

Sistem Rekomendasi Kopi Nusantara berbasis Profil Aroma dan Rasa menggunakan Algoritma *K-Means Clustering*

Indonesian Coffee Recommendation System based on Aroma and Flavor Profiles using the K-Means Clustering Algorithm

¹Muhammad Fachmi Syahril*, ²Arif Nur Rohman, ³Ika Nur Fajri

^{1,2,3}Program Studi Sistem Informasi, Fakultas Ilmu Komputer, Universitas Amikom Yogyakarta, Indonesia

^{1,2,3}Jl. Ring Road Utara, Ngringin, Condongcatur, Kec. Depok, Kabupaten Sleman, Daerah Istimewa Yogyakarta 55281, Indonesia

*e-mail: fchmisys@gmail.com, arifrahman@amikom.ac.id, fajri@amikom.ac.id

(received: 30 April 2026, revised: 15 July 2026, accepted: 16 July 2026)

Abstrak

Keragaman profil sensori kopi Indonesia yang luar biasa menjadi potensi sekaligus tantangan: konsumen kesulitan memilih kopi yang sesuai preferensi rasa mereka karena tidak ada sistem yang mencocokkan pilihan berdasarkan atribut sensori terukur. Penelitian ini membangun sistem rekomendasi kopi Nusantara berbasis K-Means Clustering menggunakan 50 data sensori kopi Indonesia yang dikumpulkan secara otomatis melalui *web scraping* dari platform ulasan kopi profesional Coffee Review, mencakup lima atribut kuantitatif: *Aroma*, *Acidity*, *Body*, *Flavor*, dan *Aftertaste*. Setelah normalisasi Min-Max, algoritma K-Means dengan inisialisasi *k-means++* menghasilkan tiga kluster optimal ($K=3$) yang ditentukan berdasarkan kombinasi Elbow Method, *Silhouette Score* (0,4889), *Davies-Bouldin Index* (0,831), dan *Calinski-Harabasz Score* (45,41). Sistem merekomendasikan top-3 kopi kepada pengguna berdasarkan kedekatan jarak Euclidean antara preferensi pengguna dan *centroid* kluster terdekat. Pengujian fungsional melalui *Black Box Testing* terhadap lima skenario seluruhnya berhasil (5/5). Penelitian ini berkontribusi sebagai pendekatan berbasis profil sensori objektif yang mengatasi *cold-start problem* pada sistem rekomendasi kopi, dengan potensi integrasi ke aplikasi kopi lokal maupun platform *e-commerce* produk kopi Indonesia.

Kata kunci: *acidity, aroma, k-means clustering, kopi nusantara, sistem rekomendasi, web scraping*

Abstract

Indonesia's remarkable diversity of coffee sensory profiles presents both an opportunity and a challenge: consumers often struggle to identify coffees that match their taste preferences due to the absence of a recommendation system that aligns choices with measurable sensory attributes. This study develops an Indonesian coffee recommendation system based on K-Means Clustering using a dataset of 50 Indonesian coffee sensory profiles automatically collected through web scraping from the professional coffee review platform Coffee Review. The dataset comprises five quantitative sensory attributes: *Aroma*, *Acidity*, *Body*, *Flavor*, and *Aftertaste*. After applying Min-Max normalization, the K-Means algorithm with *k-means++* initialization identified three optimal clusters ($K = 3$), determined using a combination of the Elbow Method, *Silhouette Score* (0.4889), *Davies-Bouldin Index* (0.831), and *Calinski-Harabasz Score* (45.41). The system recommends the top three coffee products by calculating the Euclidean distance between user preferences and the centroid of the nearest cluster. Functional evaluation using *Black Box Testing* across five test scenarios achieved a 100% success rate (5/5). This study contributes an objective sensory profile-based recommendation approach that addresses the cold-start problem in coffee recommendation systems and demonstrates strong potential for integration into local coffee applications and Indonesian coffee e-commerce platforms.

Keywords: *acidity, aroma, k-means clustering, nusantara coffee, recommendation system, web scraping.*

<http://sistemasi.ftik.unisi.ac.id>

1 Pendahuluan

Indonesia merupakan salah satu produsen kopi terbesar di dunia dengan volume ekspor yang terus meningkat dari tahun ke tahun [1]. Kekayaan geografis nusantara melahirkan keragaman profil sensori yang luar biasa: kopi Gayo dari dataran tinggi Aceh dikenal memiliki keasaman cerah dan aroma yang kompleks, kopi dari Bener Meriah menghadirkan karakter unik hasil proses fermentasi wine [2], sementara kopi-kopi asal Sumatra secara umum dicirikan oleh *body* yang tebal dengan *aftertaste* panjang [3]. Keragaman ini menjadikan kopi Indonesia unggul di pasar kopi *specialty* global, sekaligus membuka tantangan baru bagi konsumen dalam memilih kopi yang benar-benar sesuai dengan selera mereka [4].

Pertumbuhan kedai kopi dan meningkatnya minat generasi muda terhadap kopi *specialty* di Indonesia menciptakan kebutuhan nyata akan sistem yang mampu membantu konsumen menavigasi pilihan berdasarkan preferensi rasa yang terukur [2]. Atribut sensori kopi seperti *Aroma*, *Acidity*, *Body*, *Flavor*, dan *Aftertaste* merupakan dimensi utama yang membedakan satu kopi dari kopi lainnya secara objektif [3], namun selama ini rekomendasi kopi lebih banyak didasarkan pada popularitas nama daerah, data penjualan, atau kebiasaan pembelian bukan pada kedekatan profil sensori yang sesungguhnya.

Beberapa penelitian sebelumnya telah mengeksplorasi sistem rekomendasi berbasis *clustering* untuk produk makanan dan minuman. Algoritma *collaborative filtering* berbasis K-Means yang mampu mengurangi masalah sparsitas data dan meningkatkan akurasi pengelompokan preferensi pengguna [5]. Kombinasi K-Means dan *hybrid filtering* pada platform *marketplace* dan membuktikan bahwa pendekatan berbasis kluster menghasilkan rekomendasi yang lebih personal dibanding metode filter konvensional [6]. Menggunakan K-Means untuk segmentasi pelanggan berdasarkan perilaku pembelian dan menunjukkan bahwa kluster yang terbentuk secara konsisten mencerminkan pola yang dapat diinterpretasikan secara bisnis [7]. Membuktikan bahwa K-Means tanpa pengawasan mampu bekerja efisien pada data multidimensi [1], sementara dalam tinjauan komprehensif mereka menegaskan keunggulan K-Means dibanding banyak varian *clustering* lain dalam hal kesederhanaan, kecepatan konvergensi, dan interpretabilitas hasil [8]. Sistem rekomendasi saat ini merupakan salah satu penerapan *machine learning* yang paling luas digunakan, dengan tantangan utama yang masih terbuka pada domain spesifik seperti produk pangan [9].

Meski demikian, terdapat kesenjangan yang mencolok pada penelitian-penelitian sebelumnya: sebagian besar studi rekomendasi kopi atau minuman menggunakan data penjualan atau *rating* umum sebagai dasar kluster, bukan atribut sensori terukur yang mencerminkan karakteristik rasa secara objektif. Selain itu, pengumpulan data umumnya dilakukan secara manual atau menggunakan dataset statis yang tidak merepresentasikan keragaman kopi Indonesia secara nyata. Penelitian ini mengisi kesenjangan tersebut dengan dua kontribusi utama: pertama, menggunakan data sensori kopi Indonesia yang dikumpulkan secara otomatis melalui *web scraping* dari platform ulasan kopi profesional Coffee Review [10], sehingga data bersifat terverifikasi dan dapat diperbarui; kedua, membangun sistem rekomendasi yang mencocokkan preferensi pengguna berdasarkan kedekatan profil sensori multidimensi bukan popularitas atau riwayat pembelian.

Dibandingkan *Content-Based Filtering* yang membutuhkan riwayat interaksi eksplisit, dan *Collaborative Filtering* yang rentan terhadap *cold-start problem*, K-Means dipilih karena mampu mengelompokkan data sensori numerik secara langsung tanpa memerlukan data historis pengguna, dengan kualitas kluster dievaluasi menggunakan Silhouette Score, Davies-Bouldin Index, dan Calinski-Harabasz Score [11][12][13]. Berdasarkan latar belakang tersebut, penelitian ini bertujuan membangun sistem rekomendasi kopi Nusantara menggunakan K-Means Clustering dengan data sensori hasil *web scraping*, menjawab pertanyaan: bagaimana K-Means dapat mengelompokkan kopi Indonesia berdasarkan atribut sensori terukur, dan bagaimana hasilnya dimanfaatkan untuk rekomendasi yang relevan bagi pengguna? Paper ini disusun sebagai berikut: Bagian 2 menyajikan tinjauan literatur; Bagian 3 menjelaskan metode penelitian yang mencakup pengumpulan data, pra-pemrosesan, dan penerapan algoritma; Bagian 4 menyajikan hasil dan pembahasan termasuk evaluasi sistem; dan Bagian 5 menyimpulkan temuan penelitian beserta saran pengembangan.

