

ANALISIS KOMPARATIF METODE *K-MEANS* DAN *HIERARCHICAL CLUSTERING* UNTUK SEGMENTASI PELANGGAN PUSAT KEBUGARAN

Adam Gusti Andito^{1*}, Nuramaliyah²

^{1,2}Program Studi Statistika, Fakultas Sains dan Teknologi, Universitas Terbuka.

*Penulis Korespondensi: adam.gustiandito@gmail.com

ABSTRAK

Tingginya angka berhentinya berlangganan pelanggan (*churn*) menjadi tantangan utama dalam industri pusat kebugaran karena berdampak langsung pada penurunan pendapatan dan peningkatan biaya akuisisi pelanggan baru. Penelitian ini membandingkan metode *K-Means* dan *Hierarchical Clustering* untuk segmentasi pelanggan berbasis karakteristik demografis dan perilaku. Menggunakan dataset sekunder sebanyak 4000 observasi dengan 14 variabel, analisis komparatif dilakukan melalui evaluasi empat metrik internal, uji stabilitas *Adjusted Rand Index* (ARI), dan efisiensi komputasi. Hasil penelitian menunjukkan *K-Means* secara konsisten unggul dibandingkan *Hierarchical Clustering*, dengan capaian *Silhouette Coefficient* 0.151, stabilitas yang tinggi (ARI 0.905), dan waktu eksekusi yang efisien (0.081 detik). Berdasarkan segmentasi *K-Means*, teridentifikasi dua klaster utama, dengan klaster musiman yaitu “*High-Risk Casual*” (31.2%), yaitu klaster pelanggan dengan frekuensi kunjungan rendah dan durasi kontrak singkat yang memiliki tingkat *churn* sebesar 84% dan klaster inti yaitu “*Low-Risk Regular*” (68.8%), yaitu klaster pelanggan loyal dengan aktivitas rutin dan kontrak jangka panjang yang memiliki tingkat *churn* sangat rendah sebesar 0.3%. Hasil ini merekomendasikan implementasi *K-Means* untuk sistem segmentasi *real-time* guna mendukung pengembangan strategi retensi yang efektif dalam industri pusat kebugaran.

Kata kunci: *K-Means*, *Hierarchical Clustering*, Segmentasi Pelanggan, Pusat Kebugaran, *Customer Churn*

1 PENDAHULUAN

Industri pusat kebugaran global termasuk di Indonesia menghadapi tantangan bisnis yang signifikan berupa tingginya angka *customer churn* atau berhentinya keanggotaan (Yeomans et al., 2024). Retensi pelanggan menjadi faktor penentu kelangsungan finansial sebuah pusat kebugaran. Untuk mengatasi tantangan ini, pusat kebugaran memerlukan strategi yang lebih cerdas dan personal, yang dapat dicapai melalui pemahaman mendalam tentang basis pelanggan mereka yang heterogen.

Segmentasi pelanggan muncul sebagai solusi strategis untuk memecah basis pelanggan yang luas menjadi kelompok-kelompok yang lebih kecil dan homogen berdasarkan kesamaan karakteristik (Ufeli et al., 2025). Pendekatan ini memungkinkan manajemen untuk mengalokasikan sumber daya secara lebih efisien dan merancang program pemasaran serta layanan yang lebih tepat sasaran (Niu, 2025). Oleh karena itu, kemajuan dalam bidang *data mining* dan *machine learning* menawarkan metode yang objektif dan terukur untuk melakukan segmentasi, salah satunya melalui teknik *clustering* (Kumar et al., 2025).

Dua metode *clustering* yang paling banyak digunakan adalah *K-Means* dan *Hierarchical Clustering* (Wardani et al., 2023). Dalam penerapannya, salah satu tantangan dalam segmentasi adalah menentukan jumlah klaster yang optimal. *K-Means* dikenal dengan efisiensi komputasinya pada dataset besar dan kesederhanaan implementasinya (Niu, 2025). Sementara itu, *Hierarchical Clustering* memiliki keunggulan dalam tidak perlunya menentukan jumlah klaster di awal, sehingga kemampuannya menampilkan hubungan antar klaster melalui

dendrogram yang mudah diinterpretasi (Wardani et al., 2023). Namun, kinerja setiap metode dapat bervariasi tergantung pada karakteristik data dan konteks permasalahannya.

Beberapa penelitian terdahulu telah menerapkan *clustering* di sektor kebugaran, Niu (2025) berhasil menggunakan *K-Means* untuk mengelompokkan pelanggan gym, sementara Yeomans et al. (2024) mengidentifikasi kluster pelanggan yang paling berisiko *churn* dengan menggabungkan variabel perilaku dan attitudinal, menghasilkan empat kluster yang berbeda seperti “*Dissatisfied Casuals*” dan “*Enthusiastic Absentees*”. Namun, penelitian yang secara khusus membandingkan keefektifan *K-Means* dan *Hierarchical Clustering* dalam sektor pusat kebugaran masih terbatas. Penelitian perbandingan serupa lebih banyak dilakukan di sektor lain, seperti perbankan (Rahmadiana & Lestarini, 2025) dan e-commerce (Haryoko et al., 2025; Kumar et al., 2025) yang menunjukkan hasil yang beragam. Wardani et al. (2023) dan Rahmati & Wijayanto (2021) secara khusus menemukan bahwa *K-Means* unggul dalam segmentasi pasar otomotif berdasarkan nilai *Silhouette Coefficient*.

Berdasarkan hal tersebut, penelitian ini dirancang untuk menganalisis profil kluster pelanggan yang terbentuk dari karakteristik dan perilaku dengan membandingkan kualitas segmentasi yang dihasilkan kedua metode berdasarkan metrik evaluasi *clustering*. Dengan tujuan tersebut, diharapkan penelitian ini dapat memberikan bukti empiris mengenai metode yang paling optimal untuk segmentasi pelanggan pusat kebugaran.

2 METODE PENELITIAN

2.1 Data Penelitian

Data yang digunakan bersumber dari dataset publik yang berjudul *Gym Customers Features and Churn* yang diperoleh dari platform Kaggle (Vinueza, 2023). Dataset ini berisi informasi anonim mengenai pelanggan pusat kebugaran seperti usia, jenis kelamin, lama keanggotaan, dan sebagainya. Seluruh data telah disimulasikan untuk tujuan pembelajaran dan analisis data, sehingga tidak mengandung identitas pribadi responden. Data yang digunakan terdiri atas 4000 observasi dan 14 variabel, seperti yang dijelaskan pada Tabel 1 berikut.

