

Evaluating Post-Training Employment Outcomes for Workforce Policy through *Clustering*: A Comparative Study of K-Means, Hierarchical Clustering, Gaussian Mixture Model, and Fuzzy C-Means

Mayrisa Andriyani ¹, Siti Nurwilda ², Wahyu Ningtiyas Mergianti ³, Nurissaidah Ulinnuha ^{4*}

^{1,2,3,4}Mathematics, Science and Technology, UIN Sunan Ampel Surabaya, Indonesia

09010223008@uinsa.ac.id ², 09010223010@uinsa.ac.id ², 09040223058@uinsa.ac.id ³, nuris.ulinnuha@uinsa.ac.id ⁴

Article Info

Article history:

Received 2026-06-30

Revised 2026-07-30

Accepted 2026-08-13

Keyword:

*Clustering,
Davies-Bouldin Index,
Fuzzy C-means,
Post-Training Evaluation,
Workforce Policy,
Silhouette Score.*

ABSTRACT

Evaluating the effectiveness of government training programs requires more than a binary employed/unemployed indicator, since participants who find work still differ substantially in salary level, waiting time to employment, job position, gender, and employment sector. Prior studies applying clustering to workforce or socio-economic data have generally compared only two or three algorithms at a time, on datasets from other domains such as banking or health, leaving it unclear which method best segments multidimensional post-training outcome data with mixed numerical and categorical attributes and no natural hard boundaries between groups. This study addresses that gap by comparing four clustering algorithms, K-Means, Hierarchical Clustering, Gaussian Mixture Model, and Fuzzy C-Means, to segment post-training participant data from the Surabaya City Government, with the aim of providing local policymakers with an evidence-based grouping of participants that can inform which sectors and job levels most need follow-up support. The dataset consists of 509 observations with attributes for job position, monthly salary, gender, employment sector, and waiting time. Prior to clustering, the data were preprocessed through cleaning of inconsistent categorical entries, median imputation of missing numerical values, One-Hot Encoding of categorical attributes, and Min-Max normalization of numerical attributes to a common 0-1 scale. Clustering performance was evaluated using Silhouette Score and Davies-Bouldin Index (DBI) across cluster counts $k = 2-5$. Fuzzy C-Means with five clusters achieved the best overall performance, with a Silhouette Score of 0.649 and a DBI of 0.640; its closest competitor, Ward-linkage Hierarchical Clustering, achieved a comparable Silhouette Score of 0.647 but a markedly higher DBI of 0.872, indicating that FCM produced more compact, well-separated clusters overall even though the two methods separated participants almost equally well. The resulting five clusters show clear policy-relevant differences: the largest cluster, dominated by the industrial sector, has the highest average salary but also the longest waiting time to employment, suggesting that industrial-sector training would benefit from faster competency certification and job-matching support, while smaller, female-dominated clusters in the social sector show shorter waiting times but lower salaries, pointing to a need for upskilling pathways into higher-paying roles. These findings illustrate how clustering can move post-training evaluation beyond simple placement rates toward data-driven, sector-specific recommendations for workforce training policy.



This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

I. PENDAHULUAN

Ketenagakerjaan merupakan salah satu aspek penting dalam pembangunan ekonomi karena berkaitan dengan peningkatan kesejahteraan masyarakat dan pengurangan tingkat pengangguran [1]. Dalam menghadapi persaingan dunia kerja yang semakin ketat, peningkatan kualitas sumber daya manusia menjadi kebutuhan yang tidak dapat diabaikan. Salah satu upaya yang dilakukan pemerintah untuk meningkatkan kompetensi tenaga kerja adalah melalui berbagai program pelatihan keterampilan yang dirancang sesuai dengan kebutuhan pasar kerja. Pelatihan yang diselenggarakan Pemerintah Kota Surabaya merupakan salah satu upaya strategis untuk meningkatkan kompetensi peserta serta mempercepat penyerapan tenaga kerja. Namun, keberhasilan program pasca pelatihan tidak selalu dapat dinilai hanya dari status bekerja atau tidak bekerja, karena outcome pasca pelatihan dapat bervariasi dalam hal tingkat gaji, masa tunggu kerja, jabatan, dan bidang pekerjaan. Di sisi lain, perkembangan teknologi informasi telah menghasilkan data dalam jumlah besar yang dapat dimanfaatkan untuk mendukung pengambilan keputusan berbasis data. Analisis terhadap data peserta pelatihan dapat memberikan informasi mengenai karakteristik kelompok peserta dan hasil yang dicapai setelah mengikuti pelatihan. Informasi tersebut penting bagi pemerintah sebagai dasar evaluasi efektivitas program sekaligus untuk menyusun kebijakan dan strategi pengembangan pelatihan yang lebih tepat sasaran. Oleh karena itu, diperlukan pendekatan analitik yang mampu mengidentifikasi data dan mengelompokkan peserta berdasarkan kemiripan karakteristiknya, salah satunya melalui teknik clustering.

Clustering merupakan metode pengelompokan data yang bertujuan memaksimalkan kemiripan dalam satu kelompok dan meminimalkan kemiripan antar kelompok [2]. Terdapat berbagai algoritma *clustering* yang memiliki cara berbeda dalam membentuk kelompok data. K-Means merupakan salah satu metode yang paling sering digunakan, metode ini mengelompokkan data berdasarkan titik pusat (centroid) yang mewakili masing-masing *cluster* [3]. Selain itu, terdapat *Hierarchical Clustering* yang membentuk kelompok data secara bertingkat dan menghasilkan diagram pohon (dendrogram) [4]. Metode lainnya adalah Gaussian Mixture Model (GMM) yang menggunakan pendekatan probabilitas berbasis distribusi Gaussian [5]. Sementara itu, Fuzzy C-Means (FCM) merupakan metode yang memberikan derajat keanggotaan pada setiap data terhadap seluruh *cluster* [6].

Untuk memperoleh hasil clustering yang optimal diperlukan metrik evaluasi untuk menilai kualitas pengelompokan. Salah satu ukuran yang sering digunakan adalah *Silhouette Score*. *Silhouette Score* digunakan untuk mengukur seberapa baik suatu objek berada dalam *cluster* dibandingkan dengan *cluster* lain, dengan nilai yang semakin tinggi menunjukkan pemisahan *cluster* yang lebih baik. Sementara itu, *Davies-Bouldin Index (DBI)* mengukur rasio antara kekompakan *cluster* dan jarak antar *cluster*, di mana

nilai yang lebih kecil menunjukkan kualitas *cluster* yang lebih optimal [7].

Beberapa penelitian terbaru menunjukkan bahwa penggunaan lebih dari satu metrik evaluasi sangat penting dalam membandingkan algoritma clustering. Hal ini dilakukan agar hasil evaluasi menjadi lebih objektif dan mampu menggambarkan kualitas *cluster* secara lebih menyeluruh. Salah satu penelitian membandingkan algoritma K-Means, *Hierarchical Clustering*, dan Gaussian Mixture Model (GMM) menggunakan *Silhouette Score* dan *Davies-Bouldin Index (DBI)* sebagai ukuran kualitas *cluster* [8]. Hasil penelitian tersebut menunjukkan bahwa tidak ada satu algoritma yang selalu memberikan hasil terbaik pada semua jenis data. Kinerja setiap algoritma sangat dipengaruhi oleh karakteristik data yang digunakan, seperti pola penyebaran data, bentuk *cluster*, dan variabel yang dianalisis. Pada beberapa dataset, K-Means menghasilkan nilai *Silhouette Score* yang lebih tinggi sehingga mampu membentuk kelompok yang lebih jelas. Namun, pada dataset lain, GMM memberikan nilai DBI yang lebih rendah yang menunjukkan kualitas *cluster* yang lebih baik. Hasil ini mengindikasikan bahwa pemilihan algoritma *clustering* perlu disesuaikan dengan karakteristik data yang akan dianalisis.