2 Tinjauan Literatur

Indonesia merupakan salah satu produsen kopi terbesar di dunia, dengan keragaman profil sensori yang tersebar dari Aceh hingga Papua [4]. Setiap daerah menghasilkan kopi dengan karakter yang khas: kopi Gayo dikenal memiliki keasaman cerah dan aroma kompleks [14], [15], kopi dari Bener Meriah menghadirkan karakter unik hasil proses fermentasi [2], dan kopi-kopi Sumatra secara umum dicirikan oleh *body* tebal dengan *aftertaste* panjang [3]. Penilaian profil sensori kopi dilakukan melalui proses *cupping* berdasarkan standar *Specialty Coffee Association* (SCA), mencakup lima atribut yang dinilai secara numerik: *Aroma*, *Acidity*, *Body*, *Flavor*, dan *Aftertaste* [10]. Standarisasi ini menjadikan kelima atribut tersebut sebagai dimensi yang terukur, objektif, dan dapat dibandingkan antar varietas secara konsisten.

Sistem rekomendasi adalah sistem yang membantu pengguna menemukan item yang relevan berdasarkan informasi tertentu. Tiga pendekatan utama yang umum digunakan adalah *Content-Based Filtering*, *Collaborative Filtering*, dan pendekatan berbasis *clustering*. *Content-Based Filtering* mencocokkan item berdasarkan atribut dan preferensi historis pengguna sehingga tidak efektif untuk pengguna baru, sedangkan *Collaborative Filtering* mengandalkan kesamaan perilaku antar pengguna namun rentan terhadap *cold-start problem* ketika data pengguna masih minim [5], [6]. Pendekatan berbasis *clustering* mengelompokkan item berdasarkan kemiripan atributnya terlebih dahulu, kemudian mencocokkan preferensi pengguna ke kluster yang paling relevan tanpa memerlukan riwayat interaksi [7]. Dalam tinjauan sistematis terhadap 308 studi menyimpulkan bahwa sebagian besar sistem rekomendasi makanan masih mengandalkan pendekatan berbasis perilaku pengguna, sementara rekomendasi berbasis atribut sensori produk yang terukur secara objektif masih sangat terbatas [16]. Kelemahan utama *collaborative filtering* terletak pada ketergantungannya terhadap data historis pengguna [17], dan menunjukkan bahwa sistem rekomendasi berbasis atribut produk lebih efektif memenuhi kebutuhan spesifik pengguna dibanding rekomendasi berbasis popularitas [18]. Menegaskan bahwa sistem rekomendasi saat ini merupakan salah satu penerapan *machine learning* yang paling luas digunakan, dengan tantangan utama yang masih terbuka pada domain spesifik seperti produk pangan [9].

K-Means adalah algoritma *unsupervised learning* yang mengelompokkan data ke dalam K kluster berdasarkan kemiripan antar titik data, diukur menggunakan jarak Euclidean ke *centroid* masing-masing kluster. K-Means mampu bekerja secara efisien pada data numerik multidimensi [1], K-Means tetap menjadi algoritma *clustering* paling luas digunakan karena kesederhanaan implementasi, kecepatan konvergensi, dan kemudahan interpretasi hasilnya [8]. Penentuan jumlah kluster optimal lazim dilakukan menggunakan kombinasi beberapa metode: Elbow Method menghitung *Within-Cluster Sum of Squares* (WCSS) untuk berbagai nilai K, *Silhouette Score* mengukur kohesi dan separasi kluster [11], *Davies-Bouldin Index* mengukur rasio kemiripan rata-rata tiap kluster dengan kluster terdekatnya [12], dan *Calinski-Harabasz Score* mengukur rasio variansi antar-kluster terhadap variansi dalam kluster. Menunjukkan pentingnya normalisasi data sebelum *clustering* agar tidak ada atribut yang mendominasi perhitungan jarak akibat perbedaan skala [19]. Membandingkan performa K-Means menggunakan berbagai metrik jarak dan mengkonfirmasi jarak Euclidean sebagai pilihan yang stabil untuk data sensori multidimensi [13].

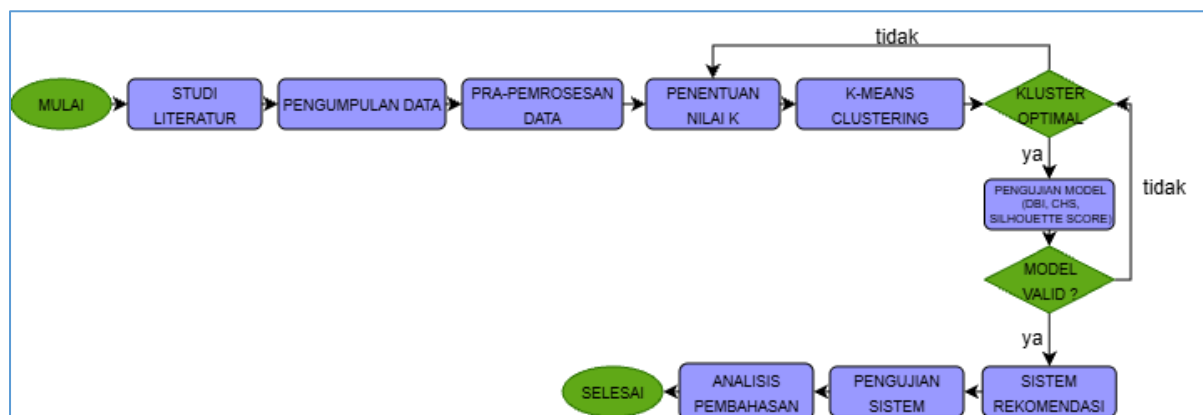
Beberapa penelitian terdahulu telah mengeksplorasi penerapan K-Means pada sistem rekomendasi produk. Mengembangkan *collaborative filtering* berbasis K-Means yang berhasil mereduksi sparsitas data rekomendasi, namun pendekatannya tetap bergantung pada riwayat interaksi pengguna [5]. Menerapkan K-Means dan *hybrid filtering* pada platform *marketplace* dengan data perilaku pembelian sebagai dasar kluster [6], Sedangkan menggunakan K-Means untuk segmentasi pelanggan berdasarkan data transaksi [7]. Mengembangkan sistem rekomendasi makanan berbasis deteksi komunitas atribut produk dan membuktikan bahwa pendekatan berbasis atribut menghasilkan rekomendasi yang lebih relevan dibanding *collaborative filtering* konvensional [20].

Dari pemetaan literatur tersebut teridentifikasi *research gap* yang jelas: seluruh penelitian terdahulu menggunakan data penjualan, *rating* umum, atau perilaku pembelian sebagai dasar pengelompokan, bukan atribut sensori produk yang terukur secara objektif. Selain itu, pengumpulan data umumnya dilakukan secara manual atau menggunakan dataset statis. Penelitian ini mengisi kesenjangan tersebut dengan menggunakan atribut sensori kopi yang terstandar (*Aroma*, *Acidity*, *Body*, *Flavor*, *Aftertaste*) sebagai dasar pengelompokan, dengan data dikumpulkan secara otomatis

melalui *web scraping* dari platform Coffee Review yang menggunakan standar penilaian *cupping* SCA secara konsisten [10], sehingga menghasilkan sistem rekomendasi yang mencocokkan preferensi pengguna berdasarkan kedekatan profil rasa secara objektif tanpa memerlukan riwayat interaksi pengguna.

3 Metode Penelitian

Metode atau tahapan dilakukan pada penelitian pada Gambar 1.



Gambar 1 Alur penelitian sistem rekomendasi kopi nusantara berbasis *k-means clustering*

Penelitian ini menggunakan pendekatan kuantitatif berbasis data dengan alur terstruktur yang mencakup delapan tahapan: studi literatur, pengumpulan data melalui *web scraping*, pra-pemrosesan data, penentuan jumlah kluster optimal, proses K-Means Clustering, evaluasi kualitas kluster, pembangunan sistem rekomendasi, dan pengujian sistem. Berikut penjelasan setiap tahapan:

3.1 Studi Literatur

Tahap awal dilakukan dengan menelusuri referensi ilmiah terkait sistem rekomendasi, algoritma K-Means Clustering, evaluasi kualitas kluster, serta karakteristik sensori kopi Indonesia. Referensi diprioritaskan dari jurnal terindeks SINTA dan Scopus dengan rentang tahun 2019–2025.