Tabel 1. Informasi Data

Kode	Variabel	Deskripsi	Tipe Data
X1	<i>Gender</i>	Jenis kelamin pelanggan (0 = Wanita, 1 = pria)	Kategorikal (Biner)
X2	<i>Near_Location</i>	Indikator lokasi pelanggan yang tinggal atau bekerja di dekat lokasi pusat kebugaran	Biner (1/0)
X3	<i>Partner</i>	Indikator pelanggan merupakan karyawan dari perusahaan mitra pusat kebugaran	Biner (1/0)
X4	<i>Promo_friends</i>	Indikator pelanggan yang mendaftar melalui program promosi “ajak teman”	Biner (1/0)
X5	<i>Phone</i>	Indikator pelanggan yang mencantumkan nomor telepon pada data pendaftaran	Biner (1/0)

X6	<i>Age</i>	Usia pelanggan dalam satuan tahun	Numerik (Integer)
X7	<i>Lifetime</i>	Durasi waktu sejak kunjungan pertama pelanggan ke pusat kebugaran (Bulan)	Numerik (Integer)
X8	<i>Contract_period</i>	Periode durasi kontrak keanggotaan yang dipilih pelanggan (1, 3, 6, atau 12 bulan)	Numerik (Ordinal)
X9	<i>Month_to_end_Contract</i>	Sisa waktu kontrak keanggotaan pelanggan berakhir. (Bulan)	Numerik (Integer)
X10	<i>Group_visits</i>	Indikator pelanggan mengikuti sesi latihan kelompok	Biner (1/0)
X11	<i>Avg_class_frequency_total</i>	Rata-rata frekuensi kunjungan pelanggan per minggu selama keseluruhan masa keanggotaan	Numerik (kontinu)
X12	<i>Avg_class_frequency_current_month</i>	Rata-rata frekuensi kunjungan per minggu selama bulan berjalan	Numerik (kontinu)
X13	<i>Avg_additional_charges_total</i>	Jumlah total pengeluaran tambahan pelanggan untuk layanan ekstra (kafe, pijat, atau produk olahraga)	Numerik (kontinu)
X14	<i>Churn</i>	Indikator status pelanggan berhenti berlangganan (0 = aktif, 1 = berhenti)	Biner (1/0)

2.2 Pra-pemrosesan Data

Proses analisis dilakukan secara sistematis menggunakan Python 3.12 pada lingkungan pengembangan Jupyter Notebook. Penelitian ini memanfaatkan pustaka Scikit-learn, dan SciPy untuk penerapan metode *clustering*, sedangkan visualisasi data memanfaatkan pustaka Matplotlib, Seaborn, dan Plotly. Terdapat beberapa Langkah sistematis dalam pra-pemrosesan data. Pertama, pemeriksaan kelengkapan data dilakukan melalui identifikasi nilai yang hilang dan observasi duplikat. Hasil pemeriksaan menunjukkan tidak terdapat data yang hilang maupun duplikasi, sehingga seluruh data dapat digunakan tanpa imputasi tambahan. Selanjutnya, penanganan pencilan (*outliers*) pada variabel numerik menggunakan metode *Interquartile Range* (IQR) dengan faktor 1.5 (Qur'ani & Wijayanto, 2023; Rahmati & Wijayanto, 2021). Metode IQR bekerja dengan menghitung kuartil pertama (Q1) dan kuartil ketiga (Q3), kemudian menentukan batas bawah dan atas menggunakan persamaan berikut.

$$x < Q_1 - 1.5 \cdot IQR \quad \text{atau} \quad x > Q_3 + 1.5 \cdot IQR$$

dimana,

$$IQR = Q_3 - Q_1$$

Observasi yang teridentifikasi sebagai pencilan kemudian ditangani menggunakan metode *Winsorization*, yaitu mengganti nilai ekstrem dengan nilai pada persentil ke-5 dan ke-95 untuk mempertahankan distribusi data tanpa menghilangkan informasi penting (Qur'ani & Wijayanto, 2023; Ufeli et al., 2025). Setelah itu, analisis multikolineritas dilakukan dengan menghitung korelasi antar variabel numerik menggunakan koefisien korelasi *Pearson* berikut.

$$r_{xy} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}}$$

Multikolineritas terjadi ketika dua atau lebih variabel memiliki korelasi yang sangat tinggi (misalnya $> 0,8$). Dalam metode klusterisasi, hal ini dapat memengaruhi hasil segmentasi karena variabel yang berkorelasi tinggi dapat memberikan bobot ganda terhadap dimensi yang sama dan dapat menggeser posisi *centroid* pada *K-Means* atau jarak antar objek pada *Hierarchical Clustering* (Abdulhafedh, 2021; Saputra et al., 2025). Seluruh variabel numerik kemudian distandardisasi menggunakan *StandardScaler* dari pustaka *Scikit-learn* agar memiliki rata-rata 0 dan simpangan baku 1 supaya variabel dengan skala besar tidak mendominasi perhitungan jarak (Wahyuningtyas et al., 2023). Standardisasi dilakukan menggunakan persamaan berikut.

$$z_i = \frac{x_i - \mu_x}{\sigma_x}$$

Dimana μ_x adalah rata-rata sampel dan σ_x adalah simpangan baku sampel. Standardisasi diperlukan agar variabel dengan skala besar tidak mendominasi perhitungan jarak (Abdulhafedh, 2021). Hal ini krusial mengingat kedua metode *clustering* yang digunakan termasuk dalam kategori *unsupervised learning*, dimana semua variabel diperlakukan setara sebagai fitur pembentuk kluster demi menghasilkan segmentasi yang representatif terhadap realitas bisnis. Sementara itu, seluruh variabel kategorik telah berbentuk biner (0/1) sehingga tidak memerlukan proses pengkodean tambahan.

2.3 Perbandingan Metode

Setelah data siap, dua metode *clustering* yaitu *K-Means* dan *Hierarchical Clustering* diterapkan dan dibandingkan secara komprehensif. Langkah awal dalam proses ini adalah penentuan jumlah kluster optimal melalui empat metrik yaitu *Elbow Method*, *Silhouette Coefficient* (SC), *Calinski-Harabasz Index* (CHI), dan *Davies-Bouldin Index* (DBI) (Abdulhafedh, 2021). *Elbow Method* memanfaatkan nilai inersia *K-Means*, yaitu *within-cluster sum of squares* (Rahmati & Wijayanto, 2021).

