tersedia data harga pangan dari 13 kota/kabupaten di Kalimantan yang dapat diunduh. Namun, tiga kota/kabupaten, yaitu Kota Bontang, Kabupaten Bulungan, dan Kabupaten Kotabaru, tidak memiliki data untuk tahun 2018 dan 2019, sehingga tidak disertakan dalam analisis ini. Sepuluh kota/kabupaten yang dianalisis adalah: Kabupaten Sintang, Kota Pontianak, Kota Singkawang, Kota Banjarmasin, Kota Tanjung, Kota Palangkaraya, Kota Sampit, Kota Balikpapan, Kota Samarinda, dan Kota Tarakan. Komoditas yang dianalisis dalam penelitian ini mencakup beras, telur ayam, daging ayam, bawang putih, dan bawang merah. Data harga dari komoditas-komoditas ini digunakan untuk mengevaluasi pola dan tren harga pangan di berbagai kota/kabupaten tersebut. *Clustering*

dilakukan pada *dataset* ini menggunakan dua algoritma, yaitu *K-Means* dan *Fuzzy C-Means*. Penelitian ini bertujuan untuk menilai dan membandingkan kinerja kedua algoritma guna menentukan algoritma yang paling efektif untuk diterapkan pada *dataset* yang digunakan. Evaluasi kinerja *clustering* diukur menggunakan *Silhouette Score* dan *Davies-Bouldin Index* untuk memberikan gambaran tentang efektivitas masing-masing algoritma.

Clustering menggunakan *K-Means* dan *Fuzzy C-Means* dilakukan dengan jumlah *cluster* $k = \{2, 3, 4, 5\}$. Kinerja *clustering* diukur menggunakan *Silhouette Score* dan *Davies-Bouldin Index*. Tabel 1 menunjukkan evaluasi kedua algoritma. Algoritma *K-Means* dengan jumlah *cluster* = 3 menghasilkan *Silhouette Score* tertinggi dengan nilai 0.2846 dan *Davies-Bouldin Index* terendah dengan nilai 0.7614, namun terdapat *cluster* yang hanya berisikan satu anggota. Hal ini tidak baik karena *cluster* tersebut tidak akan memberikan informasi berguna tentang pola atau hubungan dalam data. Oleh karena itu, akan digunakan konfigurasi yang menghasilkan *Silhouette Score* terbesar dan *Davies-Bouldin Index* terendah kedua dengan nilai berturut-turut 0.2765 dan 1.2538 yang diperoleh oleh algoritma *K-Means* dengan jumlah *cluster* = 2. Output dari algoritma *K-Means* dengan jumlah *cluster* = 2 akan digunakan untuk melanjutkan analisis.

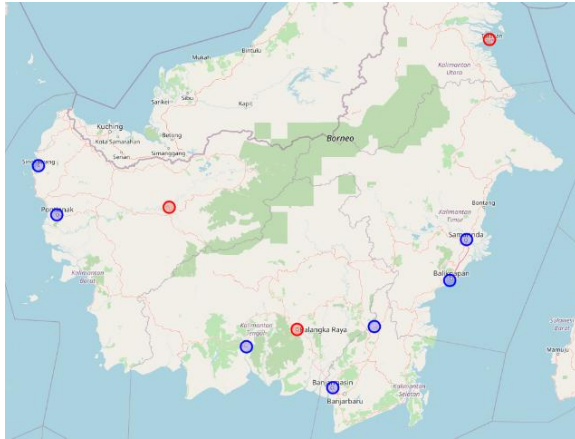
Tabel 1. Evaluasi clustering

Algoritma	Cluster	Silhouette	Davies-Bouldin
K-Means	2	0.2765	1.2538
K-Means	3	0.2846	0.7614
K-Means	4	0.1417	1.0947
K-Means	5	0.1736	0.8625
Fuzzy C-Means	2	0.2741	1.259
Fuzzy C-Means	3	0.0997	1.4976
Fuzzy C-Means	4	0.1401	1.1949
Fuzzy C-Means	5	-0.0025	1.0263

Daftar kota/kabupaten, garis lintang, garis bujur, dan *cluster* label output dari *K-Means* dengan jumlah *cluster* = 2 ditampilkan pada Tabel 2.