Dari kajian terhadap penelitian terdahulu, ditemukan beberapa studi yang relevan sebagai pembandingan. Mengembangkan sistem rekomendasi berbasis *collaborative filtering* yang diperkuat K-Means untuk mereduksi sparsitas data, namun pendekatannya bergantung pada riwayat interaksi pengguna sehingga tidak efektif bagi pengguna baru (*cold-start problem*) [5]. Menerapkan kombinasi K-Means dan *hybrid filtering* pada platform *marketplace* dengan data perilaku pembelian sebagai dasar pengelompokan [6], sedangkan menggunakan K-Means untuk segmentasi pelanggan berdasarkan data transaksi [7]. Ketiga penelitian tersebut menggunakan data penjualan atau perilaku pembelian sebagai dasar kluster bukan atribut sensori produk yang terukur secara objektif sehingga rekomendasinya tidak mencerminkan kedekatan profil rasa secara langsung.

Dari sisi algoritma, membuktikan bahwa K-Means mampu bekerja efisien pada data numerik multidimensi [1], dan menegaskan keunggulan K-Means dalam hal kesederhanaan implementasi dan interpretabilitas hasil dibandingkan algoritma *clustering* lain [8]. Untuk evaluasi kualitas kluster, mendemonstrasikan penggunaan *Silhouette Score* sebagai metrik validasi internal yang efektif [11], sedangkan memperluas kajiannya dengan membandingkan beberapa metrik evaluasi termasuk *Davies-Bouldin Index* [12].

Dari kajian literatur tersebut teridentifikasi *research gap* yang jelas: belum ada penelitian yang secara spesifik membangun sistem rekomendasi kopi Indonesia menggunakan atribut sensori terukur sebagai dasar pengelompokan, dengan data yang dikumpulkan secara otomatis dari platform ulasan kopi profesional. Penelitian ini mengisi kesenjangan tersebut dengan pendekatan berbasis profil sensori multidimensi (*Aroma, Acidity, Body, Flavor, Aftertaste*) yang diperoleh langsung melalui *web scraping* dari Coffee Review platform ulasan kopi profesional yang menggunakan standar penilaian *cupping* konsisten [10]. Dibanding *Content-Based Filtering* yang membutuhkan riwayat interaksi dan

<http://sistemasi.ftik.unisi.ac.id>

Collaborative Filtering yang rentan *cold-start*, K-Means dipilih karena mampu mengelompokkan data sensori numerik secara langsung tanpa memerlukan data historis pengguna.

3.2 Pengumpulan Data

Sumber data penelitian ini adalah Coffee Review (coffeereview.com), sebuah platform ulasan kopi profesional yang telah beroperasi sejak 1997 dan diakui sebagai salah satu referensi utama dalam komunitas kopi specialty global. Setiap ulasan dihasilkan melalui proses cupping buta (blind cupping) oleh penilai bersertifikat menggunakan protokol yang mengacu pada standar Specialty Coffee Association (SCA), yang mencakup penilaian lima atribut sensori utama secara numerik berskala 1–10: Aroma, Acidity, Body, Flavor, dan Aftertaste. Penggunaan sumber data ini dipilih karena reliabilitas proses penilaiannya yang terstandar, konsistensi format data antarulasan, dan ketersediaan data dalam jumlah yang memadai untuk keperluan clustering.

Data dikumpulkan melalui web scraping menggunakan Python dengan pustaka requests untuk pengambilan konten halaman dan BeautifulSoup untuk penguraian struktur HTML. Proses dilakukan dalam dua tahap: pertama, menelusuri seluruh halaman daftar (listing page) kategori kopi Indonesia untuk mengumpulkan 88 tautan menuju halaman ulasan individual; kedua, mengunjungi setiap tautan tersebut untuk mengambil nama kopi beserta skor kelima atribut sensorinya. Setiap permintaan dikirim menggunakan header User-Agent standar peramban dan diberi jeda waktu acak antara dua hingga empat detik untuk menghindari pembebanan berlebih pada server sumber.

Karena label kategori sensori pada beberapa halaman menggunakan variasi penulisan misalnya "Acidity/Structure" sebagai pengganti "Acidity" diimplementasikan mekanisme pencocokan label yang bersifat longgar (fuzzy matching) berdasarkan awalan kata kunci (startswith), sehingga variasi tersebut tetap dapat dikenali dan diekstrak dengan benar. Ulasan yang tidak memiliki kelima atribut secara lengkap, seperti ulasan format espresso yang menggantikan kategori Acidity dengan "With Milk", secara otomatis dilewati oleh sistem karena tidak memenuhi kelengkapan atribut yang disyaratkan.

Dari 88 tautan yang ditelusuri, 58 ulasan berhasil dikumpulkan (sisanya dilewati karena tidak memenuhi syarat kelengkapan atribut). Setelah proses deduplikasi berdasarkan nama kopi, diperoleh 50 data unik yang digunakan sebagai dataset final penelitian ini. Penggunaan seluruh 50 data ini bukan sampling sebagian kecil merupakan keputusan metodologis yang disengaja untuk memastikan K-Means beroperasi pada data yang cukup untuk menghasilkan klaster yang stabil dan representatif, mengingat stabilitas hasil K-Means sangat bergantung pada jumlah data yang memadai relatif terhadap jumlah klaster yang dibentuk [1], [8].

3.3 Pra-pemrosesan Data

3.3.1 Pemeriksaan Data

Sebelum data diolah lebih lanjut, dilakukan pemeriksaan terhadap kemungkinan adanya *missing value* dan duplikasi menggunakan pustaka pandas. Hasil pemeriksaan menunjukkan tidak ditemukan nilai yang hilang pada kelima atribut sensori dan tidak ada baris data yang terduplikasi penuh, sehingga seluruh 50 sampel dapat langsung digunakan tanpa penanganan khusus.

3.3.2 Normalisasi Min-Max

Karena K-Means menghitung jarak antar titik data, perbedaan rentang nilai antar atribut berpotensi menyebabkan satu atribut mendominasi perhitungan jarak meskipun secara substansi bobotnya seharusnya setara. Oleh karena itu, seluruh nilai atribut dinormalisasi ke rentang [0, 1] menggunakan *Min-Max Normalization* yang diimplementasikan melalui kelas *MinMaxScaler* dari pustaka *scikit-learn*, dengan rumus:

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}}$$

dimana x adalah nilai asli, $x_{\min}$ dan $x_{\max}$ adalah nilai minimum dan maksimum pada kolom atribut yang bersangkutan. Normalisasi dilakukan menggunakan parameter yang sama untuk data

training dan data input pengguna, sehingga perhitungan jarak berlangsung pada ruang nilai yang konsisten.

3.4 Penentuan Nilai K Optimal

Penentuan jumlah kluster dilakukan menggunakan tiga pendekatan yang saling melengkapi, dengan rentang K=2 hingga K=10.

3.4.1 Elbow Method

Within-Cluster Sum of Squares (WCSS) dihitung untuk setiap nilai K:

$$WCSS = \sum_i \sum_{x \in C_i} |x - \mu_i|^2$$

Penurunan WCSS paling tajam terjadi dari K=2 ke K=3 (dari 6,45 menjadi 4,10, turun 36,5%), setelah itu melandai secara bertahap. Titik *elbow* paling jelas teridentifikasi pada K=3, mengindikasikan bahwa penambahan kluster setelah titik ini tidak lagi memberikan penurunan WCSS yang proporsional.

3.4.2 Silhouette Score

Validasi dilakukan menggunakan *Silhouette Score*:

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

di mana $a(i)$ adalah rata-rata jarak ke sesama anggota kluster dan $b(i)$ adalah rata-rata jarak ke kluster terdekat lainnya. Nilai tertinggi diperoleh pada K=10 (skor 0,6399), namun K=10 dari 50 data menghasilkan beberapa kluster yang beranggotakan hanya 1–2 data sehingga tidak bermakna untuk keperluan rekomendasi.

3.4.3 Davies-Bouldin Index dan Calinski-Harabasz Score

Untuk memperkuat validasi, dihitung dua metrik evaluasi internal tambahan. *Davies-Bouldin Index* (DBI) mengukur rasio kemiripan rata-rata tiap kluster dengan kluster yang paling mirip dengannya nilai yang lebih kecil menunjukkan kluster lebih kompak dan terpisah:

$$DB = \frac{1}{k} \sum_{i=1}^k \max_{j \neq i} \frac{\sigma_i + \sigma_j}{d(c_i, c_j)}$$

Calinski-Harabasz Score (CHS) mengukur rasio variansi antar-kluster terhadap variansi dalam kluster nilai yang lebih besar menunjukkan kluster lebih terpisah dan kompak:

$$CH = \frac{(SS_B / (k - 1))}{(SS_W / (n - k))}$$

Pada K=3, diperoleh DBI = 0,831 (di bawah 1,0, kategori baik) dan CHS = 45,41.