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

Jumlah kluster optimal diidentifikasi pada titik “siku” (*elbow*), yaitu pada saat penurunan inersia tidak lagi signifikan (Rahmadhan & Wasesa, 2022).

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

Silhouette Coefficient (SC) digunakan untuk menilai pemisahan antar kluster secara objektif. Metrik ini dihitung dengan mempertimbangkan rata-rata jarak intra-kluster (a) dan jarak rata-rata ke kluster terdekat lainnya (b). Nilai SC berada pada rentang $[-1,1]$, dimana nilai yang semakin tinggi menunjukkan kualitas kluster yang lebih baik (Ramadhan et al., 2023). Oleh karena itu, SC dianggap sebagai metrik evaluasi yang *robust* (Abdulhafedh, 2021). Sementara itu, DBI mengukur rasio rata-rata dispersi intra-kluster terhadap separasi inter-kluster (Abdulpatah et al., 2025). Indeks ini dihitung sebagai berikut.

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

Nilai DBI yang lebih rendah menunjukkan kualitas kluster yang lebih baik, karena menandakan dispersi intra-kluster yang kecil dan separasi inter-kluster yang besar (Rahmadhan & Wasesa, 2022). Terakhir, CHI mengukur rasio antara dispersi antar-kluster dan dispersi dalam kluster (Ufeli et al., 2025).

$$CHI = \frac{\text{trace}(B)/(k-1)}{\text{trace}(W)/(n-k)}$$

Nilai CHI yang lebih tinggi menunjukkan kualitas kluster yang padat dan terpisah dengan jelas (Abdulhafedh, 2021). Metode *clustering* pertama yang digunakan yaitu *K-Means* yang bertujuan meminimalkan fungsi objektif inersia (Wahyuningtyas et al., 2023).

$$J = \sum_{i=1}^k \sum_{x \in C_i} |x - \mu_i|^2$$

K-Means bekerja dengan beberapa tahap, tahap pertama menginisialisasi *centroid* secara acak dengan parameter *k-means++* untuk meminimalkan bias awal dan meningkatkan kualitas *clustering* (Abdulhafedh, 2021). Tahap kedua adalah menetapkan titik data ke kluster dengan *centroid* terdekat berdasarkan jarak Euclidean seperti yang dijelaskan pada persamaan berikut.

$$d(x, \mu_i) = \sqrt{\sum_{j=1}^m (x_j - \mu_{ij})^2}$$

Tahap ketiga, menghitung ulang *centroid* sebagai rata-rata semua titik dalam kluster dengan persamaan berikut.

$$\mu_i = \frac{1}{|C_i|} \sum_{x \in C_i} x$$

Tahap kedua dan ketiga diulangi hingga mencapai konvergensi. *K-Means* dikenal efisien untuk dataset besar, penelitian terdahulu mencatat waktu komputasi hanya 0.402 detik dengan 7050 data (Wahyuningtyas et al., 2023). Metode kedua, *Hierarchical Clustering* digunakan dengan pendekatan *agglomerative* dan membandingkan tiga jenis *linkage* yaitu *Ward*, *Complete*, dan *Average* (Qur'ani & Wijayanto, 2023; Ramadhan et al., 2023). Metode *Ward* meminimalkan varians dalam kluster dengan fungsi objektif berikut.

$$d_{ward}(C_i, C_j) = \frac{|C_i||C_j|}{|C_i| + |C_j|} \|\mu_i - \mu_j\|^2$$

Sementara itu, *Complete linkage* menggunakan jarak maksimum antar titik yang dapat dihitung dengan persamaan berikut.

$$d_{complete}(C_i, C_j) = \max_{x \in C_i, y \in C_j} d(x, y)$$

Average linkage menggunakan jarak rata-rata antar titik. *Average linkage* seringkali menghasilkan struktur kluster yang lebih baik (Abdulpatah et al., 2025; Ufeli et al., 2025).

$$d_{average}(C_i, C_j) = \frac{1}{|C_i||C_j|} \sum_{x \in C_i} \sum_{y \in C_j} d(x, y)$$

Metode *linkage* optimal ditentukan dengan membandingkan nilai SC, DBI, dan CHI untuk mencari hasil yang paling seimbang. Setelah metode *linkage* terpilih, proses dilanjutkan dengan pembentukan dendrogram yang merepresentasikan struktur hierarkis lengkap (Qur'ani & Wijayanto, 2023). Pemotongan dendrogram dilakukan secara horizontal pada ketinggian *threshold* yang sesuai dengan jumlah kluster optimal. *Threshold* pemotongan dihitung berdasarkan rumus berikut dimana $Z[-(k), 2]$ merupakan jarak pada gabungan ke-(n-k) untuk jumlah kluster k dan δ merupakan *small offset* untuk memastikan pemisahan yang jelas antar kluster.

$$\text{threshold} = Z[-(k), 2] + \delta$$

Setelah kluster terbentuk, kualitas kedua metode dievaluasi kembali menggunakan empat metrik utama yaitu *Silhouette Coefficient* (SC), *Davies-Bouldin Index* (DBI), *Calinski-Harabasz Index* (CHI), dan *Dunn Index* (DI) (Abdulahfedh, 2021). *Dunn index* mengukur rasio antara jarak terdekat antar kluster dan diameter kluster terbesar, dengan persamaan sebagai berikut.

$$DI = \frac{\min_{1 \leq i < j \leq k} \delta(C_i, C_j)}{\max_{1 \leq l \leq k} \Delta_l}$$

Nilai DI yang tinggi menunjukkan kluster yang kompak (diameter kecil) dan terpisah dengan baik (Abdulahfedh, 2021). Selain kualitas kluster, stabilitas kluster dan efisiensi komputasi juga diuji. Stabilitas kluster diuji menggunakan *Adjusted Rand Index* (ARI) melalui pendekatan pengulangan (Qur'ani & Wijayanto, 2023). ARI mengukur kesamaan antara dua partisi dengan mengoreksi kesepakatan yang terjadi secara acak. Indeks ini merupakan perbaikan dari *Rand Index* dengan memperhitungkan ekspektasi similaritas. ARI dihitung berdasarkan table kontingensi:

$$ARI = \frac{\sum_{ij} \binom{n_{ij}}{2} - [\sum_i \binom{a_i}{2} \sum_j \binom{b_j}{2}] / \binom{n}{2}}{\frac{1}{2} [\sum_i \binom{a_i}{2} + \sum_j \binom{b_j}{2}] - [\sum_i \binom{a_i}{2} \sum_j \binom{b_j}{2}] / \binom{n}{2}}$$

Nilai ARI berkisar antar -1 hingga 1, dimana nilai 1 menunjukkan kedua partisi identik sempurna. Dalam penelitian ini, stabilitas diuji dengan melakukan *resampling* sebanyak 10 dengan menggunakan sekitar 80% data yang dipilih secara acak. ARI dihitung antara hasil kluster pada data lengkap dan hasil kluster pada data *resample*. Prosedur ini dilakukan untuk

membandingkan konsistensi hasil kluster terhadap variasi dalam data. Lalu efisiensi komputasi diukur melalui waktu eksekusi metode. Pengukuran dilakukan menggunakan modul *time* di Python dengan menghitung selisih antara waktu mulai dan waktu selesai proses *clustering*.