Penelitian lain yang membandingkan metode Fuzzy C-Means dan K-Means pada data tenaga kerja menunjukkan bahwa Fuzzy C-Means (FCM) mampu memberikan hasil pengelompokan yang lebih baik ketika batas antar kelompok tidak terlalu jelas [9]. Hal ini karena FCM memungkinkan suatu data memiliki tingkat keanggotaan pada lebih dari satu *cluster*. Dengan karakteristik tersebut, FCM dapat menggambarkan struktur data yang lebih kompleks dibandingkan metode K-Means yang mengharuskan setiap data hanya berada pada satu *cluster*. Selain pemilihan algoritma, penggunaan metrik evaluasi yang tepat juga menjadi faktor penting dalam analisis clustering. Sebuah penelitian menunjukkan bahwa penggunaan satu metrik evaluasi saja dapat menyebabkan hasil interpretasi menjadi kurang akurat [10]. *Silhouette Score* lebih berfokus pada seberapa baik suatu data terpisah dari *cluster* lain, sedangkan *Davies-Bouldin Index (DBI)* lebih menekankan pada tingkat kekompakan *cluster* yang terbentuk. Oleh karena itu, kombinasi kedua metrik tersebut sering direkomendasikan untuk mengevaluasi kualitas *clustering* secara lebih komprehensif.

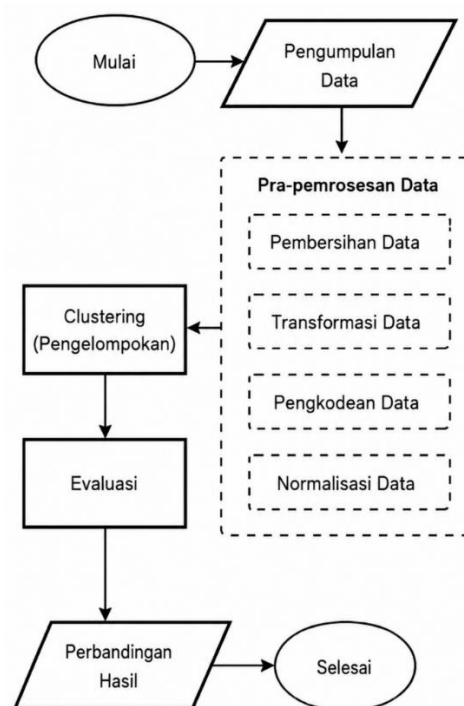
Penelitian lainnya menunjukkan bahwa setiap algoritma memiliki keunggulan dan keterbatasan masing-masing [11]. GMM cenderung memberikan hasil yang lebih baik pada data dengan pola yang kompleks, K-Means lebih efektif untuk data yang memiliki kelompok yang relatif seragam, sedangkan *Hierarchical Clustering* unggul dalam menggambarkan hubungan antar data melalui struktur bertingkat. Berdasarkan berbagai temuan tersebut, dapat disimpulkan bahwa belum ada algoritma *clustering* yang selalu unggul untuk semua jenis data. Oleh karena itu, diperlukan penelitian yang membandingkan beberapa algoritma *clustering* pada data peserta pasca pelatihan Pemerintah Kota Surabaya untuk

menentukan metode yang paling sesuai dalam menghasilkan pengelompokan yang optimal.

Dari penelitian sebelumnya yang telah membandingkan beberapa algoritma *clustering* seperti K-Means, Hierarchical Clustering, Gaussian Mixture Model (GMM), dan Fuzzy C-Means (FCM), penelitian tersebut dilakukan pada karakteristik data yang berbeda, seperti data perbankan dan kesehatan. Selain itu, penelitian terdahulu umumnya hanya membandingkan sebagian algoritma tertentu sehingga belum memberikan gambaran menyeluruh mengenai performa keempat algoritma tersebut dalam satu studi pada data ketenagakerjaan. Penelitian sebelumnya juga menunjukkan bahwa penggunaan lebih dari satu metrik evaluasi diperlukan untuk memperoleh hasil evaluasi yang lebih objektif, namun penerapannya pada data pasca pelatihan kerja pemerintah masih sangat terbatas. Oleh karena itu, masih terdapat kesenjangan penelitian terkait perbandingan K-Means, Hierarchical Clustering, Gaussian Mixture Model (GMM), dan Fuzzy C-Means (FCM) menggunakan Silhouette Score dan Davies–Bouldin Index pada data pasca pelatihan kerja. Penelitian ini berkontribusi dengan melakukan analisis komparatif terhadap keempat algoritma *clustering* tersebut pada data peserta pasca pelatihan Pemerintah Kota Surabaya menggunakan kombinasi Silhouette Score dan Davies–Bouldin Index. Hasil penelitian diharapkan dapat memberikan rekomendasi metode *clustering* yang paling sesuai untuk karakteristik data ketenagakerjaan serta mendukung evaluasi program pelatihan kerja berbasis data. Selain itu, penelitian ini bertujuan untuk membandingkan kinerja empat algoritma *clustering* K-Means, Hierarchical Clustering, Gaussian Mixture Model (GMM), dan Fuzzy C-Means (FCM) dalam mengelompokkan data pasca pelatihan Pemerintah Kota Surabaya berdasarkan variabel gaji, masa tunggu kerja (tahun), jabatan, jenis kelamin, dan bidang pekerjaan menggunakan evaluasi Silhouette Score dan Davies–Bouldin Index. Hasil penelitian diharapkan dapat memberikan masukan dalam pengambilan keputusan berbasis data serta pengembangan evaluasi program pelatihan yang lebih terstruktur.

II. METODE

Penelitian ini menggunakan pendekatan kuantitatif dengan metode *clustering* untuk mengelompokkan data pasca pelatihan Pemerintah Kota Surabaya. Tujuan penelitian adalah membandingkan kinerja beberapa algoritma *clustering* dalam menghasilkan pengelompokan data yang optimal berdasarkan karakteristik peserta. Tahapan penelitian yang dilakukan meliputi Gambar 1.



