

KOMPARASI ALGORITMA K-MEANS DAN FUZZY C-MEANS DALAM ANALISIS CLUSTERING MENU PADA PERUSAHAAN MAKANAN CEPAT SAJI MULTINASIONAL

Rahmad Hidayat, Muhammad Faturrahman, Dekka Andri Cahya,

Dimas Agil Kusuma Ken Ditha Tania, Ahmad Rifai

Sistem Informasi, Universitas Sriwijaya

Jalan Raya Palembang-Prabumulih, KM 32, Indralaya, Kabupaten Ogan Ilir, Sumatera Selatan, Indonesia.

rahmadhidayat6875@gmail.com

ABSTRAK

Tingginya konsumsi makanan cepat saji secara global menimbulkan tantangan dalam mengidentifikasi pola distribusi nutrisi secara otomatis dan objektif. Permasalahan utama yang dihadapi adalah sulitnya menentukan algoritma clustering yang paling optimal untuk memetakan menu makanan berdasarkan kandungan kalori dan lemak. Penelitian ini bertujuan untuk membandingkan performa algoritma K-Means dan Fuzzy C-Means (FCM) dalam pengelompokan menu makanan cepat saji menggunakan dataset Fast Food Nutrition dari Kaggle. Metode yang digunakan mencakup prapemrosesan data (normalisasi Z-score), penentuan jumlah kluster optimal (K-Means dengan Elbow Method menghasilkan $k=3$; FCM dengan Fuzzy Partition Coefficient menghasilkan $k=2$), serta evaluasi menggunakan Silhouette Score dan Davies-Bouldin Index (DBI). Hasil penelitian menunjukkan bahwa Fuzzy C-Means menghasilkan pengelompokan yang lebih baik dibandingkan K-Means, dengan Silhouette Score sebesar 0,5905 (lebih tinggi) dan DBI sebesar 0,6530 (lebih rendah). Temuan ini membuktikan bahwa pendekatan logika fuzzy memberikan hasil pengelompokan yang lebih solid dan efisien secara matematis untuk dataset nutrisi makanan cepat saji dibandingkan pendekatan hard clustering.

Kata kunci : *Fast Food, Clustering, K-Means, Fuzzy C-Means, Nutrition*

1. PENDAHULUAN

Industri makanan cepat saji (fast food) berskala global telah menjadi pilihan utama konsumsi masyarakat karena kemudahan akses dan variasi menu yang ditawarkan. Konsumsi makanan cepat saji yang tidak terkontrol berpotensi menimbulkan berbagai masalah kesehatan, sehingga pemetaan karakteristik nutrisi menu makanan menjadi sangat penting sebagai dasar informasi objektif bagi konsumen maupun pakar kesehatan. Pemetaan menu berdasarkan kandungan kalori, lemak, sodium, dan protein sangat diperlukan untuk memberikan gambaran yang terstruktur mengenai asupan gizi dari produk-produk fast food tersebut. Untuk memetakan data berskala besar secara otomatis dan menemukan pola yang tersembunyi, teknik machine learning khususnya clustering telah terbukti memberikan hasil segmentasi yang akurat pada klasifikasi produk makanan berdasarkan nilai nutrisinya [1]. Metode clustering yang paling umum digunakan dalam segmentasi data adalah algoritma K-Means karena memiliki tingkat kecepatan dan efisiensi komputasi yang sangat baik dalam mengelompokkan data berskala besar secara tegas [2].

Permasalahan yang muncul adalah tingginya keberagaman menu dari berbagai perusahaan fast food multinasional menyulitkan proses identifikasi pola nutrisi secara manual. Selain itu, pemilihan algoritma clustering yang tepat menjadi tantangan tersendiri karena masing-masing metode memiliki karakteristik yang berbeda dalam menangani data gizi. Meskipun K-Means sangat efisien, metode ini melakukan pembagian kelompok secara tegas (hard clustering)

yang terkadang kurang merepresentasikan realitas multidimensi data gizi, di mana batas antar-kategori nutrisi bisa saling tumpang tindih. Kinerja dan stabilitas K-Means sering kali menjadi acuan dasar yang perlu dievaluasi lebih lanjut dengan membandingkannya terhadap algoritma lain guna mendapatkan model komputasi yang paling optimal pada dataset spesifik [3]. Untuk melengkapi pendekatan hard clustering tersebut, pendekatan logika fuzzy melalui algoritma Fuzzy C-Means (FCM) digunakan karena memungkinkan setiap data menu makanan untuk memiliki nilai derajat keanggotaan pada lebih dari satu kluster secara bersamaan. Pendekatan ini sangat fleksibel dan cocok untuk menganalisis data gizi yang memiliki sifat campuran, sehingga pengelompokan yang dihasilkan menjadi lebih representatif [4].