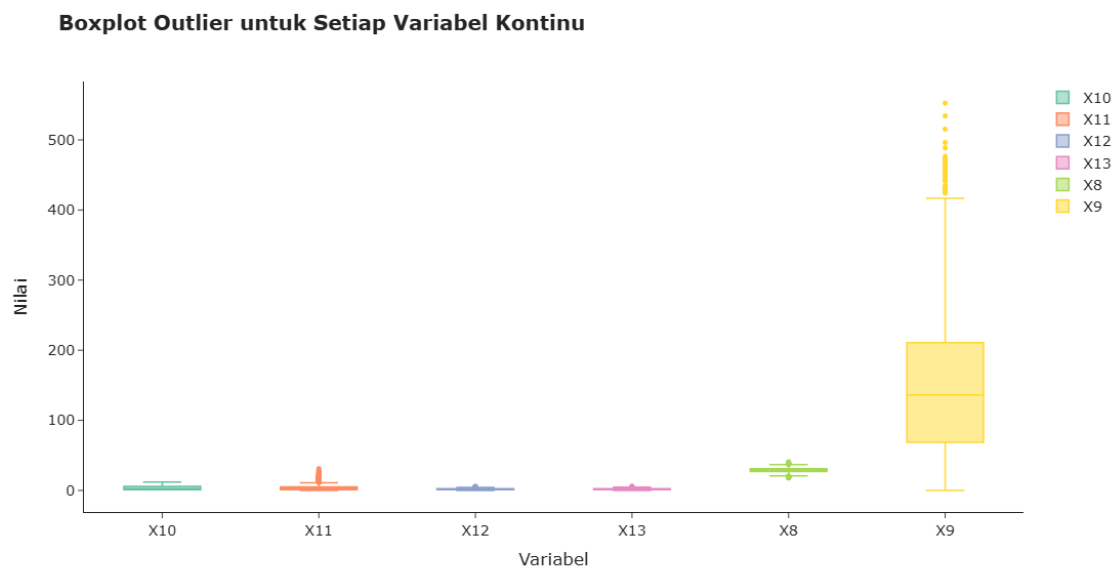
$$t_{\text{eksekusi}} = t_{\text{selesai}} - t_{\text{mulai}}$$

Waktu eksekusi mencakup seluruh proses *clustering*, termasuk inisialisasi, iterasi, dan konvergensi untuk *K-Means*, serta pembangunan dendrogram dan pemotongan untuk *Hierarchical Clustering*. Pengukuran dilakukan pada lingkungan komputasi yang sama secara perangkat keras maupun perangkat lunak untuk memastikan perbandingan yang adil antara kedua metode. Melalui evaluasi komprehensif yang mencakup kualitas internal, stabilitas (ARI), dan efisiensi komputasi, penelitian ini diharapkan memberikan perbandingan objektif antara *K-Means* dan *Hierarchical Clustering* untuk segmentasi pelanggan pusat kebugaran.

3 HASIL DAN PEMBAHASAN

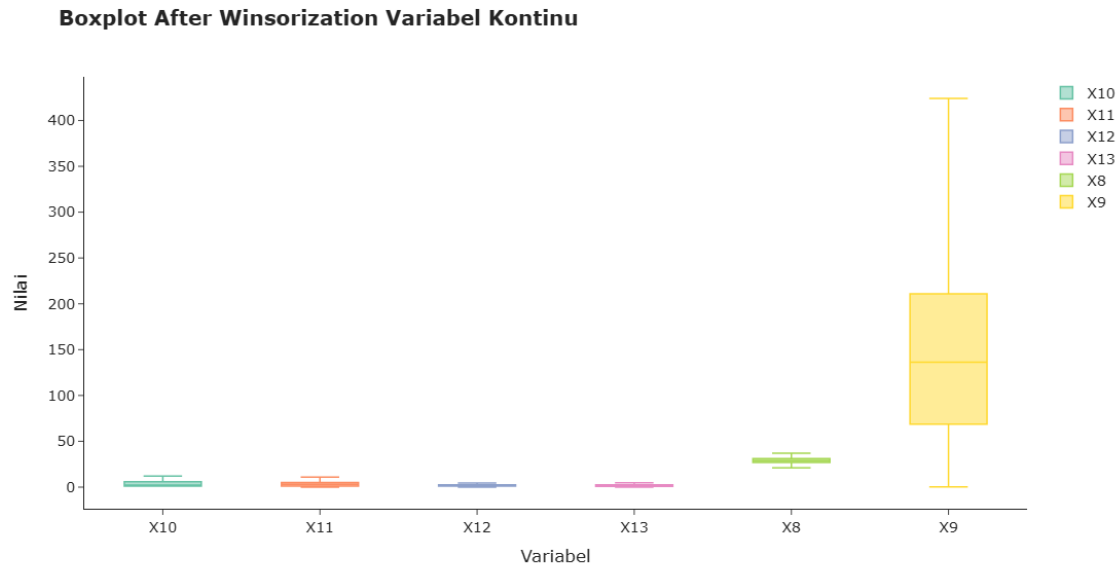
3.1 Pra-pemrosesan Data

Pemeriksaan awal terhadap dataset mengonfirmasi validitas data, dimana tidak ditemukan adanya nilai yang hilang maupun duplikasi observasi. Selain itu, variabel kategorik telah terformat dalam skema biner (0/1), sehingga tidak dibutuhkan imputasi lanjutan. Kendati demikian, identifikasi distribusi variabel numerik mengungkapkan keberadaan pencilan (*outlier*) yang signifikan, sebagaimana divisualisasikan pada Gambar 1.