Gambar 1. Tahapan Penelitian

Pemilihan empat algoritma yang tertera pada Gambar 1 yaitu K-Means, Fuzzy C-Means (FCM), Gaussian Mixture Model (GMM), dan Hierarchical Clustering didasarkan pada karakteristik pengujian komparatif dalam menangani batas kelompok yang saling tumpang tindih (*overlap*). Keempat algoritma ini dipilih dibandingkan metode clustering yang lebih baru seperti DBSCAN atau Spectral Clustering karena tiga alasan metodologis. Pertama, keempatnya merepresentasikan tiga paradigma clustering yang berbeda secara mendasar, yaitu partitional/hard clustering berbasis centroid (K-Means), soft/fuzzy clustering (FCM), model probabilistik (GMM), dan clustering berbasis hierarki (Hierarchical Clustering), sehingga perbandingan ini dapat menunjukkan paradigma mana yang paling sesuai untuk data ketenagakerjaan hasil One-Hot Encoding. Kedua, keempat metode ini tidak memerlukan penentuan parameter kepadatan (*density*) seperti pada DBSCAN, yang sulit ditentukan secara objektif pada data campuran numerik-kategorikal berdimensi tinggi hasil encoding. Ketiga, keempat algoritma ini menghasilkan keanggotaan cluster yang mudah diinterpretasikan langsung oleh pemangku kebijakan (baik sebagai partisi tegas maupun sebagai derajat keanggotaan fuzzy/probabilistik), berbeda dengan Spectral Clustering yang memerlukan reduksi dimensi tambahan dan cenderung kurang transparan untuk kebutuhan pelaporan kebijakan publik. Proses penelitian ini dilakukan mulai dari pembersihan data (*data cleaning*), yakni merapikan struktur data dan memperbaiki nilai yang tidak konsisten pada variabel Jabatan, Gaji Per Bulan, Jenis Kelamin, Bidang Pekerjaan, dan Masa Tunggu. Selanjutnya, dilakukan tahap

transformasi data dengan mengubah atribut kategorikal menjadi bentuk numerik serta melakukan normalisasi skala pada atribut numerik agar seluruh variabel berada dalam rentang nilai yang seimbang. Data yang telah siap kemudian dimasukkan sebagai variabel input ke dalam pengujian keempat algoritma clustering, hingga akhirnya dilakukan evaluasi kinerja model secara komparatif untuk menentukan algoritma yang paling optimal dalam mengelompokkan data pasca pelatihan tersebut.

1. Data Collection

Dataset yang digunakan dalam penelitian ini merupakan data pasca pelatihan yang diperoleh dari Badan Riset dan Inovasi Daerah Kota Surabaya. Data ini merepresentasikan kondisi pekerjaan dan karakteristik peserta yang telah memperoleh pekerjaan setelah mengikuti program pelatihan yang diselenggarakan oleh pemerintah daerah. Dataset terdiri dari 509 data observasi dengan 6 atribut dimana setiap baris merepresentasikan satu individu peserta pelatihan, adapun atribut yang digunakan dalam penelitian ini disajikan pada Tabel I.

Variabel bidang menunjukkan sektor usaha tempat responden bekerja. Untuk memudahkan proses analisis, atribut tersebut dikelompokkan menjadi empat kategori, yaitu Industri, Bisnis, Sosial, dan Hospitality. Definisi masing-masing kategori bidang pekerjaan disajikan pada Tabel II.

TABEL II
KLASIFIKASI BIDANG PEKERJAAN

Bidang Pekerjaan	Definisi
Industri	Bidang yang berkaitan dengan kegiatan produksi, manufaktur, teknis, dan teknologi
Bisnis	Bidang yang berkaitan dengan kegiatan perdagangan, distribusi, administrasi, keuangan, dan logistik.
Sosial	Bidang yang berfokus pada pelayanan masyarakat di sektor pendidikan, kesehatan, dan pemerintahan
Hospitality	Bidang yang bergerak dalam layanan perhotelan, pariwisata.

2. Analisis Deskriptif

Analisis deskriptif terhadap data peserta pasca pelatihan menjadi langkah awal yang dilakukan sebelum proses *preprocessing* dan pemodelan menggunakan empat algoritma clustering. Tujuannya untuk mendapatkan gambaran umum mengenai distribusi, frekuensi, serta nilai pemusatan data (rata-rata, minimum, dan maksimum) pada karakteristik demografi dan ekonomi peserta. Karakteristik kategorikal (Jabatan, Jenis Kelamin, Bidang Pekerjaan) dianalisis menggunakan persentase frekuensi, sedangkan karakteristik numerik (Gaji Per Bulan dan Masa Tunggu) dihitung menggunakan statistik ringkasan seperti rata-rata dan rentang nilai.

TABEL I
DESKRIPSI ATRIBUT DATASET

No	Nama Atribut	Tipe Data	Deskripsi	Rentang
1	No	Numerik	Nomor identitas Data	1 - 509
2	Jabatan	Kategorial	Posisi atau pekerjaan peserta	Operasional, Profesional, Manajerial
3	Gaji Per Bulan	Numerik	Pendapatan peserta per bulan	150.000 - 40.000.000
4	Jenis Kelamin	Kategorial	Jenis Kelamin peserta	Laki – laki, Perempuan
5	Bidang Pekerjaan	Kategorial	Bidang pekerjaan peserta	Industri, Sosial, Hospitality, Bisnis
6	Masa Tunggu	Numerik	Lama waktu (dalam tahun) memperoleh pekerjaan	0 – 2 tahun

3. Preprocessing data

Tahap *preprocessing* dilakukan untuk mempersiapkan dataset sebelum proses clustering. Tahapan ini meliputi data cleaning, imputasi *missing value* menggunakan median pada Persamaan (1), transformasi data kategorikal dengan *One-Hot Encoding*, serta normalisasi data menggunakan metode *Min-Max Scaling* pada Persamaan [12](2).

$$x_{ij} = \text{Median}(X_j) \quad (1)$$

dengan:

x_{ij} = nilai hasil imputasi pada atribut ke- j

X_j = himpunan nilai pada atribut ke- j

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}} \quad (2)$$

dengan:

x' = nilai hasil normalisasi

x = nilai asli data

$x_{\min}$ = nilai minimum atribut

$x_{\max}$ = nilai maksimum atribut

4. Clustering

Clustering merupakan metode untuk mengelompokkan data berdasarkan kemiripan karakteristik yang dimiliki [13]. Pada penelitian ini digunakan empat algoritma clustering, yaitu K-Means, Fuzzy C-Means (FCM), Gaussian Mixture Model (GMM), dan Hierarchical Clustering untuk mengelompokkan data pasca pelatihan.

a. K-Means

K-Means merupakan salah satu algoritma *clustering* yang banyak digunakan karena sederhana dan mudah diimplementasikan. Algoritma ini mengelompokkan data berdasarkan tingkat kemiripan karakteristik dengan

menempatkan objek yang memiliki jarak terdekat ke pusat *cluster* (centroid) yang sama [14].

Pada penelitian ini, pengukuran jarak dilakukan menggunakan metrik Euclidean. Penggunaan metrik tersebut dimungkinkan karena atribut kategorikal telah ditransformasikan menjadi representasi numerik melalui teknik *One-hot* encoding pada tahap preprocessing. Selain itu, atribut numerik seperti Gaji Per Bulan dan Masa Tunggu telah dinormalisasi ke rentang 0 hingga 1, sehingga seluruh atribut berada pada skala yang seragam dan dapat diproses secara efektif menggunakan jarak Euclidean.

b. Fuzzy C-Means

Fuzzy C-Means (FCM) merupakan algoritma *clustering* berbasis fuzzy yang mengelompokkan data menurut derajat keanggotaan, dimana setiap titik data memiliki lebih dari satu keanggotaan. Proses *clustering* dilakukan secara iteratif dengan memperbarui derajat keanggotaan dan pusat *cluster* hingga mencapai kondisi konvergen. Fungsi objektif FCM dinyatakan pada Persamaan [15](3).