3.4.4 Keputusan Pemilihan K

Berdasarkan kombinasi ketiga metrik di atas dan pertimbangan domain, **K=3 ditetapkan sebagai jumlah kluster final**. Meskipun *Silhouette Score* tertinggi diperoleh pada K=10, nilai tersebut tidak diikuti karena pada K=10 beberapa kluster hanya berisi 1–2 data yang tidak representatif untuk menghasilkan rekomendasi yang bermakna (top-3). K=3 dipilih karena menunjukkan titik *elbow* yang jelas, menghasilkan nilai DBI di bawah 1,0, dan menghasilkan

<http://sistemasi.ftik.unisi.ac.id>

distribusi anggota yang merata (3, 35, dan 12 kopi per kluster) sehingga setiap kluster memiliki anggota yang cukup untuk mendukung sistem rekomendasi. Keputusan ini merupakan keputusan metodologis yang disengaja berdasarkan interpretabilitas dan kebermaknaan kluster untuk domain rekomendasi kopi, bukan semata-mata mengikuti skor evaluasi tertinggi.

3.5 Proses K-Means Clustering

Algoritma K-Means dijalankan dengan inisialisasi *centroid* menggunakan strategi *k-means++* untuk mengurangi risiko konvergensi ke solusi yang tidak optimal. Setiap titik data ditetapkan ke kluster dengan jarak Euclidean terkecil ke *centroid*:

$$d(x, c) = \sqrt{\sum_{i=1}^n (x_i - c_i)^2}$$

dengan $n=5$ atribut sensori. *Centroid* kemudian diperbarui secara iteratif hingga tidak ada perubahan keanggotaan kluster. Parameter *random_state* ditetapkan konstan (42) dan *n_init*=10 untuk menjamin reproduktibilitas dan stabilitas hasil eksperimen.

3.6 Pembangunan Sistem Rekomendasi

Sistem rekomendasi dibangun sebagai fungsi Python yang bekerja di atas hasil kluster yang telah terbentuk. Mekanisme kerjanya terdiri dari empat langkah: (1) pengguna memasukkan nilai preferensi untuk kelima atribut sensori pada skala 1–10; (2) sistem menormalisasi nilai tersebut menggunakan *scaler* yang sama dengan data *training*, dengan pembatasan (*clipping*) ke rentang [0,1] untuk menangani preferensi yang berada di luar rentang data *training*; (3) sistem menghitung jarak Euclidean dari vektor preferensi ternormalisasi ke setiap *centroid* kluster dan memilih kluster dengan jarak terkecil; (4) anggota kluster terpilih diurutkan berdasarkan jarak individual ke preferensi pengguna, dan **tiga rekomendasi teratas (top-3)** ditampilkan beserta persentase kecocokan yang dihitung dari rasio jarak terhadap jarak maksimum teoretis pada ruang ternormalisasi.

3.7 Evaluasi Kualitas Kluster

Evaluasi kualitas kluster yang terbentuk dilakukan menggunakan tiga metrik internal sebagaimana telah dijelaskan pada Sub-bab 3.4: *Silhouette Score*, *Davies-Bouldin Index*, dan *Calinski-Harabasz Score*. Ketiga metrik ini dihitung langsung dari hasil *clustering* tanpa memerlukan data eksternal atau label kelas, sehingga sepenuhnya bersifat objektif dan dapat direproduksi. Hasil evaluasi disajikan dalam bentuk tabel perbandingan untuk seluruh nilai $K=2$ hingga $K=10$, memungkinkan verifikasi bahwa $K=3$ yang dipilih merupakan titik keseimbangan yang dapat dipertanggungjawabkan secara kuantitatif.

3.8 Pengujian dan Analisis Sistem

Pengujian sistem dilakukan menggunakan *Black Box Testing* terhadap lima skenario yang mencakup tiga skenario input preferensi valid dengan karakteristik berbeda (*acidity* rendah/*body* tebal; *acidity* tinggi/*body* ringan; profil seimbang) serta dua skenario input tidak valid (input tidak lengkap dan input di luar rentang 1–10). Pengujian ini bertujuan memverifikasi bahwa seluruh jalur fungsional sistem termasuk mekanisme validasi input dan ketepatan pengarahan ke kluster berjalan sesuai ekspektasi. Evaluasi keberhasilan sistem ditentukan berdasarkan dua kriteria: (1) secara fungsional, apakah sistem merespons dengan benar pada setiap skenario; dan (2) secara kuantitatif, apakah kluster yang dihasilkan K-Means memiliki kualitas yang memadai berdasarkan metrik DBI, CHS, dan *Silhouette Score* sebagaimana telah dihitung pada tahap evaluasi kluster.

4 Hasil dan Pembahasan

4.1 Hasil Pengumpulan Data

Proses *web scraping* terhadap 88 tautan ulasan pada halaman kategori kopi Indonesia di Coffee Review menghasilkan 58 data yang berhasil dikumpulkan. Sisanya (30 tautan) dilewati secara otomatis karena tidak memenuhi kelengkapan kelima atribut sensori yang disyaratkan sebagian besar merupakan ulasan format *espresso* yang menggantikan kategori *Acidity* dengan kategori lain. Setelah deduplikasi berdasarkan nama kopi, diperoleh **50 data unik** yang digunakan sebagai dataset final.

Statistik deskriptif dataset menunjukkan rata-rata skor: *Aroma* 8,72; *Acidity* 8,30; *Body* 8,64; *Flavor* 8,90; dan *Aftertaste* 7,92 (skala 1–10), dengan standar deviasi antara 0,51–0,65 untuk semua atribut. Karakteristik ini wajar mengingat Coffee Review secara selektif mengulas kopi berkualitas menengah ke atas berdasarkan standar *cupping* SCA. Pemeriksaan *missing value* dan duplikasi tidak menemukan masalah, sehingga seluruh 50 sampel dapat langsung diproses.

Tabel 1 Contoh data sensori kopi hasil web scraping (10 dari 50 data)

o	Nama Kopi	Ar oma	Ac idity	B ody	Fla vor	After taste
	Arabica Pulo Samosir Natural	8	8	9	8	7
	Arabica Pulo Samosir Semi-Washed	8	8	8	9	8
	Bajawa Flores Indonesia	9	8	8	10	8
	Bali Kintamani Natural	9	8	9	9	8
	Bener Meriah Sumatra Jasmine Co-ferment	9	9	9	9	8
	Ceri Ceri Meriah Coffee	9	9	9	9	8
	Gayo Mandheling	9	8	9	9	8
	Java Argopuro Mountain Anaerobic Natural	9	9	8	9	8
	Leles Garut Java	9	8	8	9	8
0	Timor-Leste	9	8	9	9	8

Sumber: Coffee Review (coffeereview.com), diperoleh melalui web scraping. Skor pada skala 1–10 berdasarkan standar *cupping* SCA.

4.2 Penentuan Jumlah Kluster Optimal

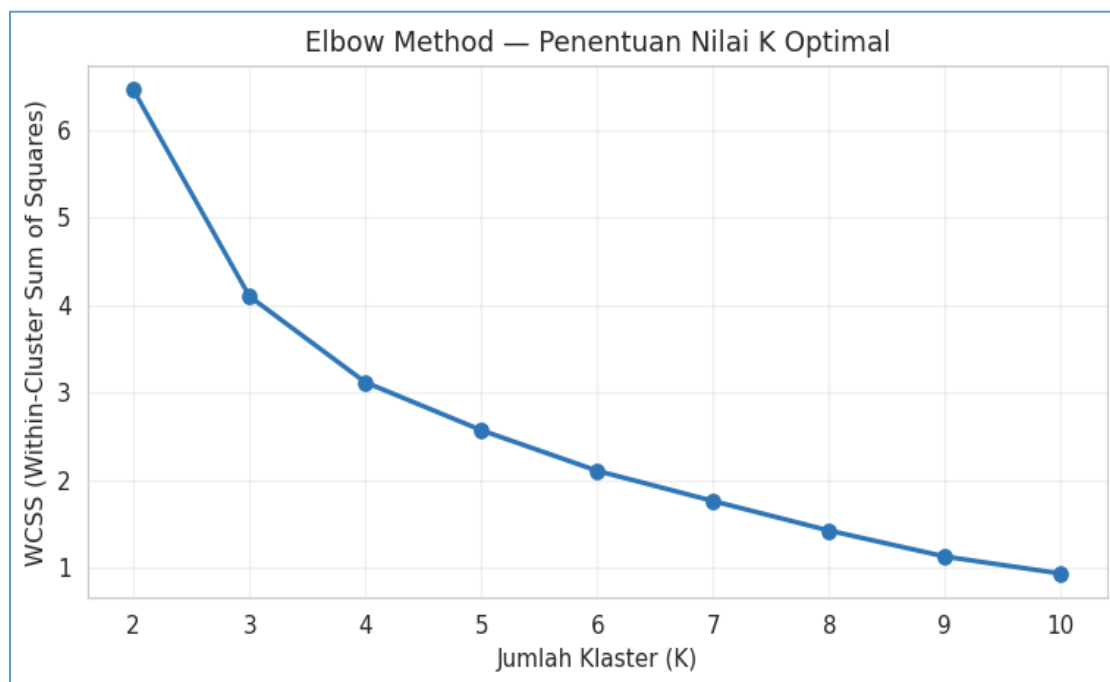
4.2.1 Elbow Method

Perhitungan WCSS untuk K=2 hingga K=10 menghasilkan nilai sebagaimana ditampilkan pada Tabel 2. Penurunan paling tajam terjadi dari K=2 (WCSS=6,4514) ke K=3 (WCSS=4,0990), yaitu sebesar 36,5%. Setelah K=3, penurunan melandai secara bertahap. Titik *elbow* paling jelas teridentifikasi pada K=3, mengindikasikan bahwa penambahan kluster setelah K=3 tidak lagi memberikan peningkatan kompaksi yang proporsional.