Tabel 2. Daftar kota/kabupaten dan label cluster

Kota/Kabupaten	Garis Lintang	Garis Bujur	Cluster
Sintang	0.1167	111.4833	0
Pontianak	-0.0226	109.3307	1
Singkawang	0.9101	108.9847	1
Banjarmasin	-3.3194	114.5944	1
Tanjung	-2.1574	115.3864	1
Palangkaraya	-2.2096	113.9213	0
Sampit	-2.5321	112.9498	1
Balikpapan	-1.2692	116.8311	1
Samarinda	-0.4952	117.1453	1
Tarakan	3.3134	117.5913	0



Gambar 2. Visualisasi persebaran anggota cluster

Visualisasi persebaran anggota *cluster* pada peta dapat dilihat pada Gambar 2. *Cluster 0* memiliki 3 anggota yang cenderung tersebar di wilayah tengah hingga utara Kalimantan, yaitu Kota Tarakan, Kota Palangkaraya, dan Kabupaten Sintang. *Cluster 1* memiliki 7 anggota yang cenderung lebih terdistribusi di wilayah pesisir dan bagian selatan Kalimantan, yaitu Kota Pontianak, Kota Singkawang, Kota Banjarmasin, Kota Tanjung, Kota Sampit, Kota Balikpapan, dan Kota Samarinda.



Gambar 3. Tren Tahunan Harga Pangan Setiap Komoditas Berdasarkan Cluster

Gambar 3 menunjukkan tren tahunan harga pangan setiap komoditas berdasarkan *cluster*. Hal-hal yang dapat dianalisis untuk setiap plot:

a. Harga Beras

- *Cluster 0* memiliki tren kenaikan harga beras yang lebih tajam dibandingkan *Cluster 1* dari tahun 2021 hingga 2023.
- Harga beras di *Cluster 1* cenderung lebih stabil dan meningkat lebih lambat.

b. Harga Telur Ayam

- *Cluster 0* menunjukkan kenaikan harga telur yang lebih tinggi dibandingkan dengan *Cluster 1*, terutama setelah tahun 2020.
- *Cluster 1* memiliki tren kenaikan yang lebih lambat dan stabil.

c. Harga Daging Ayam

- Harga daging ayam di *Cluster 0* lebih tinggi secara konsisten daripada di *Cluster 1*.

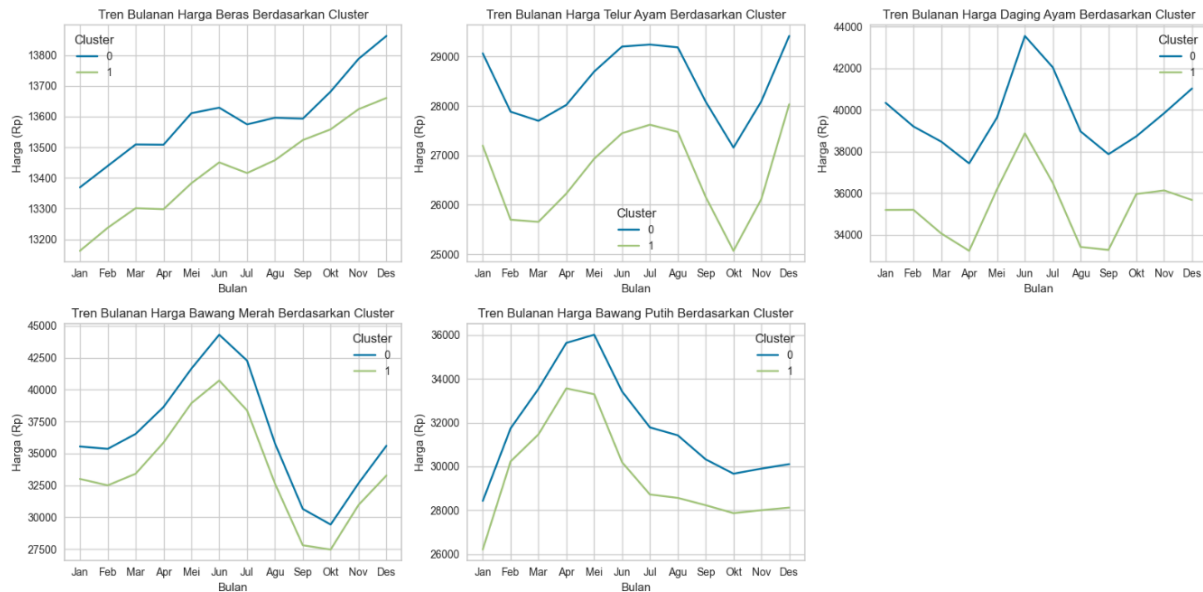
- Meskipun keduanya mengalami kenaikan harga, kenaikan di *Cluster 0* lebih signifikan.