Berdasarkan pemaparan permasalahan tersebut, tujuan penelitian ini adalah untuk membandingkan secara komprehensif performa algoritma K-Means dan Fuzzy C-Means dalam mengelompokkan menu makanan cepat saji berdasarkan kandungan kalori dan total lemak, serta menentukan algoritma yang paling optimal untuk kasus dataset nutrisi fast food. Sebagai rencana penyelesaian masalah, penelitian ini menggunakan dataset Fast Food Nutrition dari Kaggle yang mencakup menu makanan dari empat perusahaan junk food global. Data melalui tahapan prapemrosesan (pembersihan dan normalisasi Z-score), kemudian diproses menggunakan K-Means dengan penentuan kluster optimal melalui Elbow Method dan FCM dengan penentuan kluster melalui skor Fuzzy Partition

Coefficient (FPC). Evaluasi perbandingan performa antara kedua algoritma ini diukur secara objektif menggunakan indeks validitas internal (Silhouette Score dan Davies-Bouldin Index) guna memastikan model machine learning yang dihasilkan memiliki tingkat akurasi klastering terbaik dan dapat diandalkan sebagai referensi analisis data nutrisi [5].

2. TINJAUAN PUSTAKA

Bagian ini menyoroti berbagai penelitian terdahulu yang berkaitan dengan deteksi pola nutrisi serta klasifikasi pengelompokan menggunakan metode machine learning. [1] menggunakan algoritma K-Means untuk mengelompokkan makanan berdasarkan metrik nilai nutrisi. Para peneliti tersebut memanfaatkan analisis klastering untuk memberikan segmentasi yang akurat pada produk makanan. K-Means terbukti mampu mengelompokkan data berskala besar secara tegas (hard clustering) karena memiliki tingkat kecepatan dan efisiensi komputasi yang baik.

Selain algoritma tunggal, [4] telah melakukan penelitian komparatif yang menguji algoritma K-Means dan Fuzzy C-Means (FCM) dalam ranah gizi untuk mengklasifikasi status gizi balita pada puskesmas di Surabaya. Pendekatan fuzzy memberikan alternatif yang kuat ketika berhadapan dengan data multidimensi, di mana batas antar kategori gizi kerap tumpang tindih. Keandalan algoritma Fuzzy C-Means (FCM) dalam menangani data yang kompleks tersebut juga diperkuat oleh penelitian [6]. Dalam risetnya, metode FCM terbukti tangguh (robust) terhadap keberadaan data pencilon (outlier) karena inklusi data ditentukan berdasarkan derajat keanggotaan, bukan partisi mutlak. Karakteristik ini membuat metode FCM mampu menghasilkan keluaran yang luwes dan andal saat memproses himpunan data yang kerap memiliki nilai ekstrem.