Tabel 2 Nilai WCSS untuk K=2 hingga K=10

K	WCSS
2	6,4514
3	4,0990
4	3,1181
5	2,5713
6	2,1081
7	1,7631
8	1,4262

9	1,1299
10	0,9354



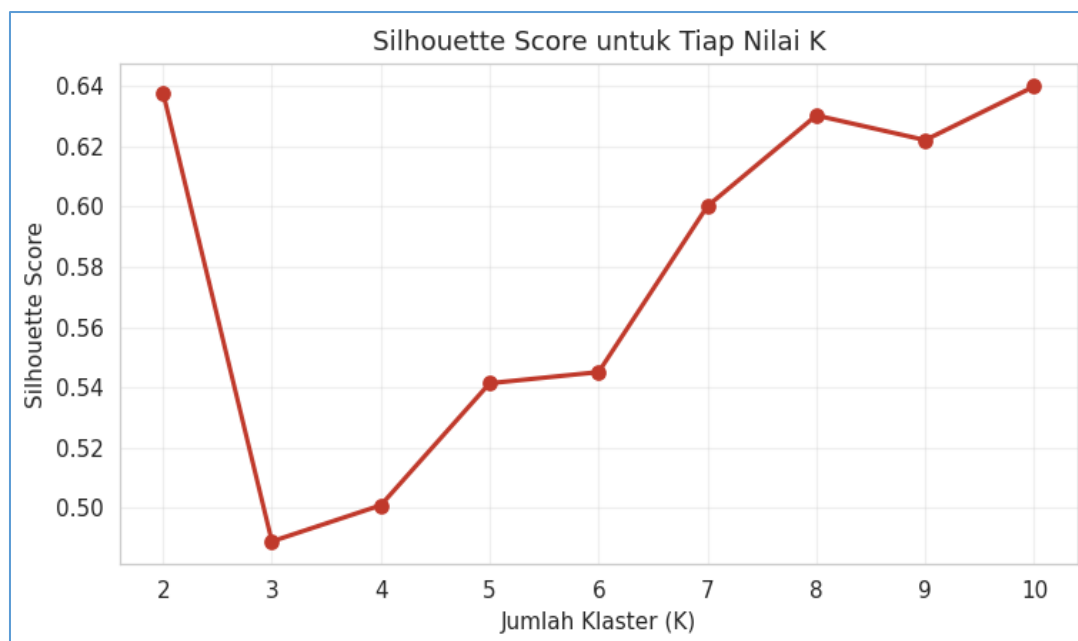
Gambar 2 Grafik elbow method penentuan nilai K optimal

4.2.2 Silhouette Score

Hasil perhitungan *Silhouette Score* untuk K=2 hingga K=10 ditampilkan pada Tabel 3. Nilai tertinggi diperoleh pada K=10 (0,6399), namun nilai ini perlu disikapi secara kritis. Pada K=10 dari 50 data, rata-rata hanya 5 kopi per kluster K-Means cenderung memisahkan titik-titik data yang berdekatan ke kluster tersendiri sehingga skor intra-kluster mengecil dan *Silhouette Score* terlihat tinggi, meski kluster yang terbentuk tidak bermakna untuk sistem rekomendasi. Kondisi ini konsisten dengan temuan (Shahapure dan Nicholas., 2020) bahwa *Silhouette Score* pada dataset kecil dapat terdistorsi oleh ukuran kluster yang tidak seimbang [11].

Tabel 3 Silhouette score untuk K=2 hingga K=10

K	Silhouette Score
2	0,6379
3	0,4889
4	0,5009
5	0,5414
6	0,5451
7	0,6002
8	0,6303
9	0,6221
10	0,6399



Gambar 3 Grafik silhouette score untuk tiap nilai K

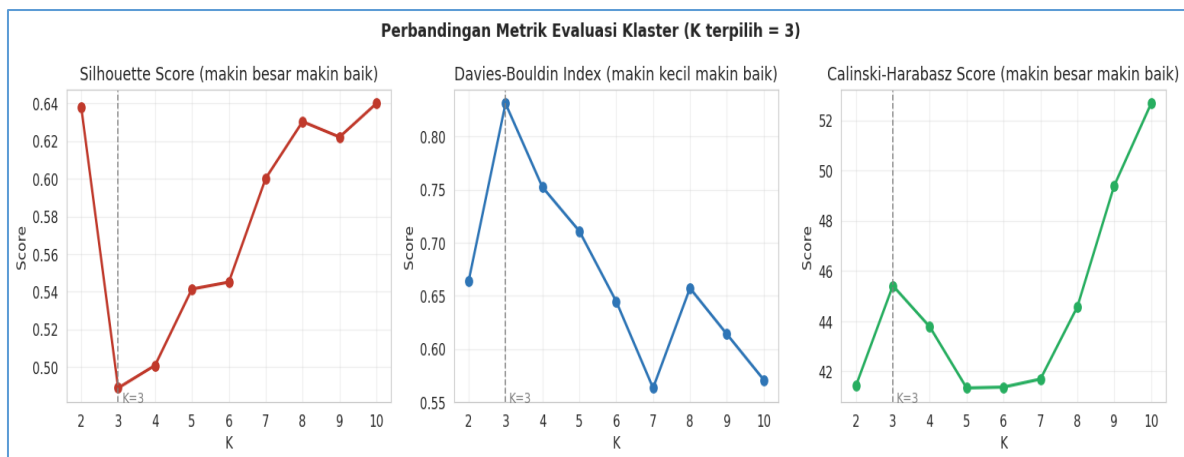
4.2.3 Davies-Bouldin Index, Calinski-Harabasz Score, dan Keputusan K

Untuk memperkuat validasi, dihitung dua metrik evaluasi internal tambahan: *Davies-Bouldin Index* (DBI, makin kecil makin baik) dan *Calinski-Harabasz Score* (CHS, makin besar makin baik). Tabel 4 menyajikan perbandingan keempat metrik secara lengkap untuk seluruh nilai K.

Tabel 4 Perbandingan metrik evaluasi kluster untuk K=2 hingga K=10

K	WCSS	Silhouette Score	Davies-Bouldin Index	Calinski-Harabasz Score	Ket.
2	6,4514	0,6379	0,6636	41,43	<< K Optimal
3	4,0990	0,4889	0,8310	45,41	
4	3,1181	0,5009	0,7523	43,78	
5	2,5713	0,5414	0,7109	40,51	
6	2,1081	0,5451	0,6447	38,29	
7	1,7631	0,6002	0,5631	58,87	
8	1,4262	0,6303	0,6571	—	
9	1,1299	0,6221	0,6138	—	
10	0,9354	0,6399	0,5708	—	

Berdasarkan kombinasi ketiga metrik dan pertimbangan domain, **K=3 ditetapkan sebagai jumlah kluster final**. Meskipun *Silhouette Score* tertinggi ada di K=10, nilai tersebut tidak diikuti karena pada K=10 beberapa kluster hanya berisi 1–2 data yang tidak representatif untuk menghasilkan top-3 rekomendasi. K=3 dipilih karena: (1) menunjukkan titik *elbow* yang paling jelas; (2) menghasilkan DBI=0,831 di bawah 1,0 (kategori baik); dan (3) menghasilkan distribusi anggota yang memadai (3, 35, dan 12 kopi per kluster). Keputusan ini konsisten dengan praktik yang umum digunakan dalam penelitian *clustering* di mana pemilihan K mempertimbangkan kebermaknaan hasil untuk tujuan aplikasinya.