$$J_m = \sum_{i=1}^c \sum_{j=1}^n (u_{ij})^m d^2(X_{ij}, V_k) \quad (3)$$

dengan:

u_{ij} = Derajat keanggotaan data ke-j dalam *cluster* ke-i
 $d(X_{ij}, V_k)$ = Jarak antara data ke-j dan pusat *cluster* ke-i

Derajat keanggotaan u_{ij} untuk data ke-j pada *cluster* ke-i dihitung dengan menggunakan Persamaan (4):

$$u_{ij} = \left[\sum_{k=1}^c \left(\frac{d_{ij}}{d_{kj}} \right)^{\frac{2}{m-1}} \right]^{-1} \quad (4)$$

dengan:

d_{ij} = jarak antara data ke-j dan pusat *cluster* ke-i
 d_{kj} = jarak antara data ke-j dan pusat *cluster* ke-k
 m = parameter fuzzifikasi

c. Gaussian Mixture Model (GMM)

Gaussian Mixture Model (GMM) merupakan algoritma *clustering* probabilistik yang mengelompokkan data berdasarkan campuran beberapa distribusi Gaussian [16]. GMM dapat membuat setiap data untuk memiliki keanggotaan probabilistik dalam lebih dari satu *cluster* dengan derajat keanggotaan yang bervariasi. Model GMM dinyatakan pada Persamaan [17](5).

$$p(x_i | \pi, \mu, \Sigma) = \sum_{k=1}^K \pi_k N(x_i | \mu_k, \Sigma_k) \quad (5)$$

dengan:

$p(x_i | \pi, \mu, \Sigma)$ = probabilitas data ke-i
 π_k = koefisien campuran Gaussian ke-k
 $N(x_i | \mu_k, \Sigma_k)$ = fungsi densitas probabilitas dari distribusi Gaussian ke-k, dengan μ_k sebagai rata-rata dan Σ_k sebagai kovarians
 K = jumlah komponen Gaussian

Parameter model pada Gaussian Mixture Model (GMM) diestimasi menggunakan algoritma Expectation-Maximization (EM). Algoritma ini bekerja secara iteratif melalui dua tahapan utama, yaitu Expectation Step (E-Step) dan Maximization Step (M-Step) untuk memperbarui parameter rata-rata (μ_k), matriks kovarians (Σ_k), dan koefisien campuran (π_k) hingga mencapai kondisi konvergen. Detail dari kedua tahapan tersebut adalah sebagai berikut:

1. Expectation Step (E-Step)

nilai fungsi keanggotaan posterior atau probabilitas tanggung jawab (γ_{nk}) dihitung menggunakan parameter yang ada saat ini berdasarkan nilai densitas distribusi Gaussian melalui Persamaan (6):

$$\gamma_{nk} = \frac{\pi_k N(x_n | \mu_k, \Sigma_k)}{\sum_{k=1}^K \pi_k N(x_n | \mu_k, \Sigma_k)} \quad (6)$$

dengan:

γ_{nk} = probabilitas data ke-n berasal dari *cluster* ke-k
 π_k = koefisien campuran Gaussian ke-k
 μ_k = rata-rata Gaussian ke-k
 Σ_k = matriks kovarians Gaussian ke-k

2. Maximization Step (M-Step)

Setelah nilai γ_{nk} diperoleh pada tahap E-step, seluruh parameter model diperbarui untuk memaksimalkan fungsi log-likelihood dengan menggunakan persamaan-persamaan berikut:

$$\mu_k^{new} = \frac{\sum_{n=1}^N \gamma_{nk} x_n}{\sum_{n=1}^N \gamma_{nk}} \quad (7)$$

$$\Sigma_k^{new} = \frac{\sum_{n=1}^N \gamma_{nk} (x_n - \mu_k^{new})(x_n - \mu_k^{new})}{\sum_{n=1}^N \gamma_{nk}} \quad (8)$$

$$\pi_k^{new} = \frac{\sum_{n=1}^N \gamma_{nk}}{N} \quad (9)$$

Dimana N menunjukkan jumlah total data observasi dalam penelitian ini yaitu sebanyak 509 data, dan K adalah jumlah komponen distribusi Gaussian. Proses iterasi antara E-step dan M-step ini terus diulang hingga perubahan nilai log-likelihood berada di bawah ambang batas yang ditentukan (konvergen).

d. Hierarchical Clustering

Hierarchical Clustering merupakan metode *clustering* yang membentuk struktur *cluster* secara bertingkat (hierarki) tanpa menentukan jumlah *cluster* di awal [19]. Proses pengelompokan dapat dilakukan secara Agglomerative (bottom-up) maupun Divisive (top-down) dan umumnya divisualisasikan dalam bentuk dendrogram [20]. Pada penelitian ini digunakan metode Average Linkage dan Ward Linkage untuk menentukan kedekatan antar *cluster* [21].

Average Linkage menghitung jarak antar *cluster* berdasarkan rata-rata jarak seluruh pasangan objek dari dua *cluster*, sebagaimana ditunjukkan pada Persamaan [22](10).

$$d_{(uv)w} = \frac{N_u}{N_{uv}} d_{uw} + \frac{N_v}{N_{uv}} d_{vw} \quad (10)$$

dengan:

d_{UW} = jarak antar kluster U dengan kluster W sebelum digabungkan

d_{VW} = jarak antar kluster V dengan kluster W sebelum digabungkan

N_{UV} = jumlah objek pada cluster UV

N_W = jumlah objek pada cluster W

Sementara itu, *Ward Linkage* menggabungkan *cluster* dengan meminimalkan variasi dalam *cluster* (*within-cluster variance*) sehingga menghasilkan *cluster* yang lebih homogen [23]. Perhitungan dinyatakan pada Persamaan (11).

$$SSE = \sum_{j=1}^p \left(\sum_{i=1}^n X_{ij}^2 - \frac{1}{n} \left(\sum_{i=1}^n X_{ij} \right)^2 \right) \quad (11)$$

dengan:

SSE = total variasi dalam cluster

X_{ij} = nilai objek ke- i pada variabel ke- j

p = jumlah variabel

n = jumlah objek

Pada metode Ward, kriteria pengelompokan ditentukan berdasarkan fungsi peningkatan jumlah kuadrat kekeliruan atau *Error Sum of Squares Increase* (ΔSSE) ketika dua kluster (misalkan kluster U dan kluster V) bergabung menjadi satu kluster baru (UV). Nilai peningkatan *error* tersebut dihitung menggunakan Persamaan (12):

$$\Delta SSE_{UV} = SSE_{UV} - (SSE_U + SSE_V) \quad (12)$$

dengan:

ΔSSE_{UV} = nilai peningkatan error akibat penggabungan U dan V

SSE_{UV} = nilai total error setelah kluster U dan V bergabung menjadi satu

SSE_U = nilai total error pada kluster U sebelum penggabungan

SSE_V = nilai total error setelah kluster V sebelum penggabungan

Algoritma Ward Linkage akan mensimulasikan seluruh kemungkinan penggabungan kluster dan memilih pasangan kluster yang menghasilkan nilai peningkatan *error* (ΔSSE) paling minimum agar kluster tetap homogen [24].

5. Parameter pengujian

Parameter merupakan nilai yang digunakan untuk mengontrol proses pembelajaran atau pemodelan data [25]. Parameter berfungsi sebagai pengatur utama dalam menentukan bagaimana suatu algoritma bekerja, termasuk dalam mengatur struktur model, proses iterasi, hingga mekanisme optimasi yang digunakan.