d. Harga Bawang Merah

- Kedua *cluster* menunjukkan fluktuasi harga bawang merah yang cukup besar, tetapi *Cluster 0* cenderung memiliki harga yang lebih tinggi dan fluktuasi yang lebih besar.
- *Cluster 1* menunjukkan tren yang lebih stabil meskipun juga memiliki fluktuasi.

e. Harga Bawang

- Harga bawang putih di *Cluster 0* lebih tinggi secara konsisten daripada di *Cluster 1*.
- Keduanya menunjukkan kenaikan harga yang signifikan pada tahun 2023, tetapi *Cluster 0* memiliki harga yang lebih tinggi dan peningkatan yang lebih signifikan.



Gambar 4. Tren Bulanan Harga Pangan Setiap Komoditas Berdasarkan Cluster

Gambar 4 menunjukkan tren bulanan harga pangan setiap komoditas berdasarkan *cluster*. Hal-hal yang dapat dianalisis untuk setiap plot:

a. Harga Beras

- *Cluster 0* menunjukkan harga beras yang lebih tinggi dibandingkan dengan *Cluster 1* sepanjang tahun.
- Tren harga beras cenderung meningkat secara bertahap dari Januari hingga Desember pada kedua *cluster*, tetapi dengan peningkatan yang lebih tajam di *Cluster 0*.

b. Harga Telur Ayam

- *Cluster 0* memiliki harga telur ayam yang lebih tinggi secara konsisten dibandingkan dengan *Cluster 1*.
- Kedua *cluster* menunjukkan pola musiman yang serupa, dengan harga yang sedikit menurun di bulan Maret dan mulai meningkat lagi setelah itu.

c. Harga Daging Ayam

- *Cluster 0* menunjukkan harga daging ayam yang lebih tinggi dibandingkan dengan *Cluster 1*.
- Tren harga menunjukkan puncak di sekitar bulan Mei dan Juni, yang mungkin terkait dengan perayaan tertentu atau faktor musiman lainnya.

d. Harga Bawang Merah

- *Cluster 0* menunjukkan harga bawang merah yang lebih tinggi dan fluktuasi yang lebih besar dibandingkan dengan *Cluster 1*.
- Ada puncak harga yang signifikan sekitar bulan Mei hingga Juli di *Cluster 0*, sedangkan *Cluster 1* menunjukkan puncak yang lebih rendah pada periode yang sama.

e. Harga Bawang

- *Cluster 0* juga memiliki harga bawang putih yang lebih tinggi dibandingkan dengan *Cluster 1*.

- Kedua *cluster* menunjukkan tren harga yang meningkat dari awal tahun hingga sekitar bulan Mei, kemudian mengalami penurunan hingga akhir tahun.