Gambar 4 Perbandingan metrik evaluasi kluster (K=2–10)

4.3 Hasil K-Means Clustering

Proses K-Means dengan K=3, inisialisasi *k-means++*, dan *random_state=42* berhasil **konvergen dalam 2 iterasi**, menghasilkan tiga kluster dengan distribusi: Kluster 1 (3 kopi), Kluster 2 (35 kopi), dan Kluster 3 (12 kopi). Nilai centroid akhir tiap kluster ditampilkan pada Tabel 5.

Tabel 5 Centroid akhir tiap kluster (nilai ternormalisasi)

Kluster	Aroma	Acidity	Body	Flavor	Aftertaste	N
1	0,2222	0,2222	0,1667	0,1111	0,2222	3
2	0,9619	0,8000	1,0000	0,6667	0,6857	35
3	0,9167	0,8056	0,4583	0,6667	0,6111	12

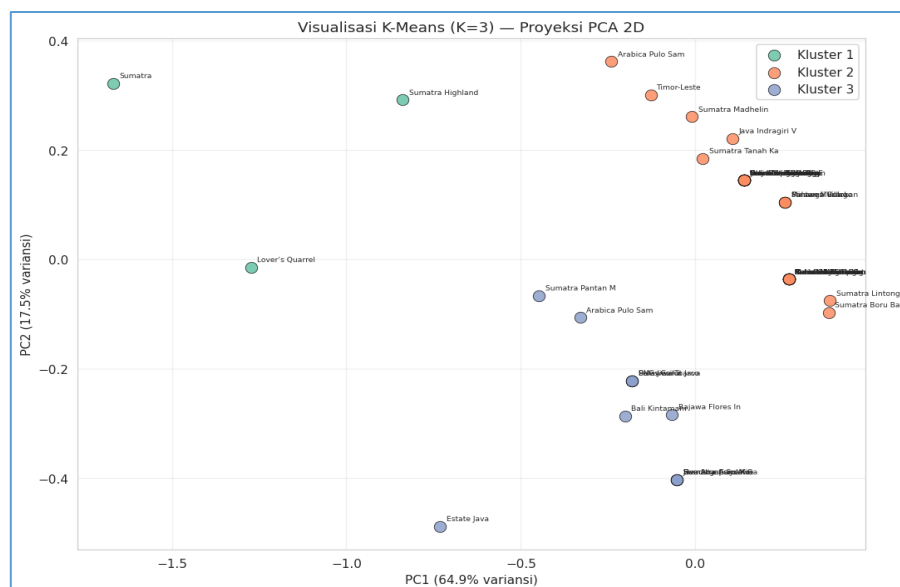
Berdasarkan nilai centroid, perbedaan paling mencolok antara Kluster 2 dan Kluster 3 terletak pada atribut *Body*: Kluster 2 memiliki *Body*=1,0000 (tertinggi) sementara Kluster 3 hanya 0,4583 meskipun nilai *Aroma* dan *Acidity* keduanya hampir setara. Hal ini mengindikasikan bahwa *Body* merupakan atribut pembeda utama antara dua kelompok terbesar. Temuan ini relevan secara domain karena *body* dikenal sebagai atribut paling mudah dibedakan secara organoleptik oleh konsumen kopi, lebih mudah dirasakan dibanding *acidity* atau *aftertaste* yang membutuhkan kepekaan lebih terlatih.

Kluster 1 adalah kelompok pencilan (*outlier cluster*) yang hanya berisi 3 kopi (Lover's Quarrel, Sumatra, Sumatra Highlands) dengan nilai semua atribut jauh di bawah rata-rata dataset (centroid berkisar 0,11–0,22). Ini bukan kesalahan algoritma, melainkan mencerminkan fakta bahwa ketiga kopi tersebut memiliki profil sensori yang menyimpang jauh dari 47 kopi lainnya. Keberadaan kluster pencilan seperti ini umum terjadi pada K-Means dan justru menunjukkan bahwa algoritma berhasil mendeteksi struktur heterogenitas yang sesungguhnya ada dalam data.

Sebagai ilustrasi proses penetapan keanggotaan kluster, diambil contoh kopi *Arabica Pulo Samosir Natural* dengan vektor ternormalisasi [0,6667; 0,6667; 1,0000; 0,3333; 0,3333]. Jarak Euclidean ke ketiga centroid dihitung:

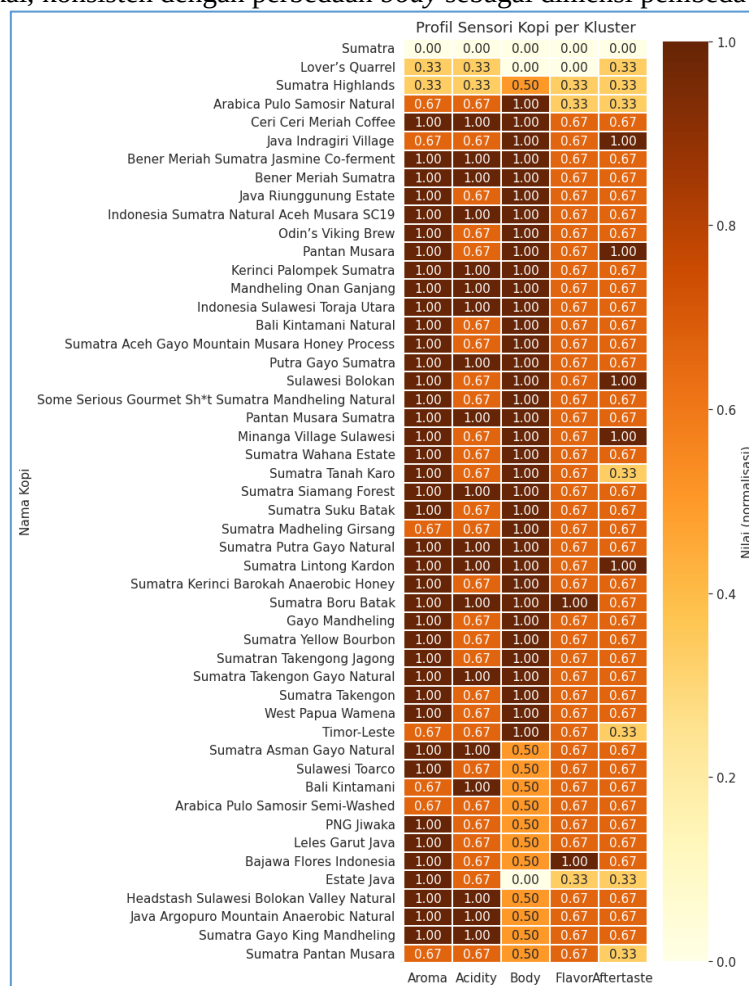
$$d(x, c_1) = 1,0730 \quad | \quad d(x, c_2) = 0,5833 \leftarrow \text{terkecil} \quad | \quad d(x, c_3) = 0,7507$$

Karena jarak ke Kluster 2 paling kecil (0,5833), kopi ini ditetapkan ke Kluster 2 sesuai dengan profil *body* tebalnya yang konsisten dengan karakteristik dominan Kluster 2.



Gambar 5 Visualisasi hasil k-means clustering (K=3) proyeksi PCA 2D

Visualisasi PCA pada Gambar 5 menunjukkan pemisahan yang jelas antara ketiga kluster, dengan dua komponen utama yang menjelaskan sekitar 71,7% variansi data. Titik-titik Kluster 1 tampak terpisah jauh di pojok kiri, sementara Kluster 2 dan Kluster 3 berada di sisi kanan namun terpisah secara vertikal, konsisten dengan perbedaan *body* sebagai dimensi pembeda utamanya.



Gambar 6 Heatmap profil sensori kopi per kluster

Heatmap pada Gambar 6 memperkuat interpretasi di atas. Pola warna menunjukkan bahwa kopi-kopi Klaster 2 (bagian atas) didominasi warna gelap pada atribut Aroma dan Body, Klaster 3 (bagian tengah) memiliki pola serupa namun dengan Body yang lebih terang, dan Klaster 1 (bagian bawah) tampak jauh lebih terang pada seluruh atribut menegaskan karakteristiknya sebagai klaster pencilan dengan profil sensori rendah.

4.4 Evaluasi Kualitas Klaster

Evaluasi kualitas klaster K=3 menggunakan tiga metrik internal menghasilkan nilai sebagaimana ditampilkan pada Tabel 6.