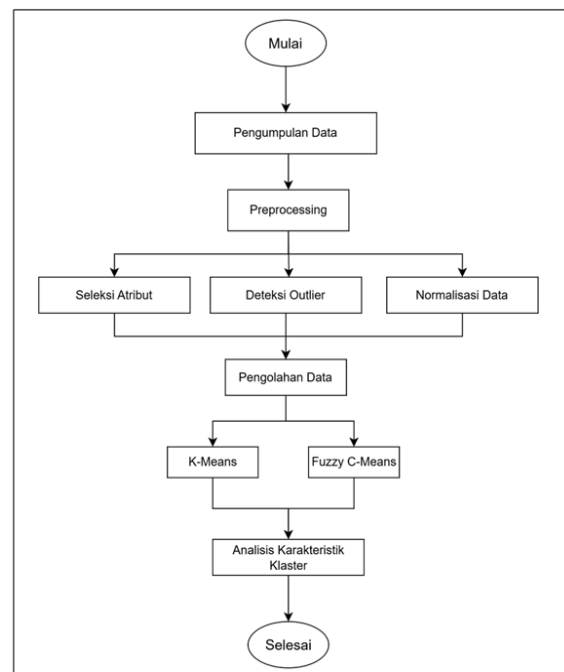
Tinjauan literatur di atas menunjukkan bahwa berbagai penelitian terkait klastering gizi telah banyak dilakukan secara luas, namun sebagian besar studi berfokus pada segmentasi kesehatan umum terbuka atau hanya menerapkan satu metode evaluasi dasar. Pekerjaan kami melengkapi pendekatan hard clustering dan soft clustering pada ranah yang lebih terfokus, yakni pada menu perusahaan makanan cepat saji multinasional (junk food). Penelitian kami membandingkan secara langsung algoritma K-Means dan Fuzzy C-Means untuk memetakan dataset nutrisi berdasarkan total kalori dan total lemak. Lebih jauh, tidak seperti beberapa riset sebelumnya, evaluasi kami secara spesifik mendefinisikan batas partisi menggunakan metrik objektif; pembentukan tiga klaster pada K-Means divalidasi dengan Elbow Method, sedangkan pembentukan dua klaster pada FCM divalidasi menggunakan skor Fuzzy Partition Coefficient (FPC).

Kebaruan (novelty) dari penelitian ini dibuktikan melalui evaluasi performa komputasinya, di mana Fuzzy C-Means menghasilkan pengelompokan yang

lebih solid dan efisien dibandingkan K-Means dalam kasus data makanan cepat saji. Keunggulan tersebut terekam dalam perolehan Silhouette Score FCM sebesar 0,5905 yang lebih tinggi dan nilai Davies-Bouldin Index (DBI) sebesar 0,6530 yang lebih rendah. Oleh karena itu, pendekatan logika fuzzy dalam riset kami menawarkan pemahaman dan pemetaan distribusi nutrisi fast food yang lebih representatif dibandingkan batas kategori kaku (hard clustering).

3. METODE PENELITIAN

Penelitian ini dilaksanakan melalui serangkaian tahapan yang sistematis untuk mencapai tujuan yang telah ditetapkan. Secara garis besar, metodologi penelitian ini mencakup tahap pengumpulan data, prapemrosesan data (yang meliputi seleksi atribut, deteksi outlier, dan normalisasi), proses pengolahan data menggunakan algoritma K-Means dan Fuzzy C-Means, serta diakhiri dengan analisis karakteristik klaster. Keseluruhan alur kerja penelitian tersebut divisualisasikan melalui diagram alir pada Gambar 1.



Gambar 1. Tahapan Penelitian

3.1. Pengumpulan Data

Pada penelitian ini dataset yang didapat bersumber dari situs dataset Kaggle (<https://www.kaggle.com/code/joebeachcapital/fast-food-nutrition-eda>) yang diunggah oleh Joakim Arvidsson yang didapatkan dari situs Kaggle. Gambar 2 adalah cuplikan dari dataset yang diperoleh dari Kaggle.

Company	Item	Calories	Total Fat	Sodium
McDonald's	Hamburger	250	9	520
Burger King	Whopper® Sandwich	660	40	980
Wendy's	Baconator	950	62	1810
KFC	Limited Time Cinnabon Dessert Biscuits	290	14	310
Taco Bell	Bacon Club Chalupa	440	25	1000
Pizza Hut	Detroit Double Cheesy Pizza Slice	280	9	580

Gambar 2. Dataset Nutrisi Situs Kaggle

3.2. Seleksi Atribut

Pada tahapan ini, seleksi atribut difokuskan pada variabel utama yang relevan, dan tahap prapemrosesan data seperti pembersihan data (data cleaning) sangat krusial dilakukan sebelum menerapkan algoritma klustering [7]. Data mentah dibersihkan dengan cara mengekstraksi nilai numerik, dan item data yang memiliki nilai kosong (Missing Values) dihapus secara selektif agar tahapan klustering dapat memproses data dengan valid [7]. Proses ini bertujuan mereduksi dimensi yang tidak relevan dan mengoptimalkan kualitas segmentasi pada model algoritma yang digunakan [7].