5. KESIMPULAN DAN SARAN

Kesimpulan dari penelitian ini menunjukkan bahwa kedua metode *clustering*, *K-Means* dan *Fuzzy C-Means*, berhasil mengelompokkan kota/kabupaten di Kalimantan berdasarkan data harga pangan. Algoritma *K-Means* terbukti lebih efektif dengan *Silhouette Score* sebesar 0.2765 dan *Davies-Bouldin Index* sebesar 1.2538 untuk jumlah *cluster* = 2. *Cluster 0* terdiri dari Kabupaten Sintang, Kota Palangkaraya, dan Kota Tarakan, sedangkan *Cluster 1* mencakup Kota Pontianak, Kota Singkawang, Kota Banjarmasin, Kota Tanjung, Kota Sampit, Kota Balikpapan, dan Kota Samarinda. Analisis geografis menunjukkan bahwa kota/kabupaten di *Cluster 0* lebih terkonsentrasi di wilayah tengah dan utara Kalimantan, dengan harga pangan yang lebih tinggi dan fluktuasi yang lebih besar, kemungkinan disebabkan oleh faktor logistik dan aksesibilitas. Sebaliknya, *Cluster 1*, yang mencakup lebih banyak kota di pesisir dan selatan Kalimantan, memiliki harga pangan yang lebih stabil dan lebih rendah dengan fluktuasi musiman yang lebih kecil. Hasil ini dapat membantu dalam mengidentifikasi faktor-faktor yang mempengaruhi harga pangan dan mendukung perencanaan kebijakan yang lebih baik untuk mengelola harga pangan di wilayah Kalimantan.

DAFTAR PUSTAKA

- [1] Y. H. Siska, M. S. Anwari, and A. Yani, "KEANEKARAGAMAN JENIS IKAN AIR TAWAR DI SUNGAI KEPARI DAN SUNGAI EMPERAS DESA KEPARI KECAMATAN SUNGAI LAUR KABUPATEN KETAPANG," *JURNAL HUTAN LESTARI*, 2020.