Pengujian algoritma *clustering* dalam penelitian ini dibatasi dengan jumlah kelompok antara 2 hingga 5 *cluster* ($k = 2, 3, 4, 5$). Penentuan rentang rentang jumlah *cluster* ini didasarkan pada dua pertimbangan utama. Secara metodologis, batasan ini digunakan untuk mempermudah analisis komparatif performa algoritma (seperti *Silhouette Score* dan *Davies-Bouldin Index*) dalam menemukan struktur pengelompokan yang paling kontras dan stabil tanpa risiko terjadinya *over-clustering* yang membuat interpretasi data

menjadi terlalu bias. Secara praktis, rentang 2 hingga 5 *cluster* dinilai paling ideal dan relevan bagi Pemerintah Kota Surabaya dalam merumuskan strategi tindak lanjut pasca pelatihan agar intervensi kebijakan ketenagakerjaan berikutnya dapat dilakukan secara lebih fokus dan tepat sasaran. Perlu dicatat bahwa nilai performa terbaik pada penelitian ini diperoleh pada $k = 5$, yaitu batas atas rentang yang diuji, sehingga terdapat kemungkinan performa *clustering* dapat terus meningkat pada nilai k yang lebih besar. Penelitian selanjutnya disarankan memperluas rentang pengujian jumlah *cluster* melebihi lima agar dapat dipastikan bahwa $k = 5$ benar-benar merupakan solusi optimal dan bukan sekadar nilai terbaik dalam rentang yang dibatasi. Pendekatan *hierarchical fuzzy clustering* yang dapat menentukan jumlah *cluster* optimal secara otomatis tanpa bergantung pada indeks validitas tertentu, sebagaimana diusulkan pada [28], dapat menjadi alternatif metodologis untuk mengatasi keterbatasan ini.

6. Evaluasi

Kualitas hasil *clustering* dievaluasi menggunakan *Silhouette Score* dan *Davies-Bouldin Index* (DBI). *Silhouette Score* digunakan untuk mengukur tingkat pemisahan dan kekompakan *cluster*, sedangkan *Davies-Bouldin Index* digunakan untuk mengukur kemiripan antar *cluster*.

Kedua metrik ini bersifat internal, artinya kualitas *cluster* dinilai hanya berdasarkan struktur data itu sendiri tanpa label acuan eksternal, sehingga hasil evaluasi tetap perlu dimaknai secara hati-hati. *Silhouette Score* dan DBI juga mengasumsikan jarak Euclidean sebagai ukuran kemiripan yang relevan, sebuah asumsi yang sesuai untuk K-Means dan Ward Linkage, namun kurang selaras dengan basis probabilistik GMM maupun basis keanggotaan derajat parsial pada FCM. K-Means dan Ward Linkage mengasumsikan *cluster* berbentuk relatif bulat (*spherical*) dengan varians yang homogen antar *cluster*, sedangkan GMM mengasumsikan setiap *cluster* mengikuti distribusi Gaussian multivariat, sebuah asumsi yang kurang sesuai untuk atribut kategorikal hasil One-Hot Encoding yang bersifat biner. FCM tidak mensyaratkan bentuk distribusi tertentu dan memungkinkan derajat keanggotaan parsial pada data yang berada di batas antar kelompok, sehingga secara teoritis lebih toleran terhadap tumpang tindih yang muncul akibat kombinasi atribut numerik dan kategorikal pada dataset ini. Sebagai pelengkap, penelitian ini menyertakan *Silhouette Plot* pada Gambar 2 untuk memeriksa konsistensi nilai *silhouette* antar anggota *cluster*, sekaligus menjadi bentuk validasi tambahan di luar nilai rata-rata *Silhouette Score* dan DBI. Penelitian selanjutnya disarankan menambahkan metrik atau analisis stabilitas lain, seperti Calinski-Harabasz Index atau validasi bootstrap terhadap keanggotaan *cluster*, untuk memperkuat kesimpulan yang diperoleh.

a. Silhouette Score

Silhouette Score mengukur seberapa baik suatu data berada pada *cluster* yang sesuai dibandingkan dengan *cluster*

lainnya. Nilai Silhouette Score berada pada rentang -1 hingga 1 , di mana nilai yang lebih tinggi menunjukkan kualitas *clustering* yang lebih baik [26]. Berikut perhitungan evaluasi Silhouette index dalam persamaan (13).

$$S_i = \frac{b_i - a_i}{\max(a_i, b_i)} \quad (13)$$

dengan:

S_i = Nilai silhouette untuk data ke- i

a_i = Rata-rata jarak data ke cluster yang sama

b_i = Rata-rata jarak data ke *cluster* terdekat

b. Davies-Bouldin Index

Davies-Bouldin Index (DBI) adalah metrik yang mengukur rasio antar jarak tiap *cluster* dengan ukuran *cluster* itu sendiri [27]. Semakin rendah nilai DBI, semakin baik kualitas yang dihasilkan cluster. DBI dihitung dengan memperhitungkan dua komponen utama, yakni *separation* (pemisahan antar cluster) dan *compactness* (kepadatan dalam cluster). Sehingga didapatkan rumus persamaan (14).

$$DBI = \frac{1}{C} \sum_{i=1}^C \max_{j \neq i} \left(\frac{s_i + s_j}{d_{ij}} \right) \quad (14)$$

dengan:

C = jumlah cluster

s_i, s_j = rata-rata jarak setiap data ke centroid *cluster* i/j

d_{ij} = jarak antar pusat *cluster* i dan j

$\max_{j \neq i} \left(\frac{s_i + s_j}{d_{ij}} \right)$ = Nilai kemiripan maksimum antara *cluster* ke- i dengan *cluster* lain

7. Comparison Result

Seluruh hasil evaluasi dibandingkan untuk menentukan algoritma *clustering* terbaik berdasarkan nilai Silhouette Score dan Davies-Bouldin Index.

III. HASIL DAN PEMBAHASAN

Hasil yang disajikan pada bab ini meliputi gambaran dataset yang digunakan, hasil *preprocessing* data, konfigurasi proses clustering, hasil pengelompokan data menggunakan masing-masing algoritma, serta hasil evaluasi kinerja dari setiap metode *clustering* yang digunakan dalam penelitian ini.

A. Gambaran Dataset

Dataset yang digunakan dalam penelitian ini merupakan data peserta pelatihan yang diperoleh dari Badan Riset dan Inovasi Daerah Kota Surabaya. Atribut yang terdapat pada dataset meliputi Jabatan, Gaji Per bulan, Jenis Kelamin, Bidang, dan Masa Tunggu. Atribut-atribut tersebut digunakan sebagai dasar dalam proses pengelompokan data menggunakan beberapa algoritma *clustering* yang diterapkan dalam penelitian ini. Data peserta pelatihan yang digunakan dalam penelitian ini ditampilkan pada Tabel III.

TABEL III
DATASET PESERTA PELATIHAN

No	Jabatan	Gaji Per Bulan	Jenis Kelamin	Bidang Pekerjaan	Masa Tunggu
1	Operasional	4.525.479	Perempuan	Industri	1
2	Manajerial	2.700.000	Laki-laki	Bisnis	0
⋮	⋮	⋮	⋮	⋮	⋮
507	Profesional	2.000.000	Perempuan	Sosial	0
508	Operasional	1.900.000	Perempuan	Sosial	0
509	Operasional	3.200.000	Laki-laki	Industri	0

B. Analisis Deskriptif Data

Sebanyak 509 peserta pasca-pelatihan Pemerintah Kota Surabaya menjadi sampel yang dianalisis secara deskriptif dalam penelitian ini. Analisis tersebut dilakukan untuk memberikan gambaran umum mengenai karakteristik demografi dan kondisi ekonomi mereka sebelum melangkah ke proses pengelompokan dengan algoritma *clustering*.