3.3. Deteksi Outlier

Berdasarkan peninjauan karakteristik data, penelitian ini mempertahankan nilai ekstrem dan tidak memotong outlier karena mempertimbangkan karakteristik dari algoritma Fuzzy C-Means (FCM) itu sendiri [6]. Metode FCM menentukan inklusi data berdasarkan derajat keanggotaan, yang terbukti tangguh (robust) terhadap keberadaan data outlier [6]. Oleh karena itu, metode FCM sering digunakan dalam analisis kluster karena mampu menghasilkan keluaran yang halus serta memiliki keandalan yang tinggi dalam menangani himpunan data yang mengandung nilai ekstrem [6].

3.4. Normalisasi Data

Untuk mengatasi ketimpangan skala pada rentang data, penelitian ini menerapkan teknik normalisasi berdasarkan metode Z-score [8]. Sebelum dilakukan proses pengujian menggunakan algoritma klustering seperti K-Means, data perlu ditransformasikan agar proses perhitungan jarak berjalan optimal [8]. Di samping itu, tahap pembersihan dan normalisasi data merupakan standar krusial dalam prapemrosesan agar setiap fitur memiliki skala yang seragam, sehingga bobot perhitungan jarak sentroid menjadi akurat.

3.5. Algoritma K-Means

K-Means merupakan salah satu metode data clustering non-hierarki yang mempartisi data ke dalam bentuk satu atau lebih cluster [9]. Algoritma ini mengelompokkan data berdasarkan karakteristik yang sama ke dalam satu cluster yang sama, sedangkan data dengan karakteristik berbeda dikelompokkan ke dalam kelompok yang lain. K-Means menggunakan pendekatan *hard clustering*, di mana setiap objek data secara mutlak hanya menjadi anggota dari satu cluster tertentu [9].

Teknik partisi ini berbasis sentroid (centroid-based), di mana sentroid mewakili titik pusat dari sebuah cluster. Perbedaan atau tingkat kemiripan antara sebuah objek dengan perwakilan cluster (sentroid) diukur menggunakan jarak Euclidean. Kualitas pengelompokan dievaluasi dengan meminimalkan variasi within-cluster, yaitu jumlah kuadrat kesalahan (Sum of Squared Errors atau SSE) antara semua objek dengan sentroidnya [9]. Fungsi objektif ini didefinisikan sebagai berikut:

$$SSE = \sum_{i=1}^k \sum_{x \in C_i} \|x - \mu_i\|^2 \quad (1)$$

menggunakan rumus Euclidean Distance:

$$d(x, \mu) = \sqrt{\sum_{j=1}^m (x_j - \mu_j)^2} \quad (2)$$

Langkah-langkah dari algoritma K-Means adalah sebagai berikut:

- Menentukan nilai k sebagai jumlah cluster yang ingin dibentuk.
- Menentukan titik pusat cluster (centroid) awal secara acak sebanyak k
- Menghitung jarak setiap titik data input terhadap masing-masing centroid menggunakan rumus Euclidean Distance
- Mengelompokkan setiap data berdasarkan kedekatannya dengan centroid (mencari jarak terkecil)
- Memperbarui nilai centroid baru dengan menghitung nilai rata-rata (mean) dari semua data yang berafiliasi pada masing-masing cluster
- Mengulangi langkah 3 dan 5 sampai posisi centroid konvergen (tidak ada lagi data yang berpindah cluster)

3.6. Algoritma Fuzzy C-Means

Fuzzy C-Means (FCM) adalah metode *clustering* yang mengadopsi model pengelompokan berbasis logika fuzzy (*soft clustering*). Berbeda dengan K-Means, keberadaan setiap titik data dalam suatu cluster pada FCM tidak bersifat mutlak, melainkan ditentukan oleh nilai derajat keanggotaan (membership degree) yang berada pada rentang 0 hingga 1 [10]. Pendekatan ini dinilai sangat tangguh untuk menangani kumpulan data yang memiliki karakteristik tumpang tindih (*overlapping*), karena satu titik data dapat berafiliasi dengan beberapa cluster sekaligus dengan bobot proporsi yang berbeda [10].