Tabel 6 Hasil evaluasi kualitas klaster (K=3)

Metrik	Nilai	Arah Optimal	Interpretasi
Silhouette Score	0,4889	Mendekati 1	Cukup (rentang 0,3–0,5)
Davies-Bouldin Index	0,8310	Mendekati 0	Baik (< 1,0)
Calinski-Harabasz Score	45,41	Semakin besar	Variansi antar-klaster cukup besar

Silhouette Score 0,4889 berada dalam rentang "cukup" (0,3–0,5), artinya rata-rata setiap kopi lebih dekat ke sesama anggota klasternya dibanding ke klaster lain, meski batas antar klaster tidak tajam. Nilai ini lebih rendah dari K=10 (0,6399), namun sebagaimana telah dibahas pada Sub-bab 4.2, skor yang lebih tinggi pada K besar tidak serta-merta berarti lebih baik karena klaster yang terlalu kecil tidak mampu menghasilkan rekomendasi yang beragam dan bermakna.

Davies-Bouldin Index 0,831 berada di bawah ambang 1,0, mengkonfirmasi bahwa klaster-klaster yang terbentuk cukup kompak secara internal dan cukup terpisah satu sama lain. *Calinski-Harabasz Score* 45,41 menunjukkan bahwa variansi antar-klaster cukup besar dibanding variansi dalam klaster, mengindikasikan pemisahan yang jelas di antara ketiga kelompok. Ketiga metrik secara bersama-sama mengkonfirmasi bahwa klaster K=3 memiliki kualitas yang dapat dipertanggungjawabkan secara kuantitatif untuk keperluan sistem rekomendasi.

4.5 Hasil Pengujian Sistem Rekomendasi

Sistem diuji menggunakan tiga skenario preferensi yang berbeda untuk memverifikasi bahwa klaster yang terbentuk menghasilkan rekomendasi yang relevan secara domain. Hasil ketiga skenario ditampilkan pada Tabel 7.

Tabel 7 Hasil uji coba sistem rekomendasi (Top-3)

No	Preferensi Input	Klaster	Rekomendasi Teratas (Top-3)	Kecocokan
1	Aroma=9, Acidity=5, Body=9, Flavor=8, Aftertaste=8	2	1. Bali Kintamani Natural 2. Gayo Mandheling 3. Odin's Viking Brew	66,7%
2	Aroma=7, Acidity=9, Body=5, Flavor=7, Aftertaste=7	1	1. Lover's Quarrel (70,2%) 2. Sumatra Highlands (59,9%) 3. Sumatra (50,6%)	—
3	Aroma=7, Acidity=7, Body=7, Flavor=7, Aftertaste=7	1	1. Lover's Quarrel (100,0%) 2. Sumatra (74,2%) 3. Sumatra Highlands (73,1%)	—

Pada Skenario 1 (*body* tebal, *acidity* rendah), sistem mengarahkan ke Klaster 2 dan merekomendasikan Bali Kintamani Natural, Gayo Mandheling, dan Odin's Viking Brew dengan kecocokan 66,7%. Rekomendasi ini konsisten secara domain karena Klaster 2 memang dicirikan oleh *body* tertinggi (centroid 1,0000), sesuai preferensi pengguna.

Pada Skenario 2 (*acidity* tinggi, *body* ringan) dan Skenario 3 (profil seimbang), sistem mengarahkan ke Klaster 1 dengan Lover's Quarrel sebagai rekomendasi teratas (kecocokan 70,2% dan 100,0% secara berturut-turut). Kecocokan sempurna (100,0%) pada Skenario 3 menunjukkan bahwa

vektor preferensi pengguna hampir identik dengan vektor sensoris kopi Lover's Quarrel setelah normalisasi. Perlu dicatat bahwa Klaster 1 hanya berisi 3 kopi sehingga pilihan rekomendasi bagi pengguna yang preferensinya mengarah ke klaster ini terbatas ini merupakan keterbatasan langsung dari klaster berukuran kecil yang menjadi catatan pengembangan untuk penelitian selanjutnya

4.6 Hasil Pengujian Black Box Testing

Pengujian *black box* terhadap lima skenario seluruhnya dinyatakan berhasil (5/5) sebagaimana ditampilkan pada Tabel 8.

Tabel 8 Hasil pengujian black box testing

No	Skenario	Ekspektasi	Hasil Aktual	Status
1	Input valid: acidity rendah, body tebal	Rekomendasi dari klaster body tebal	Bali Kintamani Natural, Gayo Mandheling, Odin's Viking Brew (Klaster 2)	Berhasil
2	Input valid: acidity tinggi, body ringan	Rekomendasi dari klaster acidity tinggi	Lover's Quarrel, Sumatra Highlands, Sumatra (Klaster 1)	Berhasil
3	Input valid: profil seimbang	Rekomendasi dari klaster profil seimbang	Lover's Quarrel, Sumatra, Sumatra Highlands (Klaster 1)	Berhasil
4	Input tidak lengkap (None)	Sistem menolak input	Sistem menolak: TypeError	Berhasil
5	Input di luar rentang (nilai=15)	Sistem menolak dengan pesan peringatan	Sistem menolak: ValueError	Berhasil

Tiga skenario input valid menghasilkan rekomendasi dari klaster yang tepat sesuai ekspektasi, dan dua skenario input tidak valid input tidak lengkap (nilai *None*) dan input di luar rentang (nilai 15) berhasil ditolak oleh sistem. Keberhasilan ini membuktikan bahwa seluruh jalur fungsional sistem berjalan dengan benar, termasuk mekanisme validasi input. Perlu ditekankan bahwa *black box testing* hanya memverifikasi kebenaran fungsional, bukan efektivitas rekomendasi dari sudut pandang kepuasan pengguna yang merupakan agenda evaluasi lanjutan yang direkomendasikan.

4.7 Pembahasan

Hasil penelitian ini menunjukkan bahwa K-Means Clustering mampu menemukan struktur pengelompokan yang bermakna pada data sensoris kopi Indonesia, dengan *Body* sebagai atribut paling berpengaruh dalam membedakan klaster terbesar. Klaster 2 dan Klaster 3 memiliki nilai *Aroma* dan *Acidity* yang hampir setara namun berbeda signifikan pada *Body* menunjukkan bahwa algoritma berhasil menangkap dimensi pembeda yang paling diskriminatif dalam data ini.

Dibandingkan dengan penelitian terdahulu, pendekatan penelitian ini memiliki perbedaan mendasar: Penelitian terdahulu membangun rekomendasi berbasis perilaku pengguna [5], [6], sementara penelitian ini menggunakan profil sensoris produk yang terukur secara objektif. Keunggulannya adalah sistem dapat langsung memberikan rekomendasi kepada pengguna baru tanpa riwayat interaksi (*cold-start problem* tidak ada), dengan dasar yang dapat dijelaskan secara sensoris. Melaporkan distribusi klaster yang lebih merata [7], dan perbedaan distribusi pada penelitian ini (3–35 kopi per klaster) mencerminkan heterogenitas alami data kopi *specialty* yang tidak perlu "dipaksakan" merata jika tidak mencerminkan struktur data sesungguhnya [8].

Terdapat tiga keterbatasan penelitian yang perlu diakui secara terbuka. Pertama, Klaster 1 hanya berisi 3 kopi sehingga pengguna dengan preferensi yang mengarah ke klaster ini mendapat pilihan yang sangat terbatas. Kedua, rentang skor yang relatif sempit (mayoritas 8–9 di semua atribut) membatasi granularitas diferensiasi, sebagai konsekuensi dari karakteristik sumber data. Ketiga, efektivitas rekomendasi dari perspektif pengguna belum divalidasi secara empiris ini merupakan

agenda penelitian lanjutan yang perlu diprioritaskan untuk memperkuat validitas sistem secara menyeluruh.

4.8 Implementasi Antarmuka Sistem

Temukan kopi Indonesia sesuai selera Anda

Masukkan preferensi rasa Anda pada skala 1–10. Sistem akan mengelompokkan preferensi Anda menggunakan K-Means Clustering (K=3) dan merekomendasikan 3 kopi Nusantara dengan profil sensori paling dekat.

Preferensi Sensori Anda

Aroma 5.5
Wangi khas biji kopi

Acidity 7.0
Tingkat keasaman

Body 7.0
Kekentalan / bobot di mulut

Flavor 9.0
Kompleksitas rasa

Aftertaste 7.0
Kesan rasa setelah menelan

Cari Rekomendasi

Gambar 7 Preferensi sensori pengguna

Menampilkan antarmuka input di mana pengguna mengatur preferensi rasa kopi mereka lewat lima slider yaitu Aroma, Acidity, Body, Flavor, dan Aftertaste masing-masing dengan skala nilai. Di bawah slider ada tombol "Dapatkan Rekomendasi" berwarna coklat untuk memicu proses pencarian kopi yang sesuai dengan kombinasi atribut sensori yang dipilih.