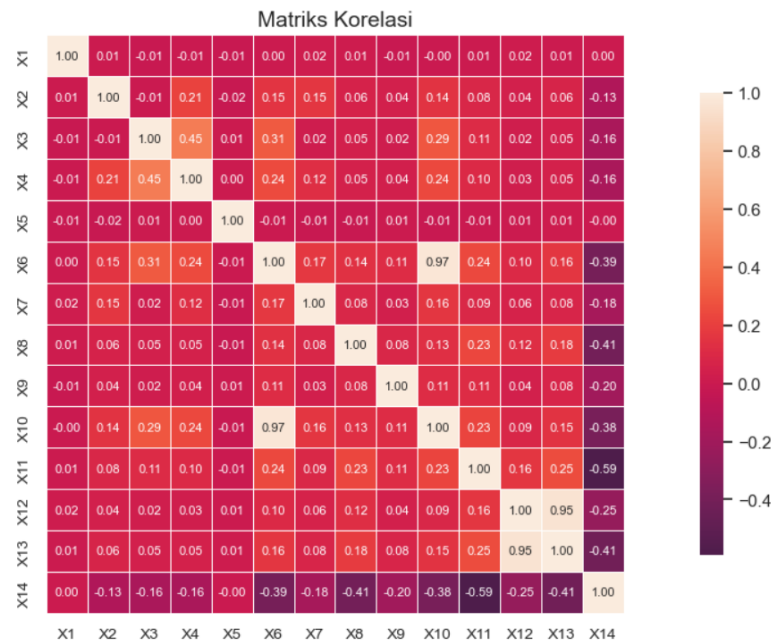
Gambar 1. Grafik *Outlier* Variabel Kontinu

Mitigasi dampak pencilan dilakukan melalui metode *Winsorization* dengan membatasi nilai ekstrem pada persentil ke-5 dan ke-95. Pendekatan ini diterapkan untuk mempreservasi integritas distribusi data tanpa mereduksi informasi esensial, sekaligus meminimalkan distorsi terhadap pembentukan kluster (Qur'ani & Wijayanto, 2023; Rahmati & Wijayanto, 2021). Hasilnya tidak terdapat pencilan sebagaimana divisualisasikan pada Gambar 2.



Gambar 2. Grafik *After Winsorization* Variabel Kontinu

Pasca-penanganan pencilan, evaluasi multikolinearitas menggunakan matriks korelasi *Pearson* (Gambar 3) mengindikasikan adanya redundansi data yang kuat (>0.9) pada beberapa pasangan variabel.

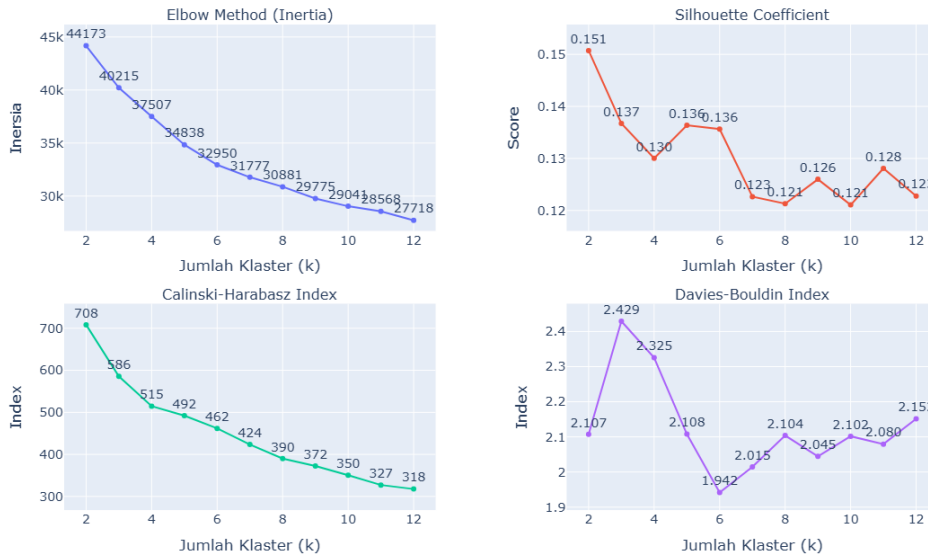


Gambar 3. Grafik korelasi antar variabel

Merujuk pada Gambar 3, koefisien korelasi ekstrem terdeteksi antara periode kontrak (X6) dan sisa waktu kontrak (X10) sebesar 0.97, serta antara rata-rata kunjungan mingguan akumulatif (X12) dan bulan berjalan (X13) sebesar 0.95. Berdasarkan pertimbangan substantif, variabel X6 dieliminasi akibat redundansi informasi dengan X10 yang dinilai lebih relevan dalam merepresentasikan sisa masa kontrak pelanggan saat ini. Sebaliknya, kedua variabel frekuensi kunjungan (X12 dan X13) tetap dipertahankan karena merefleksikan dimensi perilaku yang berbeda, yakni pola historis jangka panjang versus intensitas aktivitas terkini.

3.2 Penentuan Jumlah Kluster Optimal

Penentuan jumlah kluster optimal dilakukan dengan menggunakan empat metrik evaluasi yaitu *Elbow Method*, SC, DBI, dan CHI. Keempat metrik diaplikasikan dengan menghitung nilai masing-masing metrik untuk jumlah kluster (k) dari 2 hingga 12 dengan memanfaatkan pustaka Scikit-learn. Hasil perhitungan keempat metrik disajikan Gambar 4.



Gambar 4. Grafik Penentuan Jumlah Kluster dengan Empat Metrik

Berdasarkan evaluasi keempat metrik, $k=2$ dipilih sebagai konfigurasi optimal. Hal ini didukung oleh penurunan inersia pada Gambar 3, serta nilai tertinggi pada SC (0.151) dan CHI (708.384). Walaupun DBI merujuk $k=6$, konsistensi mayoritas indikator memvalidasi $k=2$ sebagai representasi struktur data yang paling *robust*.

3.3 K-Means

Metode *K-Means* dengan inisialisasi *k-means++* diterapkan untuk menjamin stabilitas konvergensi. Distribusi yang dihasilkan mengonfirmasi terbentuknya dua kluster pelanggan yang distingtif, sebagaimana tertera pada Tabel 2.

Tabel 2. Jumlah Kluster dengan *K-Means*

Kluster	Jumlah	Persentase
Kluster 0	1249	31.22%
Kluster 1	2751	68.77%

Distribusi kluster menunjukkan disparitas signifikan antara Kluster 1 (68.78%) dan Kluster 0 (31.22%). Evaluasi metrik pada Tabel 3 menunjukkan model $k=2$ menghasilkan SC terbaik, yang mengindikasikan pemisahan data yang paling distingtif.