Karakteristik kategorikal peserta dari segi jenis kelamin menunjukkan dominasi kuat oleh laki-laki pada bidang-bidang teknis, sedangkan peserta perempuan lebih banyak terkonsentrasi di sektor layanan sosial serta profesional seperti pendidikan dan kesehatan. Posisi pekerjaan tingkat operasional menjadi status jabatan yang ditempati oleh mayoritas mutlak responden dalam penelitian ini. Kondisi tersebut selaras dengan tujuan utama pelatihan kerja pemerintah yang memang menyasar penguatan keterampilan teknis siap kerja, mengingat posisi profesional dan manajerial memiliki frekuensi yang jauh lebih kecil serta sebagian besar diisi oleh tenaga kerja berlatar belakang sektor sosial atau pendidikan. Sektor industri dan bisnis akhirnya tercatat sebagai wadah penyerapan tenaga kerja terbesar bagi para alumni setelah menyelesaikan pelatihan, yang kemudian disusul oleh sektor sosial dan sektor *Hospitality* (perhotelan/pariwisata).

Data pertama yang termasuk numerik yakni Pendapatan bulanan peserta. Pendapatan bulanan peserta menunjukkan rentang yang sangat lebar, yaitu berkisar antara Rp150,000 hingga Rp40,000,000 per bulan. Ketimpangan yang tinggi ini mengindikasikan adanya variasi jenis pekerjaan yang sangat kontras di lapangan, di mana sektor industri secara umum menawarkan rata-rata gaji yang lebih kompetitif dibandingkan dengan sektor *Hospitality* dan sosial pada tingkat operasional. Karakter data numerik berikutnya terdapat rentang waktu penyerapan lulusan pelatihan ke dunia kerja tercatat berada pada kurun waktu antara 0 hingga 2 tahun. Mayoritas peserta dari seluruh sektor berhasil mendapatkan pekerjaan dalam waktu di bawah 1 tahun, meskipun sektor tertentu seperti industri menunjukkan waktu tunggu yang relatif lebih lama akibat ketatnya proses seleksi kompetensi.

C. Hasil Data Preprocessing

Pada tahap ini dilakukan *preprocessing* untuk mempersiapkan data sebelum digunakan dalam proses clustering. Proses diawali dengan pembersihan data (data

cleaning) melalui perapian struktur dataset, pengaturan ulang indeks, serta pembersihan atribut kategorikal seperti Jabatan, Jenis Kelamin, dan Bidang. Pembersihan dilakukan dengan menghapus spasi berlebih (*trimming*), mengubah seluruh teks menjadi huruf kecil (*lowercase*), serta memperbaiki kesalahan penulisan (*typo*) pada beberapa data.

Selanjutnya dilakukan transformasi data dengan mengonversi atribut numerik ke format yang sesuai. Atribut Gaji Per Bulan dibersihkan dari karakter non-numerik dan dikonversi ke tipe data numerik (*float*), sedangkan atribut Masa Tunggu dikonversi ke bentuk numerik. Setelah itu, atribut kategorikal yang tidak memiliki urutan yakni Jabatan

dan Bidang ditransformasikan menjadi bentuk numerik menggunakan metode *One-Hot Encoding*.

Tahap berikutnya adalah normalisasi pada atribut numerik menggunakan metode Min-Max Scaling untuk menyamakan rentang nilai antar variabel ke dalam skala 0 hingga 1. Atribut yang dinormalisasi meliputi Gaji Per Bulan dan Masa Tunggu. Hasil *preprocessing* data ditunjukkan pada Tabel IV.

Berdasarkan Tabel IV, atribut numerik telah dinormalisasi sehingga berada pada rentang nilai 0 hingga 1, sedangkan atribut kategorikal telah ditransformasikan menggunakan One-Hot Encoding. Hasil *preprocessing* menghasilkan 11 fitur yang selanjutnya digunakan sebagai input pada proses clustering.

TABEL IV
SAMPel HASIL *PREPROCESSING* DATA

No	Gaji Per bulan	Masa Tunggu	Jabatan Manajerial	Jabatan Operasional	Jabatan Profesional	Jenis Kelamin	Bidang Bisnis	Bidang Hospitality	Bidang Industri	Bidang Sosial
1	0,109799	0,25	0	1	0	0	0	0	1	0
2	0,063990	0,00	1	0	0	1	1	0	0	0
3	0,011292	0,25	0	1	0	0	0	0	0	1
4	0,015056	0,00	0	0	1	0	0	0	0	1

TABEL V
HASIL CLUSTERING

Cluster	FCM (Sil/DBI)	Hierarchical (Sil/DBI)		K-Means (Sil/DBI)	GMM (Sil/DBI)
		Ward	Average		
<i>Cluster 2</i>	0,366/1,153	0,348/1,123	0,366/1,000	0,395/1,279	0,395/1,279
<i>Cluster 3</i>	0,511/1,093	0,410/1,141	0,332/0,948	0,430/1,151	0,350/1,168
<i>Cluster 4</i>	0,613/0,787	0,590/0,954	0,412/0,982	0,467/1,065	0,445/1,125
<i>Cluster 5</i>	0,649/0,640	0,647/0,872	0,555/0,956	0,562/1,047	0,486/0,954

Keterangan: Sil = Silhouette Score, DBI = Davies-Bouldin Index

TABEL VI
KARAKTERISTIK CLUSTERING

Cluster	Jumlah Data	Gaji Rata-rata	Masa Tunggu	Bidang Dominan	Gender Dominan	Jabatan Dominan
C1	78	3.573.410	0,23	Sosial	50%P	Operasional
C2	173	3.793.639	0,43	Industri	79% L	Operasional
C3	149	3.186.376	0,33	Bisnis	52% L	Operasional
C4	72	2.830.583	0,22	<i>Hospitality</i>	54%L	Operasional
C5	37	2.919.108	0,32	Sosial	89% P	Profesional

TABEL VII
PENGELOMPOKAN BIDANG PEKERJAAN

Bidang	Sub Bidang	Contoh Jabatan
Industri	Engineering, Logistik, Keamanan, TI	Teknisi, <i>Welder</i> , Operator, IT <i>Support</i> , <i>Security</i>
Bisnis	Marketing Bisnis, Perbankan, Perkantoran	<i>Sales</i> , <i>Teller</i> , Admin, <i>Customer Service</i>
<i>Hospitality</i>	F&B, Perhotelan	Barista, <i>Cook</i> , <i>Waiter</i> , <i>Resepsionis</i> , <i>Housekeeping</i>
Sosial	Pendidikan, Kesehatan	Guru, Perawat, Bidan, Tutor

D. Hasil Clustering

Hasil evaluasi *clustering* menggunakan metode FCM, Hierarchical Clustering, K-Means, dan GMM pada jumlah *cluster* 2 hingga 5 ditunjukkan pada Tabel V. Evaluasi dilakukan menggunakan Silhouette Score dan DBI untuk menentukan performa *clustering* terbaik.