Tujuan utama dari algoritma FCM adalah meminimalkan fungsi objektif yang menggabungkan jarak data ke pusat cluster dengan derajat keanggotaannya [10]. Fungsi objektif FCM dirumuskan sebagai berikut:

$$J_m = \sum_{i=1}^N \sum_{j=1}^C u_{ij}^m \|x_i - c_j\|^2 \quad (3)$$

Pembaruan derajat keanggotaan matriks (u_{ij})

dan pembaruan letak pusat cluster (c_j) pada setiap proses iterasi dihitung menggunakan dua persamaan berikut:

$$u_{ij} = \frac{1}{\sum_{k=1}^C \left(\frac{\|x_i - c_j\|}{\|x_i - c_k\|} \right)^{\frac{2}{m-1}}} \quad (4)$$

Langkah-langkah operasional algoritma Fuzzy C-Means adalah sebagai berikut:

- Menentukan parameter awal pengujian: Jumlah cluster (C), Nilai Fuzzifier pembobot (m), maksimum iterasi, dan batas error toleransi ($threshold \epsilon$)
- Menginisialisasi matriks keanggotaan fuzzy awal $U = [u_{ij}]$ secara acak. dengan syarat bahwa total jumlah derajat keanggotaan sebuah titik data pada seluruh cluster harus bernilai 1 ($\sum_{j=1}^C u_{ij} = 1$)
- Menghitung posisi pusat cluster (c_j) berdasarkan matriks keanggotaan yang ada
- Memperbarui matriks keanggotaan U yang baru berdasarkan kedekatan jarak ke pusat cluster yang baru dihitung
- Memeriksa kondisi penghentian (stopping condition): Jika selisih perubahan matriks U pada iterasi saat ini dengan iterasi sebelumnya lebih kecil dari batas error ϵ , atau iterasi telah mencapai batas maksimum, maka komputasi dihentikan. Jika tidak, proses berulang kembali ke langkah 3.

3.7. Evaluasi Silhouette Score

Silhouette Score digunakan untuk mengukur seberapa baik suatu objek ditempatkan dalam klasternya sendiri dibandingkan dengan klaster lainnya. Nilai Silhouette berkisar antara -1 hingga 1, di mana nilai yang mendekati 1 menunjukkan bahwa data telah terkelompok dengan sangat baik. Rumus untuk menghitung nilai silhouette pada satu objek adalah:

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

3.8. Evaluasi Davies-Bouldin Index (DBI)

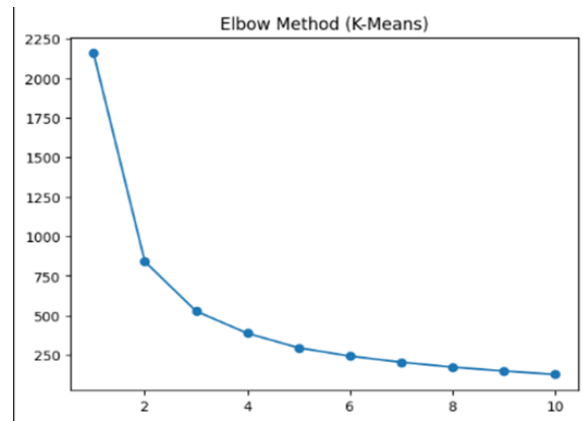
Davies-Bouldin Index merupakan metrik evaluasi yang menghitung rasio jumlah penyebaran di dalam klaster terhadap jarak antar-klaster. Berbeda dengan Silhouette, nilai DBI yang lebih kecil (mendekati 0) menunjukkan kualitas klustering yang lebih baik karena menandakan klaster yang padat dan terpisah secara optimal. Rumus DBI didefinisikan sebagai:

$$DBI = \frac{1}{k} \sum_{i=1}^k \max_{j \neq i} \left(\frac{s_i + s_j}{d_{ij}} \right) \quad (6)$$

4. HASIL DAN PEMBAHASAN

4.1. Hasil K-Means Clustering

Tahap pertama dalam penerapan algoritma K-Means adalah menentukan jumlah klaster (k) yang paling optimal. Pada penelitian ini, penentuan jumlah klaster dilakukan menggunakan metode Elbow berdasarkan nilai Within-Cluster Sum of Squares (WCSS). Evaluasi penurunan nilai WCSS pada berbagai jumlah klaster divisualisasikan pada Gambar 3.