Tabel 3. Evaluasi Kualitas Kluster

Metrik	Nilai
<i>Silhouette Score</i>	0.1507
<i>Calinski-Harabasz Index</i>	708.3844
<i>Davies-Bouldin Index</i>	2.1073

3.4 Hierarchical Clustering

Penentuan metode *linkage* optimal pada metode *Hierarchical Clustering* melalui komparasi metode Ward, Complete, dan Average. Evaluasi terhadap ketiga metode tersebut didasarkan pada empat metrik validasi internal untuk memastikan objektivitas dan stabilitas hasil sebagaimana dirangkum pada Tabel 4 berikut.

Tabel 4. Perbandingan Metode *Linkage* pada *Hierarchical Clustering*

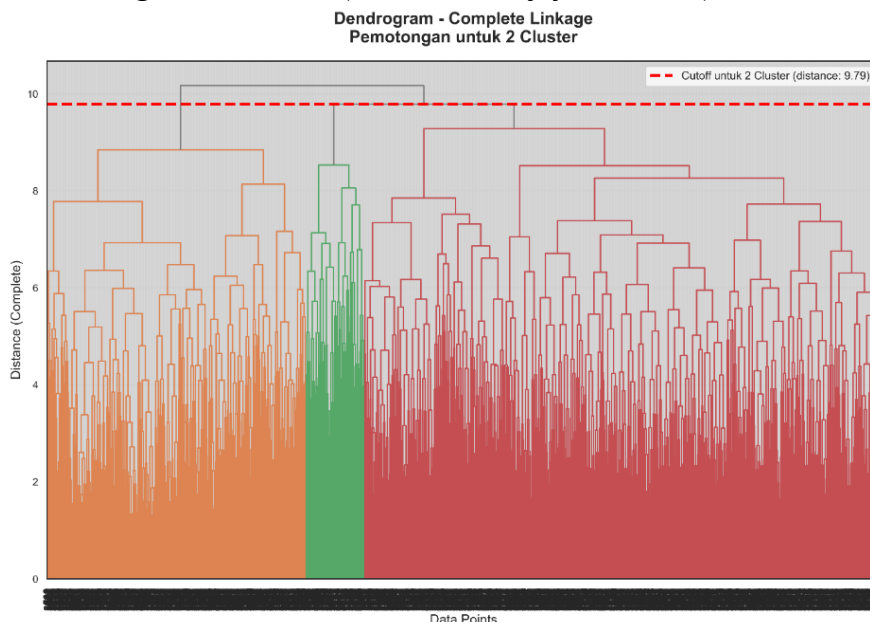
Metode	SC	CHI	DBI	Distribusi Klaster
<i>Ward</i>	0.1248	521.4415	2.5082	2702 vs 1298
<i>Complete</i>	0.1314	564.3539	2.3877	2755 vs 1245
<i>Average</i>	0.2102	4.3860	0.9229	3998 vs 2

Merujuk pada hasil komparasi, metode *Complete* menunjukkan performa paling seimbang dengan nilai SC 0.1314 dan CHI 564.3539. Kendati metode *Average* mencatatkan metrik kompetitif, varian ini dikesampingkan akibat ketimpangan distribusi klaster yang ekstrem. Oleh karena itu, metode *Complete* ditetapkan sebagai pendekatan optimal. Implementasi selanjutnya menggunakan konfigurasi $k=2$ menghasilkan struktur klaster yang disajikan pada Tabel 5 berikut.

Tabel 5. Distribusi Anggota Cluster *Hierarchical Clustering*

Klaster	Jumlah	Persentase
Klaster 0	2755	68.88%
Klaster 1	1245	31.12%

Distribusi *Hierarchical Clustering* menunjukkan konsistensi pola dengan metode *K-Means*, dimana Klaster 0 mendominasi (68.88%) dan Klaster 1 (31.12%). Struktur aglomerasi data selanjutnya divisualisasikan melalui dendrogram guna memetakan urutan penggabungan klaster berdasarkan tingkat similaritas (Qur'ani & Wijayanto, 2023).



Gambar 5. Dendrogram – *Hierarchical Clustering* (Complete Linkage)

Pada Gambar 5, penetapan garis potong (*cut-off*) pada jarak Euclidean 9.79 menghasilkan dua klaster utama dengan separasi yang tegas. Tingginya jarak pemotongan ini merefleksikan

derajat disimilaritas yang besar, sehingga adanya perbedaan karakteristik yang signifikan antara kedua klaster pelanggan (Qur'ani & Wijayanto, 2023).

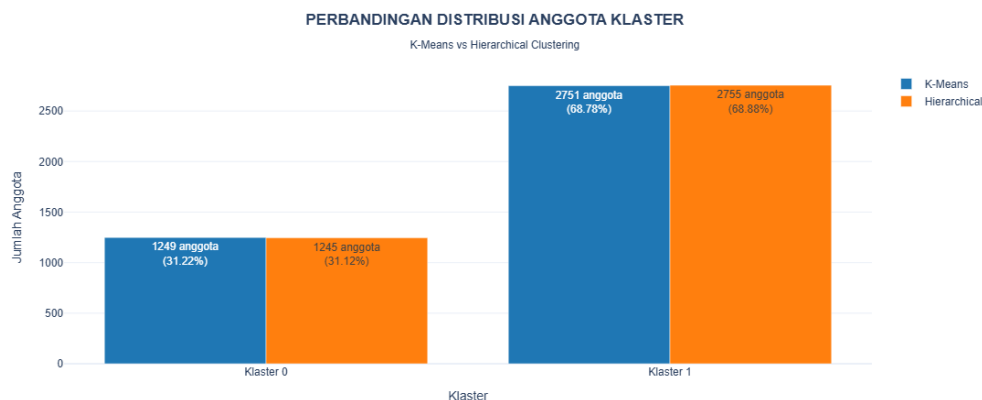
3.5 Perbandingan Kinerja Metode

Perbandingan kinerja kedua metode klusterisasi dilakukan dengan keempat metrik internal untuk kualitas klaster: uji stabilitas ARI dengan sepuluh kali pengulangan pada 80% sampel data, serta waktu komputasi yang dibutuhkan kedua metode. Dengan memanfaatkan pustaka Sklearn dan Scipy, hasil perbandingan metode *K-means* dengan *Hierarchical Clustering* disajikan pada Tabel 6 berikut.