Berdasarkan Tabel VI, metode FCM dengan jumlah *cluster* lima menghasilkan performa *clustering* terbaik dibandingkan metode lainnya, dengan nilai *Silhouette Score* tertinggi sebesar 0,649 dan nilai DBI terendah sebesar 0,640. Oleh karena itu, metode FCM dengan lima *cluster* dipilih sebagai model *clustering* terbaik pada penelitian ini. Hasil *clustering* FCM menghasilkan lima kelompok data dengan karakteristik yang berbeda berdasarkan atribut numerik dan kategorikal. Ringkasan karakteristik masing-masing *cluster* ditampilkan pada Tabel VI.m

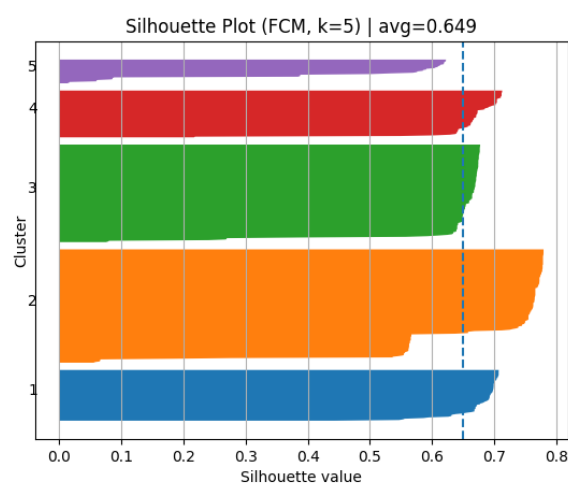
Tabel VII menunjukkan pengelompokan sub-bidang pekerjaan ke dalam empat kategori bidang yang digunakan sebagai atribut kategorikal dalam proses *clustering*. Berdasarkan Tabel VI, *Cluster* C2 menjadi kelompok yang paling potensial untuk diprioritaskan dalam perencanaan pelatihan. *Cluster* ini didominasi oleh bidang Industri dengan jumlah anggota terbesar, yaitu 173 data, serta memiliki gaji rata-rata tertinggi sebesar 3,793,639. Namun, *cluster* ini juga memiliki nilai lama menunggu tertinggi sebesar 0,43. Kondisi tersebut menunjukkan bahwa bidang Industri memiliki prospek pendapatan yang baik, tetapi masih memerlukan dukungan pelatihan yang lebih optimal agar proses penyerapan kerja dapat berlangsung lebih cepat. Oleh karena itu, Pemerintah Kota Surabaya dapat mempertimbangkan peningkatan kuota pelatihan pada bidang Industri, khususnya untuk jabatan operasional, karena bidang ini memiliki potensi ekonomi yang tinggi dan jumlah peserta yang besar. Pendekatan segmentasi berbasis *clustering* semacam ini sejalan dengan penelitian [29] yang menunjukkan bahwa pengelompokan data ketenagakerjaan dapat memberikan rekomendasi kebijakan pelatihan dan reskilling yang lebih tepat sasaran per sektor.

Keempat *cluster* lainnya juga menunjukkan pola yang relevan bagi perumusan kebijakan pelatihan. *Cluster* C1 (78 data, didominasi bidang Sosial dengan komposisi gender seimbang) memiliki gaji rata-rata menengah (Rp3.573.410) dengan masa tunggu relatif singkat (0,23), yang mengindikasikan bahwa jalur masuk kerja di sektor sosial relatif mudah diakses namun kurang kompetitif dari sisi upah, sehingga peserta pada *cluster* ini dapat diarahkan pada pelatihan lanjutan atau sertifikasi keahlian khusus (misalnya keperawatan atau pendidikan bersertifikat) untuk meningkatkan daya tawar gaji. *Cluster* C3 (149 data, didominasi bidang Bisnis dan laki-laki) menunjukkan gaji menengah (Rp3.186.376) dan masa tunggu sedang (0,33), mencerminkan sektor bisnis sebagai jalur yang relatif stabil namun tidak menonjol; pelatihan tambahan pada keterampilan digital atau administrasi dapat membantu peserta *cluster* ini naik ke posisi yang lebih tinggi. *Cluster* C4

(72 data, didominasi bidang Hospitality) memiliki gaji rata-rata terendah (Rp2.830.583) meskipun masa tunggu tergolong singkat (0,22), menunjukkan bahwa sektor Hospitality mudah menyerap tenaga kerja tetapi dengan kompensasi yang kurang kompetitif, sehingga kebijakan pelatihan perlu diarahkan pada peningkatan kompetensi ke posisi yang lebih bernilai tambah, seperti manajemen perhotelan atau pelayanan tamu tingkat lanjut. *Cluster* C5 (37 data, didominasi perempuan pada jabatan Profesional di bidang Sosial) memiliki gaji terendah kedua (Rp2.919.108) dengan masa tunggu sedang (0,32) meskipun menempati jabatan profesional, sebuah pola yang mengindikasikan kesenjangan upah pada jabatan setara dan menjadi perhatian khusus bagi Pemerintah Kota Surabaya untuk mendorong pemerataan kompensasi maupun jenjang karier bagi peserta perempuan pada posisi profesional. Secara keseluruhan, pola lintas *cluster* ini menegaskan bahwa kebijakan pelatihan berbasis data perlu dirancang secara spesifik per sektor, bukan sebagai program yang seragam bagi seluruh peserta.

E. Hasil Evaluasi Kinerja Clustering

Visualisasi Silhouette Plot untuk algoritma FCM dengan jumlah *cluster* sebanyak 5 ditampilkan pada Gambar 2. Berdasarkan Gambar 2, sebagian besar data pada setiap *cluster* memiliki nilai silhouette yang cukup tinggi dengan rata-rata sebesar 0,649. Hal ini menunjukkan bahwa hasil *clustering* menggunakan algoritma FCM memiliki kualitas yang baik, di mana data dalam setiap *cluster* cenderung homogen dan memiliki pemisahan yang cukup jelas antar *cluster*. *Cluster* 2 dan *Cluster* 4 menunjukkan distribusi nilai silhouette yang relatif tinggi dan stabil, sehingga kedua *cluster* tersebut dapat dikatakan terbentuk dengan baik. Selain itu, tidak terdapat nilai silhouette negatif pada seluruh *cluster*, yang menunjukkan bahwa tidak ada data yang salah penempatan atau lebih cocok berada pada *cluster* lain.



Gambar 2. Silhouette Plot Fuzzy C-Means (5 Cluster)

Evaluasi *clustering* berdasarkan nilai DBI menunjukkan bahwa algoritma FCM menghasilkan nilai DBI terendah sebesar 0,640, yang mengindikasikan bahwa *cluster* yang

terbentuk memiliki tingkat kekompakan yang tinggi dan pemisahan antar *cluster* yang optimal.

Perlu dicermati bahwa keunggulan FCM terutama tampak jelas pada nilai DBI (0,640 berbanding 0,872 pada Hierarchical Ward), sementara selisih Silhouette Score antara FCM dan Hierarchical Ward pada lima cluster tergolong sangat kecil (0,649 berbanding 0,647). Karena penelitian ini tidak memiliki label kelompok eksternal sebagai acuan, pengujian signifikansi statistik konvensional (misalnya uji-t) tidak dapat diterapkan secara langsung pada perbandingan skor evaluasi cluster. Sebagai gantinya, kesimpulan bahwa FCM adalah metode terbaik pada penelitian ini didasarkan pada konsistensi keunggulannya pada kedua metrik evaluasi sekaligus (Silhouette Score tertinggi dan DBI terendah), bukan semata-mata pada besarnya selisih nilai Silhouette Score, yang perlu diinterpretasikan secara hati-hati mengingat kedekatannya dengan performa Hierarchical Ward.