Gambar 3. Elbow method

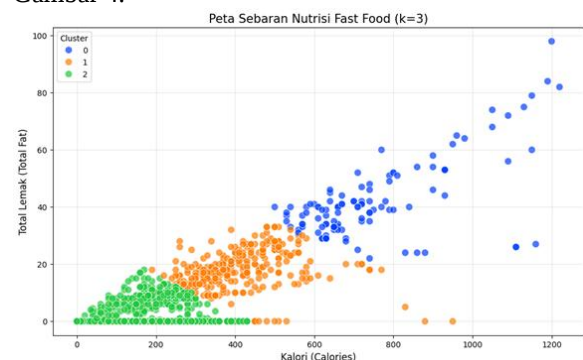
Berdasarkan grafik pada Gambar 1, terlihat bahwa penurunan nilai WCSS mulai melandai secara signifikan setelah titik $k = 3$. Oleh karena itu, jumlah klaster sebanyak 3 (tiga) dipilih sebagai parameter optimal. Berdasarkan nilai parameter tersebut, algoritma K-Means mempartisi data nutrisi menu makanan cepat saji menjadi tiga kelompok dengan karakteristik rata-rata (centroid) yang berbeda.

Tabel 1. Rata-rata per cluster k-means

Cluster	Calories	Total Fat
0	749.203540	42.176991
1	408.918495	17.884013
2	157.712519	3.299073

Berdasarkan Tabel 1, hasil *clustering* k-means membagi data menjadi tiga kelompok. Cluster 0 memiliki rata-rata kalori dan total fat tertinggi, diikuti cluster 1 dengan nilai sedang, dan cluster 2 dengan nilai terendah. Hal ini menunjukkan bahwa data berhasil dikelompokkan berdasarkan tingkat kandungan kalori dan lemak.

Pemetaan secara spasial dari ketiga klaster tersebut diilustrasikan melalui scatter plot pada Gambar 4.



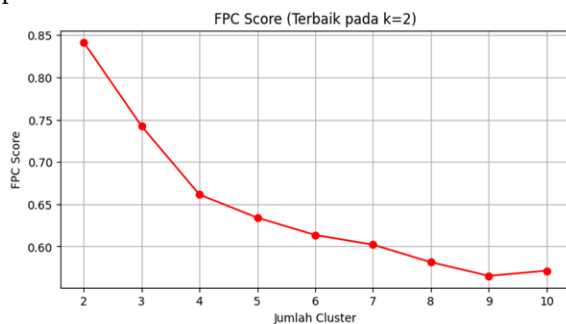
Gambar 4. Pemetaan k-means clustering

Melalui visualisasi pada Gambar 4, dapat diamati bahwa metode K-Means mampu memisahkan titik-titik data ke dalam batas-batas kelompok (partisi) yang tegas (*hard clustering*). Pemisahan ini memberikan gambaran yang jelas mengenai segmentasi menu

makanan cepat saji dari segi densitas energi dan kandungan lemaknya.

4.2. Hasil Fuzzy C-Means

Berbeda dengan K-Means, algoritma Fuzzy C-Means (FCM) mengakomodasi ambiguitas data melalui pendekatan derajat keanggotaan (nilai probabilitas 0 hingga 1). Untuk menentukan partisi data yang paling stabil pada FCM, digunakan evaluasi metrik Fuzzy Partition Coefficient (FPC). Nilai FPC yang mendekati 1 menunjukkan tingkat tumpang tindih (overlapping) antar kluster yang semakin kecil. Hasil pengujian nilai FPC pada dataset ini ditunjukkan pada Gambar 5.



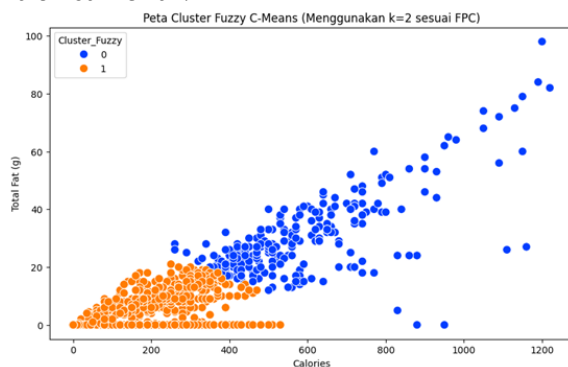
Gambar 5. FPC score

Dengan menetapkan jumlah kluster sebanyak 2, proses klustering menggunakan FCM dilakukan. Pengelompokan didasarkan pada nilai derajat keanggotaan tertinggi dari setiap item menu terhadap pusat kluster (centroid). Hasil komputasi FCM menunjukkan profil kelompok yang berdekatan namun terdistribusi secara fleksibel.