Tabel 6. Perbandingan Kinerja *K-Means* dan *Hierarchical Clustering*

Metrik	<i>K-Means</i>	<i>Hierarchical</i>
<i>Silhouette Score</i>	0.1507	0.1314
<i>Calinski-Harabasz Index</i>	708.3844	564.3539
<i>Davies-Bouldin Index</i>	2.1073	2.3877
<i>Dunn Index</i>	0.0528	0.0459
Stabilitas (ARI)	0.9051	0.0901
Waktu Komputasi (s)	0.0808	0.6442

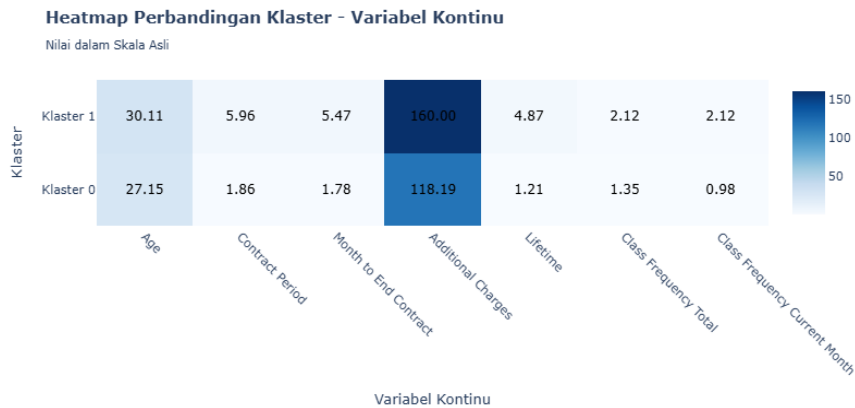
Berdasarkan Tabel 6, *K-Means* unggul dalam seluruh metrik evaluasi mulai dari kualitas klaster, stabilitas yang tinggi dengan ARI 0.9051 dan kecepatan waktu eksekusi dengan waktu 0.0808 detik. Visualisasi pada Gambar 6 menunjukkan kemiripan yang sangat tinggi dalam distribusi anggota klaster antara kedua metode yang menghasilkan pola distribusi yang hampir identik. Metode *K-Means* direkomendasikan secara kuat untuk implementasi dalam sistem manajemen pelanggan pusat kebugaran dalam mendukung pengembangan strategi pemasaran yang efektif.



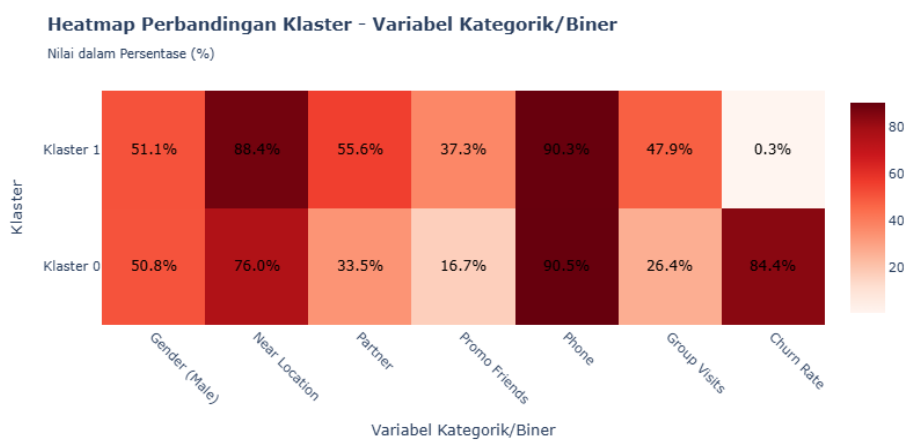
Gambar 6. Perbandingan Distribusi Klaster

3.6 Interpretasi Profil Klaster

Berdasarkan evaluasi komprehensif, metode *K-Means* terbukti paling unggul dalam mengidentifikasi dua klaster pelanggan dengan karakteristik distingtif. Hasil ini menjadi landasan empiris bagi perumusan strategi manajemen pelanggan. Profilisasi mendalam terhadap karakteristik tiap klaster divisualisasikan melalui *heatmap* variabel kategorik dan kontinu, sebagaimana disajikan pada Gambar 7 dan Gambar 8.



Gambar 7. Perbandingan Kluster - Variabel Kontinu



Gambar 8. Perbandingan Kluster – Variabel Kategorik

Merujuk visualisasi tersebut, kluster musiman (Kluster 0) diklasifikasikan sebagai “*High-Risk Casual*” yaitu kluster pelanggan dengan tingkat *engagement* yang rendah dan risiko berhenti berlangganan yang sangat tinggi. Kluster ini beranggotakan 1249 pelanggan dan didominasi oleh demografi usia muda rata-rata 27 tahun dengan frekuensi kunjungan rendah (1 kali/bulan) dan partisipasi grup yang terbatas (26.4%). Anggota kluster ini memiliki indikator loyalitas tergolong lemah, ditandai dengan durasi kontrak singkat (1-2 bulan) serta tingkat *churn* yang ekstrem mencapai 84%. Dengan kontribusi pendapatan tambahan sebesar \$118.19, kluster ini memerlukan intervensi retensi agresif untuk memitigasi risiko berhenti berlangganan.

Sebaliknya, kluster inti (Kluster 1) diidentifikasi sebagai “*Low-Risk Regular*”, yaitu kluster pelanggan setia yang memiliki rutinitas latihan konsisten dan risiko berhenti berlangganan yang sangat rendah. Kluster ini terdiri dari 2751 pelanggan yang merepresentasikan kelompok usia lebih matang (rata-rata 30 tahun). Anggota kluster ini menunjukkan *engagement* tinggi melalui kunjungan rutin (2 kali/bulan), partisipasi grup aktif (47.9%), serta durasi kontrak panjang (6 bulan). Dengan tingkat *churn* minimal (0.3%) dan pengeluaran tambahan signifikan sebesar \$160.00, kluster ini merupakan asset prioritas yang perlu dioptimalkan melalui program loyalitas dan peningkatan nilai ekonomi berkelanjutan.