- [2] M. Christyanto and H. Mayulu, "Pentingnya Pembangunan Pertanian dan Pemberdayaan Petani Wilayah Perbatasan Dalam Upaya Mendukung Ketahanan Pangan Nasional: Studi Kasus di Wilayah Perbatasan Kalimantan," *Journal of Tropical AgriFood*, vol. 3, no. 1, p. 1, Jul. 2021, doi: 10.35941/jtaf.3.1.2021.5041.1-14.
- [3] R. Sandi and D. Trisnawarman, "DESAIN DASBOARD UNTUK ANALISIS HARGA PANGAN DI INDONESIA DASHBOARD DESIGN FOR FOOD PRICE ANALYSIS IN INDONESIA," *Journal of Information Technology and Computer Science (INTECOMS)*, vol. 7, no. 3, 2024.
- [4] A. Winata, M. D. Lauro, and T. Handhayani, "Analysis and Prediction of Foodstuffs Prices in Tasikmalaya Using ELM and LSTM," *SISTEMASI*, vol. 12, no. 3, p. 874, Sep. 2023, doi: 10.32520/stmsi.v12i3.3145.
- [5] T. Ericko, M. D. Lauro, A. Winata, and T. Handhayani, "An Analysis And Forecasting The Foodstuffs Prices In Surabaya Traditional Market Using LSTM," *Jurnal CoreIT: Jurnal Hasil Penelitian Ilmu Komputer dan Teknologi Informasi*, vol. 10, no. 2, p. 99, Dec. 2024, doi: 10.24014/coreit.v10i2.27855.
- [6] J. S. Masyarakat, W. K. Atem, and N. Niko, "Persoalan Kerawanan Pangan pada Masyarakat Miskin di Wilayah Perbatasan Entikong (Indonesia-Malaysia) Kalimantan Barat Food Security at Low-Income Community in the Border Region of Entikong (Indonesia-Malaysia)," vol. 2, no. 2, pp. 94–104, 2020, doi: 10.26714/jsm.2.1.2019.94-104.
- [7] C. D. Munialo and D. D. Mellor, "A review of the impact of social disruptions on food security and food choice," *Food Sci Nutr*, vol. 12, no. 1, pp. 13–23, Jan. 2024, doi: 10.1002/fsn3.3752.
- [8] L. Lasminigrat and D. Efriza, "Pembangunan Lumbung Pangan Nasional: Strategi Antisipasi Krisis Pangan Indonesia," 2020. [Online]. Available: <https://www.>
- [9] Hari Hermawan and Harmi Andrianyta, "Respons Kebijakan Terhadap Potensi Krisis Pangan Global," *Jurnal Kebijakan Publik*, vol. 14, 2023.
- [10] H. Tayah and R. Priskila, "RANCANG BANGUN APLIKASI PENGELOLAAN DATA HARGA PANGAN DI DINAS KETAHANAN PANGAN PROVINSI KALIMANTAN TENGAH BERBASIS WEB," *Journal of Information Technology and Computer Science*, vol. 1, 2021.
- [11] T. Handhayani and I. Lewenusa, "An Intelligent Clustering Approach For Analyzing A Multivariate Time Series Dataset, Case Study COVID-19 Outbreak in Indonesia," in *2023 17th International Conference on Telecommunication Systems, Services, and Applications (TSSA)*, IEEE, Oct. 2023, pp. 1–6. doi: 10.1109/TSSA59948.2023.10367007.
- [12] T. Handhayani and Z. Rusdi, "K-Means Using Dynamic Time Warping For Clustering Cities in Java Island According to Meteorological Conditions," 2023. doi: 10.1109/ICIC60109.2023.10381899.
- [13] G. Andrian, D. Arisandi, and T. Handhayani, "CLUSTERING DATA METEOROLOGI WILAYAH INDONESIA TIMUR DENGAN METODE K-MEANS DAN FUZZY C-MEANS," *INTI Nusa Mandiri*, vol. 18, no. 2, pp. 100–106, Feb. 2024, doi: 10.33480/inti.v18i2.5039.
- [14] Y. A. Ishak, T. Handhayani, M. D. L. Sitorus, W. William, J. Pragantha, and I. Lewenusa, "Advanced Clustering Approach for Mapping Regions of Paddy Productivity in Indonesia Using Intelligent K-Means," in *IEEE Asia-Pacific Conference on Geoscience, Electronics and Remote Sensing Technology (AGERS)*, Manado: IEEE, Mar. 2024, pp. 269–274.
- [15] William, L. Bayuaji, N. J. Perdana, and T. Handhayani, "Mapping Indonesia's Regions Based on Carbon Emissions Using the K-Means Algorithm," in *2025 International Conference on Computer Sciences, Engineering, and Technology Innovation (ICoCSETI)*, IEEE, Jan. 2025, pp. 200–205. doi: 10.1109/ICoCSETI63724.2025.11020099.
- [16] T. Handhayani and I. Lewenusa, "An Analysis of Meteorological Data in Sumatra and Nearby using Agglomerative Clustering," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 8, no. 2, pp. 234–241, Apr. 2024, doi: 10.29207/resti.v8i2.5663.
- [17] C. F. Palembang and S. P. Palembang, "PENGELOMPOKKAN TINGKAT HARGA CABAI RAWIT BERDASARKAN PROVINSI DI INDONESIA MENGGUNAKAN ALGORITMA K-MEANS CLUSTERING (per Juli 2020-Juli 2021)," *VARIANCE: Journal of Statistics and Its Applications*, vol. 3, no. 2, pp. 48–60, Dec. 2021, doi: 10.30598/variancevol3iss2page48-60.
- [18] A. Cluster Provinsi Indonesia Berdasarkan Produksi Bahan Pangan Menggunakan Algoritma K-Means, T. Tendean, and W. Purba, "Analisis Cluster Provinsi Indonesia Berdasarkan Produksi Bahan Pangan Menggunakan Algoritma K-Means," *Jurnal Sains dan Teknologi*, vol. 1, no. 2, pp. 5–11.
- [19] D. Nurmin, M. N. Hayati, and R. Goejantoro, "Application of the Fuzzy C-Means Method in the Grouping of Regencies/Cities in Kalimantan Island Based on People's Welfare Indicators in 2020".
- [20] I. M. Nur, A. N. L. Syifa, M. Kharis, and S. Heidya Permatasari, "IMPLEMENTASI METODE FUZZY C-MEANS DALAM