Tabel 2. Rata-rata per cluster fuzzy c-means

Cluster	Calories	Total Fat
0	588.651079	30.262590
1	191.635456	5.234082

Berdasarkan Tabel 2, hasil *clustering* fuzzy c-means membagi data menjadi dua kelompok. Cluster 0 memiliki rata-rata kalori dan total fat yang lebih tinggi, sedangkan cluster 1 memiliki nilai yang lebih rendah. Hal ini menunjukkan bahwa data berhasil dikelompokkan berdasarkan perbedaan kandungan kalori dan lemak.



Gambar 6. Pemetaan fuzzy c-means clustering

Visualisasi pada Gambar 6 menunjukkan karakteristik *soft clustering* khas dari logika fuzzy, di mana beberapa titik data yang posisinya berdekatan dengan batas antar-kluster memiliki keanggotaan ganda. Meskipun batas partisi lebih luwes dibandingkan K-Means, FCM tetap berhasil mengisolasi kelompok menu makanan berat dan makanan ringan sesuai dengan pola distribusi nutrisinya.

4.3. Analisa Karakteristik Klustering

Setelah melakukan pengelompokan menggunakan algoritma K-Means dan Fuzzy C-Means, tahap selanjutnya adalah melakukan evaluasi untuk membandingkan performa kedua model tersebut. Pengujian ini menggunakan dua metrik evaluasi klustering, yaitu Silhouette Score untuk mengukur seberapa baik objek berada di dalam klusternya sendiri dibandingkan kluster lain (nilai mendekati 1 semakin baik), dan Davies-Bouldin Index (DBI) untuk mengukur rasio jarak antarpusat kluster (nilai mendekati 0 semakin baik). Hasil perbandingan skor evaluasi dari kedua algoritma tersebut dapat dilihat pada Table 3.

Tabel 3. Perbandingan antar algoritma

Metode	Silhouette Score	Davies-Bouldin Index
K-Means	0.4979	0.7035
Fuzzy C-Means	0.5905	0.6530

Berdasarkan hasil pengujian yang disajikan, terdapat perbedaan performa yang cukup signifikan antara kedua model. Algoritma Fuzzy C-Means ($k=2$) terbukti menghasilkan kualitas pengelompokan yang lebih baik dibandingkan K-Means ($k=3$). Hal ini dibuktikan dengan nilai Silhouette Score FCM yang lebih tinggi (0.5905) dan nilai Davies-Bouldin Index yang lebih rendah (0.6530). Hasil ini menunjukkan bahwa karakteristik nutrisi pada menu makanan cepat saji cenderung lebih stabil dan terdefinisi dengan baik ketika dipartisi menjadi dua kelompok utama menggunakan pendekatan keanggotaan fuzzy.

5. KESIMPULAN DAN SARAN

Penelitian ini telah berhasil membandingkan penerapan algoritma K-Means dan Fuzzy C-Means dalam mengelompokkan data nutrisi makanan cepat saji berdasarkan parameter kalori dan lemak. Berdasarkan hasil pengujian, algoritma Fuzzy C-Means ($k=2$) terbukti sebagai model yang lebih optimal dibandingkan K-Means ($k=3$) untuk dataset ini. Keunggulan tersebut ditunjukkan oleh perolehan metrik Silhouette Score yang lebih tinggi, yaitu 0.5905, yang menandakan kepadatan objek di dalam kluster lebih solid, serta nilai Davies-Bouldin Index yang lebih rendah, yaitu 0.6530, yang mengindikasikan pemisahan jarak antar-kluster yang lebih optimal.