4 KESIMPULAN

Berdasarkan hasil analisis dan evaluasi model, metode *K-Means* lebih unggul dibandingkan *Hierarchical Clustering* pada seluruh dimensi evaluasi. Tingginya stabilitas

algoritma yang divalidasi oleh nilai ARI sebesar 0.9051, menjadikannya instrument yang reliabel untuk penerapan praktis. Penelitian ini berhasil memetakan dua klaster utama pelanggan pusat kebugaran yaitu pelanggan klaster inti (68.8%) yang loyal dan klaster musiman (31.2%) yang berisiko tinggi berhenti berlangganan. Anggota Klaster musiman memiliki kesamaan dalam durasi kontrak singkat dan interaksi rendah, sedangkan anggota Klaster inti memiliki kesamaan dalam loyalitas tinggi dan partisipasi aktif. Sebagai implikasi manajerial, implementasi sistem segmentasi *real-time* menggunakan metode *K-Means* direkomendasikan bagi manajemen pusat kebugaran.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Allah SWT atas kelancara yang diberikan dalam penyelesaian penelitian ini. Penghargaan yang tulus disampaikan kepada Ibu Nuramaliyah, M.Si., selaku dosen pembimbing atas arahan, diskusi, dan masukan berharga selama proses penyusunan artikel ini. Penulis juga mengucapkan terima kasih kepada Fakultas Sains dan Teknologi, Universitas Terbuka yang telah memfasilitasi pengembangan akademik penulis, serta kepada penyedia dataset publik di platform Kaggle yang telah menyediakan akses data terbuka sehingga mendukung terlaksananya penelitian komparatif ini.

DAFTAR PUSTAKA

- Abdulhafedh, A. (2021). Incorporating K-means, Hierarchical Clustering and PCA in Customer Segmentation. *Journal of City and Development*, 3(1), 12–30. <https://doi.org/10.12691/jcd-3-1-3>
- Abdulpatah, T., Sari, B. N., & Susilawati. (2025). Perbandingan Algoritma K-Means dan Agglomerative Hierarchical Clustering untuk Pengelompokan Daerah Penghasil Padi di Indonesia. *Jurnal Informatika Dan Teknik Elektro Terapan*, 13(3). <https://doi.org/10.23960/jitet.v13i3.7251>
- Haryoko, A., Putra, A., Arifia, A., & Nurlifa, A. (2025). Penerapan Algoritma K-Means dengan Pendekatan RFM untuk Analisis Loyalitas Pelanggan pada Bisnis Online. *CURTINA (Computer Science or Informatic Journal)*, 6(1), 9–17. <https://doi.org/10.55719/curtina.v6i1.1745>
- Kumar, S., Rani, R., Pippal, S. K., & Agrawal, R. (2025). Customer segmentation in e-commerce: K-means vs hierarchical clustering. *Telkomnika (Telecommunication Computing Electronics and Control)*, 23(1), 119–128. <https://doi.org/10.12928/TELKOMNIKA.v23i1.26384>
- Niu, Z. (2025). Customer Segmentation and Management Strategy Optimization for Gyms Using K-Means Clustering. *SCITEPRESS – Science and Technology Publications*, 1, 239–242. <https://doi.org/10.5220/0013214000004568>
- Qur'ani, N. S., & Wijayanto, A. W. (2023). Implementation of K-Means and Hierarchical Clustering in Determining Levels of Smart City 2022 Based on Motion Index. *SISTEMASI*, 12(2), 346–360. <https://doi.org/10.32520/stmsi.v12i2.2457>
- Rahmadhan, R., & Wasesa, M. (2022). Segmentation using Customers Lifetime Value: Hybrid K-means Clustering and Analytic Hierarchy Process. *Journal of Information Systems*

- Engineering and Business Intelligence*, 8(2), 130–141.
<https://doi.org/10.20473/jisebi.8.2.130-141>
- Rahmadiana, R. N., & Lestarini, D. (2025). Implementation of the K-Means Algorithm for Customer Churn Segmentation in Developing Bank Marketing Strategies. *SISTEMASI*, 14(4), 1909–1919. <https://doi.org/10.32520/stmsi.v14i4.5341>
- Rahmati, R., & Wijayanto, A. W. (2021). Analisis Cluster dengan Algoritma K-Means, Fuzzy C-Means dan Hierarchical Clustering (Studi Kasus: Indeks Pembangunan Manusia Tahun 2019). *JIKO (Jurnal Informatika Dan Komputer)*, Vol 5, No. 2, 73–80. <https://ejournal.akakom.ac.id/index.php/jiko/article/view/422>
- Ramadhan, H., Rizal Abdan Kamaludin, M., Alfian Nasrullah, M., & Rolliawati, D. (2023). Comparison of Hierarchical, K-Means and DBSCAN Clustering Methods for Credit Card Customer Segmentation Analysis Based on Expenditure Level. In *Journal of Applied Informatics and Computing (JAIC)* (Vol. 7, Issue 2). <https://doi.org/10.30871/jaic.v7i2.4939>
- Saputra, M. A., Dinata, R. K., & Afrillia, Y. (2025). Clustering of Aquaculture Productivity Villages in East Aceh Using the K-Means Algorithm. *Journal of Applied Informatics and Computing*, 9(5), 2382–2390. <https://doi.org/10.30871/jaic.v9i5.10102>
- Ufeli, C. P., Sattar, M. U., Hasan, R., & Mahmood, S. (2025). Enhancing Customer Segmentation Through Factor Analysis of Mixed Data (FAMD)-Based Approach Using K-Means and Hierarchical Clustering Algorithms. *Information*, 16(6), 441. <https://doi.org/10.3390/info16060441>
- Vinueza, A. (2023). *Gym customers features and churn*. <https://www.kaggle.com/datasets/adrianvinueza/gym-customers-features-and-churn>
- Wahyuningtyas, F. D., Arafat, A., Stiawan, A., & Rolliawati, D. (2023). Komparasi Algoritma Hierarchical, K-Means, dan DBSCAN pada Analisis Data Penjualan Melalui Facebook. *Explore: Jurnal Sistem Informasi Dan Telematika*, 14(1), 7–16. <https://doi.org/10.36448/jsit.v14i1.2931>
- Wardani, S. D. K., Ariyanto, A. S., Umroh, M., & Rolliawati, D. (2023). Perbandingan Hasil Metode Clustering K-Means, DB Scanner & Hierarchical untuk Analisa Segmentasi Pasar. *JIKO (Jurnal Informatika Dan Komputer)*, 7(2), 191. <https://doi.org/10.26798/jiko.v7i2.796>
- Yeomans, C., Karg, A., & Nguyen, J. (2024). Who churns from fitness centres? Evidence from behavioural and attitudinal segmentation. *Managing Sport and Leisure*, 1–17. <https://doi.org/10.1080/23750472.2024.2305896>