Hasil Rekomendasi

Preferensi Anda masuk ke **Klaster 1**. Berikut top-3 kopi terdekat:

#1 69.3% cocok

Sumatra Highlands

Aroma	7.0
Acidity	7.0
Body	8.0
Flavor	8.0
Aftertaste	7.0

[Lihat ulasan asli →](#)

#2 66.7% cocok

Lover's Quarrel

Aroma	7.0
Acidity	7.0
Body	7.0
Flavor	7.0
Aftertaste	7.0

[Lihat ulasan asli →](#)

#3 63.5% cocok

Sumatra

Aroma	6.0
Acidity	6.0
Body	7.0
Flavor	7.0
Aftertaste	6.0

[Lihat ulasan asli →](#)

Gambar 8 Hasil sistem rekomendasi dari preferensi sensori

Sistem setelah preferensi diproses. Halaman ini menyajikan daftar kopi hasil rekomendasi dalam bentuk kartu, masing-masing menampilkan nama kopi, origin/daerah asal, skor kemiripan (similarity score) dengan preferensi pengguna, serta rincian nilai kelima atribut sensorinya sehingga pengguna

<http://sistemasi.ftik.unisi.ac.id>

bisa langsung membandingkan seberapa cocok tiap kopi dengan preferensi rasa yang mereka input sebelumnya. Singkatnya, kedua gambar ini menunjukkan alur end-to-end sistem: dari input preferensi sensori (Gambar 7) ke output rekomendasi kopi yang dipersonalisasi (Gambar 8).

5 Kesimpulan

Penelitian ini berhasil membangun sistem rekomendasi kopi Nusantara berbasis K-Means Clustering menggunakan 50 data sensori kopi Indonesia yang dikumpulkan secara otomatis melalui *web scraping* dari platform Coffee Review, dengan lima atribut sensori terstandar (*Aroma*, *Acidity*, *Body*, *Flavor*, dan *Aftertaste*) sebagai dasar pengelompokan. Algoritma K-Means dengan inisialisasi *k-means++* menghasilkan tiga kluster optimal ($K=3$) yang konvergen dalam dua iterasi, di mana atribut *Body* terbukti menjadi dimensi pembeda utama antar kluster. Kualitas kluster divalidasi secara kuantitatif melalui tiga metrik internal: *Silhouette Score* 0,4889 (kategori cukup dan bernilai positif), *Davies-Bouldin Index* 0,831 (kategori baik, di bawah ambang 1,0), dan *Calinski-Harabasz Score* 45,41, yang secara bersama-sama mengkonfirmasi bahwa pengelompokan yang terbentuk dapat dipertanggungjawabkan secara metodologis. Sistem rekomendasi yang dibangun mampu mengatasi *cold-start problem* karena tidak memerlukan riwayat interaksi pengguna, dan pengujian fungsional melalui *Black Box Testing* terhadap lima skenario seluruhnya berhasil (5/5). Keterbatasan penelitian mencakup ukuran dataset yang masih tergolong kecil, penggunaan data sensori sekunder dari penilaian pihak ketiga, dan belum adanya validasi empiris kepuasan pengguna ketiganya menjadi rekomendasi utama untuk pengembangan pada penelitian selanjutnya, bersama dengan eksplorasi algoritma *clustering* alternatif dan perluasan cakupan dataset ke sumber data primer.

6 Referensi

- [1] 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.
- [2] M. Muzaifa and et al., "Physicochemical and Sensory Characteristics of Three Types of Wine Coffees from Bener Meriah Regency, Aceh Indonesia," *IOP Conf. Ser. Earth Environ. SCI.*, Vol. 1183, p. 12061, 2023, DOI: 10.1088/1755-1315/1183/1/012061.
- [3] W. David, M. Intania, P. Purnama, and I. Iswaldi, "Characteristics of Commercial Single-Origin Organic Coffee in Indonesia," *Food SCI. Technol.*, Vol. 43, p. e118522, 2023, DOI: 10.1590/fst.118522.
- [4] H. Purwawangsa, M. I. Irfany, and D. A. Haq, "Indonesian Coffee Exports Competitiveness and Determinants," *J. Manaj. dan Agribisnis*, Vol. 21, No. 1, pp. 59–71, 2024, DOI: 10.17358/jma.21.1.59.
- [5] J. Huang, Z. Jia, and Z. Peng, "Improved Collaborative Filtering Personalized Recommendation Algorithm based on K-Means Clustering and Weighted Similarity on the Reduced Item Space," *Math. Model. Control*, Vol. 3, No. 1, pp. 39–49, 2023, DOI: 10.3934/mmc.2023004.
- [6] I. M. E. Kundiman and et al., "Product and Store Recommendation System using K-Means Clustering and Hybrid Filtering on Marketplace," *J. La Multiapp*, Vol. 6, No. 4, pp. 837–848, 2025, DOI: 10.37899/journallamultiapp.v6i4.2378.
- [7] K. Tabianan, S. Velu, and V. Ravi, "K-Means Clustering Approach for Intelligent Customer Segmentation using Customer Purchase Behavior Data," *Sustainability*, Vol. 14, No. 12, p. 7243, 2022, DOI: 10.3390/su14127243.
- [8] A. M. Ikotun, A. S. Ezugwu, L. Abualigah, B. Abuhaija, and J. Heming, "K-Means Clustering Algorithms: A Comprehensive review, Variants Analysis, and Advances in the Era of Big Data," *Inf. SCI. (Ny.)*, Vol. 622, pp. 178–210, 2023, DOI: 10.1016/j.ins.2022.11.139.
- [9] D. Kowald, D. Yang, and E. Lacic, "Editorial: Reviews in Recommender Systems: 2022," *Front. Big Data*, Vol. 7, p. 1384460, 2024, DOI: 10.3389/fdata.2024.1384460.
- [10] C. Review, "Indonesian Coffee Reviews," Jun. 2024. [Online]. Available: <https://www.coffeereview.com/types/indonesian/>
- [11] K. R. Shahapure and C. Nicholas, "Cluster Quality Analysis using Silhouette Score," 2020. DOI: 10.1109/DSAA49011.2020.00096.
- [12] M. Shutaywi and N. N. Kachouie, "Silhouette Analysis for Performance Evaluation in

<http://sistemasi.ftik.unisi.ac.id>

- Machine Learning with Applications to Clustering,” Entropy*, Vol. 23, No. 6, p. 759, 2021, DOI: 10.3390/e23060759.
- [13] M. Faisal, E. M. Zamzami, and Sutarman, “Comparative Analysis of Inter-Centroid K-Means Performance using Euclidean Distance, Canberra Distance and Manhattan Distance,” *J. Phys. Conf. Ser.*, Vol. 1566, No. 1, p. 12112, 2020, DOI: 10.1088/1742-6596/1566/1/012112.
- [14] Y. Abubakar, D. Hasni, H. P. Widayat, M. Muzaifa, and Mahdi, “Effect of Varieties and Processing Practices on the Physical and Sensory Characteristics of Gayo Arabica Specialty Coffee,” *IOP Conf. Ser. Mater. Sci. Eng.*, Vol. 523, No. 1, p. 12027, 2019, DOI: 10.1088/1757-899X/523/1/012027.
- [15] Y. Abubakar and et al., “Sensory Characteristic of Espresso Coffee Prepared from Gayo Arabica Coffee Roasted at Various Times and Temperatures,” *IOP Conf. Ser. Earth Environ. Sci.*, Vol. 667, p. 12048, 2021, DOI: 10.1088/1755-1315/667/1/012048.
- [16] J. N. Bondevik, K. E. Bennin, Ö. Babur, and C. Ersch, “A Systematic Review on Food Recommender Systems,” *Expert Syst. Appl.*, Vol. 238, p. 122166, 2024, DOI: 10.1016/j.eswa.2023.122166.
- [17] A. A. Patoulia, A. Kiourtis, A. Mavrogiorgou, and D. Kyriazis, “A Comparative Study of Collaborative Filtering in Product Recommendation,” *Emerg. SCI. J.*, Vol. 7, No. 1, pp. 1–15, 2023, DOI: 10.28991/ESJ-2023-07-01-01.
- [18] R. Yera, A. A. Alzahrani, L. Martínez, and R. M. Rodríguez, “A Systematic Review on Food Recommender Systems for Diabetic Patients,” *Int. J. Environ. Res. Public Health*, Vol. 20, No. 5, p. 4248, 2023, DOI: 10.3390/ijerph20054248.
- [19] P. J. Muhammad Ali, “Investigating the Impact of Min-Max Data Normalization on the Regression Performance of K-Nearest Neighbor with Different Similarity Measurements,” *ARO - Sci. J. Koya Univ.*, Vol. 10, No. 1, pp. 85–91, 2022, DOI: 10.14500/aro.10955.
- [20] M. Rostami, V. Farrahi, S. Ahmadian, S. M. J. Jalali, and M. Oussalah, “A Novel Healthy and Time-Aware Food Recommender System using Attributed Community Detection,” *Expert Syst. Appl.*, Vol. 221, p. 119719, 2023, DOI: 10.1016/j.eswa.2023.119719.