Secara praktis, pemetaan ini menunjukkan bahwa penggunaan logika fuzzy dengan dua kelompok utama memberikan gambaran distribusi nutrisi yang lebih stabil dan efisien secara matematis dibandingkan pembagian tiga kelompok secara kaku menggunakan K-Means. Sebagai prospek pengembangan selanjutnya, disarankan untuk mengeksplorasi penggunaan algoritma berbasis kepadatan atau mengintegrasikan fitur nutrisi tambahan seperti sodium dan gula guna menghasilkan analisis profil kesehatan menu makanan cepat saji yang lebih komprehensif.

DAFTAR PUSTAKA

- [1] P. Alga Vredizon, H. Firmansyah, N. Shafira Salsabila, and W. Eko Nugroho, "Penerapan Algoritma K-Means Untuk Mengelompokkan Makanan Berdasarkan Nilai Nutrisi," *J. Technol. Informatics*, vol. 5, no. 2, pp. 108–115, 2024, doi: 10.37802/joti.v5i2.577.
- [2] M. Wahyudi, S. Solikhun, and L. Pujiastuti, "Komparasi K-Means Clustering dan K-Medoids Clustering dalam Mengelompokkan Produksi Susu Segar di Indonesia Berdasarkan Nilai DBI," *J. Bumigora Inf. Technol.*, vol. 4, no. 2, pp. 243–254, 2022, doi: 10.30812/bite.v4i2.2104.
- [3] H. A. Ulvi and M. Ikhsan, "Comparison of K-Means and K-Medoids Clustering Algorithms for Export and Import Grouping of Goods in Indonesia," *Sinkron*, vol. 8, no. 3, pp. 1671–1685, 2024, doi: 10.33395/sinkron.v8i3.13815.
- [4] A. Kurnia, "Perbandingan Algoritma K-Means dan Fuzzy C-Means Untuk Clustering Puskesmas Berdasarkan Gizi Balita Surabaya," *J. Process.*, vol. 18, no. 1, pp. 83–88, 2023, doi: 10.33998/processor.2023.18.1.696.
- [5] N. Ulinuha and D. C. R. Novitasari, "Penerapan Fuzzy C-Means Untuk Pengelompokkan Tingkat Kualitas Pendidikan Di Jawa Timur," *Simetris J. Tek. Mesin, Elektro dan Ilmu Komput.*, vol. 14, no. 2, pp. 419–426, 2023, doi: 10.24176/simet.v14i2.9442.
- [6] I. J. Panjaitan, I. Indahwati, and F. M. Afendi, "The Implementation of the Fuzzy C-Means Method in Handling Outlier Data in the 2021 Village Potential Data of Bengkulu Province," *ComTech Comput. Math. Eng. Appl.*, vol. 16, no. 1, pp. 23–33, 2025, doi: 10.21512/comtech.v16i1.12274.
- [7] I. P. Satria, D. Wibawa, and M. Agung, "Analisis Perbandingan K-Means ++ , Mini Batch K-Means , dan Fuzzy C-Means pada Segmentasi Pelanggan," vol. 4, no. November, pp. 171–180, 2025.
- [8] S. E. Saqila, I. P. Ferina, and A. Iskandar, "Analisis Perbandingan Kinerja Clustering Data Mining Untuk Normalisasi Dataset," *J. Sist. Komput. dan Inform.*, vol. 5, no. 2, p. 356, 2023, doi: 10.30865/json.v5i2.6919.
- [9] Faozan Dwiki Ramadana, Wahyu Putra Pratama, Cannes Lingga Yogario, Abdul Khohar, and Ito Setiawan, "Implementasi Algoritma K-Means Clustering terhadap Tingkat Kepuasan Peserta LKP Multi Talenta Komputer Purwokerto," *Mars J. Tek. Mesin, Ind. Elektro Dan Ilmu Komput.*, vol. 3, no. 1, pp. 184–193, 2025, doi: 10.61132/mars.v3i1.675.
- [10] P. S. R.S, P. S. R.S, Y. Reswan, Y. Apridiansyah, and D. Sunardi, "PENERAPAN FUZZY C-Means CLUSTERING sebagai PENDUKUNG SISTEM KEPUTUSAN SELEKSI PENERIMA BANTUAN SOSIAL: Desa Sukau Kayo, Lebong, Bengkulu, Indonesia," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3S1, 2024, doi: 10.23960/jitet.v12i3s1.5164.