- PENGELOMPOKKAN HASIL PANEN PADI DI PROVINSI BALI,” *VARIANCE: Journal of Statistics and Its Applications*, vol. 5, no. 1, pp. 13–24, Sep. 2023, doi: 10.30598/variancevol5iss1page13-24.
- [21] Q. I. Mawarni and E. S. Budi, “Implementasi Algoritma K-Means Clustering Dalam Penilaian Kedisiplinan Siswa,” *Jurnal Sistem Komputer dan Informatika (JSON)*, vol. 3, no. 4, p. 522, Jun. 2022, doi: 10.30865/json.v3i4.4242.
- [22] Y. Diputra, E. Mirshad, J. Hamka Kampus UNP, and A. Tawar Padang, “Jurnal Vocational Teknik Elektronika dan Informatika,” vol. 11, no. 3, 2023, [Online]. Available: <http://ejournal.unp.ac.id/index.php/voteknika/>
- [23] S. Kurniawan, A. M. Siregar, and H. Y. Novita, “Penerapan Algoritma K-Means dan Fuzzy C-Means Dalam Mengelompokan Prestasi Siswa Berdasarkan Nilai Akademik,” vol. IV, no. 1, 2023.
- [24] T. D. Khang, N. D. Vuong, M. K. Tran, and M. Fowler, “Fuzzy C-means clustering algorithm with multiple fuzzification coefficients,” *Algorithms*, vol. 13, no. 13, Jul. 2020, doi: 10.3390/A13070158.
- [25] K. R. Shahapure and C. Nicholas, “Cluster Quality Analysis Using Silhouette Score,” 2020. [Online]. Available: <https://www.>
- [26] E. Biabiany, D. C. Bernard, V. Page, and H. Paugam-Moisy, “Design of an expert distance metric for climate clustering: The case of rainfall in the Lesser Antilles,” *Comput Geosci*, vol. 145, Dec. 2020, doi: 10.1016/j.cageo.2020.104612.
- [27] 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 Journal*, vol. 10, no. 3, pp. 1099–1103, Aug. 2021, doi: 10.18421/TEM103-13.
- [28] F. Ros, R. Riad, and S. Guillaume, “PDBI: A partitioning Davies-Bouldin index for clustering evaluation,” *Neurocomputing*, vol. 528, pp. 178–199, Apr. 2023, doi: 10.1016/j.neucom.2023.01.043.
- [29] G. Andrian, D. Arisandi, and T. Handhayani, “CLUSTERING DATA METEOROLOGI WILAYAH INDONESIA TIMUR DENGAN METODE K-MEANS DAN FUZZY C-MEANS,” *INTI Nusa Mandiri*, vol. 18, no. 2, pp. 100–106, Feb. 2024, doi: 10.33480/inti.v18i2.5039.
- [30] D. D. Purwanto and E. S. Honggara, “Klasifikasi Kategori Hasil Perhitungan Indeks Standar Pencemaran Udara dengan Gaussian Naïve Bayes (Studi Kasus: ISPU DKI Jakarta 2020),” *Journal of Intelligent System and Computation*, vol. 4, no. 2, pp. 102–108, Oct. 2022, doi: 10.52985/insyst.v4i2.259.
- [31] C. Yuan and H. Yang, “Research on K-Value Selection Method of K-Means Clustering Algorithm,” *J (Basel)*, vol. 2, no. 2, pp. 226–235, Jun. 2019, doi: 10.3390/j2020016.
- [32] K. P. Sinaga and M. S. Yang, “Unsupervised K-means clustering algorithm,” *IEEE Access*, vol. 8, pp. 80716–80727, 2020, doi: 10.1109/ACCESS.2020.2988796.
- [33] M. Ahmed, R. Seraj, and S. M. S. Islam, “The k-means algorithm: A comprehensive survey and performance evaluation,” Aug. 01, 2020, *MDPI AG*. doi: 10.3390/electronics9081295.
- [34] “Pusat Informasi Harga Pangan Nasional.” Accessed: Jun. 26, 2024. [Online]. Available: <https://www.bi.go.id/hargapangan/TabelHargaPasarTradisionalDaerah>