Keunggulan FCM pada penelitian ini dipengaruhi oleh pendekatan *soft clustering*, yang memungkinkan setiap data memiliki derajat keanggotaan pada lebih dari satu *cluster*. Karakteristik dataset yang terdiri atas kombinasi atribut numerik dan kategorikal dengan batas antar kelompok yang kurang tegas membuat FCM lebih mampu mengakomodasi data yang berada di area perbatasan cluster dibandingkan metode *hard clustering* seperti K-Means. Sementara itu, Hierarchical Clustering kurang optimal karena pengelompokan berbasis hierarki sensitif terhadap kemiripan antar data, sedangkan GMM mengasumsikan distribusi Gaussian yang tidak sepenuhnya sesuai dengan karakteristik dataset hasil *one-hot encoding*.

Hasil penelitian ini sejalan dengan penelitian sebelumnya [9] yang menyatakan bahwa FCM cenderung memberikan kualitas *clustering* yang lebih baik pada data dengan batas antar *cluster* yang tidak tegas. Selain itu, penggunaan Silhouette Score dan DBI memberikan evaluasi yang saling melengkapi dalam menilai kualitas *clustering*, sehingga semakin memperkuat bahwa FCM merupakan metode *clustering* terbaik pada penelitian ini. Hal ini sesuai dengan penelitian [10] yang menyatakan bahwa Silhouette Score lebih menekankan pada aspek pemisahan antar *cluster*, sedangkan DBI lebih berfokus pada tingkat kekompakan cluster. Dalam penelitian ini, konsistensi hasil yang ditunjukkan oleh kedua metrik tersebut semakin memperkuat bahwa FCM merupakan metode yang paling optimal.

Hasil penelitian ini juga mendukung temuan pada penelitian sebelumnya [8] dan [11] yang menyatakan bahwa tidak ada algoritma *clustering* yang selalu unggul pada semua jenis data. Performa algoritma sangat dipengaruhi oleh karakteristik dataset yang digunakan. Pada penelitian ini, struktur data yang cenderung memiliki variasi antar kelompok menyebabkan algoritma berbasis probabilistik dan *soft clustering* seperti FCM lebih unggul dibandingkan metode lain seperti K-Means, Hierarchical Clustering, maupun GMM. Berbeda dengan penelitian [8] dan [9] yang membandingkan algoritma *clustering* pada domain lain seperti perbankan dan kesehatan, dan berbeda pula dari [29]

yang berfokus pada klasterisasi jenis pekerjaan berdasarkan adopsi AI, penelitian ini secara khusus memposisikan diri pada evaluasi program pelatihan kerja pemerintah dengan membandingkan keempat algoritma pada satu dataset pasca pelatihan yang sama, sehingga memberikan bukti empiris langsung mengenai algoritma *clustering* mana yang paling sesuai untuk mendukung evaluasi dan perumusan kebijakan pelatihan berbasis data di lingkup pemerintah daerah, sebuah konteks yang belum banyak dieksplorasi pada literatur *clustering* sebelumnya.

IV. KESIMPULAN

Hasil penelitian menunjukkan bahwa algoritma FCM menghasilkan performa terbaik dibandingkan K-Means, Hierarchical Clustering, dan GMM dengan nilai Silhouette Score tertinggi sebesar 0,649 dan DBI terendah sebesar 0,640. Keunggulan tersebut menunjukkan bahwa FCM lebih efektif dalam mengelompokkan data yang memiliki batas antar cluster kurang tegas, sehingga mampu menghasilkan cluster yang lebih kompak dan terpisah dengan baik. Penelitian selanjutnya dapat mengeksplorasi metode *clustering* lainnya atau menerapkan teknik seleksi fitur untuk meningkatkan kualitas pengelompokan pada dataset dengan karakteristik yang lebih kompleks. Hasil segmentasi juga menunjukkan bahwa bidang Industri memberikan potensi gaji tertinggi namun dengan masa tunggu kerja terlama, sedangkan bidang *Hospitality* dan sebagian jabatan Profesional perempuan di sektor Sosial memiliki gaji terendah meskipun masa tunggu relatif singkat, sehingga Pemerintah Kota Surabaya disarankan merancang intervensi pelatihan yang berbeda untuk tiap kelompok, alih-alih program yang seragam. Penelitian ini memiliki keterbatasan pada cakupan variabel yang digunakan, karena analisis hanya menggunakan atribut yang tersedia dari Badan Riset dan Inovasi Daerah (BRIDA) Kota Surabaya. Penelitian selanjutnya disarankan menyertakan variabel yang lebih beragam dan representatif, seperti jenis pelatihan, tingkat pendidikan, usia, lama pelatihan, maupun kompetensi peserta, agar evaluasi keberhasilan program pelatihan menjadi lebih menyeluruh. Selain itu, validasi hasil *clustering* bersama pakar atau pihak pengelola program pelatihan juga disarankan pada penelitian selanjutnya, agar interpretasi karakteristik cluster yang dihasilkan memiliki validitas dan tingkat kepercayaan yang lebih tinggi sebelum diterapkan sebagai dasar pengambilan kebijakan.

DAFTAR PUSTAKA

- [1] S. G. Pane, W. Pramudya, R. C. Amalya, S. N. Aulia, dan P. Nabila Pebriani, "Analisis Pengaruh Tenaga Kerja dan Pengangguran Terhadap Pertumbuhan Ekonomi di Indonesia," *Econ. Rev. J.*, vol. 3, hal. 1204–1214, 2024, doi: 10.56709/mrj.v3i4.401.
- [2] P. Sari, Efan, dan R. Syahri, "Algoritma K-Means Clustering: Sebuah Studi Literatur," *J. Inform.*, hal. 1–7, 2023, doi: 10.12345/juri.
- [3] R. Rahmati, A. W. Wijayanto, P. Studi, K. Statistik, dan P. Sains, "Analisis Cluster dengan Algoritma K-Means, Fuzzy C-means dan

- Hierarchical *Clustering* (Studi Kasus: Indeks Pembangunan Manusia Tahun 2019),” *J. Inform. dan Komput.*, vol. 5, no. 2, hal. 73–80, 2021, doi: <https://dx.doi.org/10.26798/jiko.v5i2.422>.
- [4] H. Yang, C. Wang, H. Zhang, Y. Zhou, dan B. Luo, “Recognition of Maize Seed Varieties based on Hyperspectral Imaging Technology and Integrated Learning Algorithms,” *PeerJ Comput. Sci.*, hal. 1–20, 2023, doi: [10.7717/peerj-cs.1354](https://doi.org/10.7717/peerj-cs.1354).
- [5] H. Huang, Z. Liao, X. Wei, dan Y. Zhou, “Combined Gaussian Mixture Model and Pathfinder Algorithm,” *entropy*, hal. 1–20, 2023, doi: <https://doi.org/10.3390/e25060946>.
- [6] H. Firdaus dan A. Sofro, “Analisa *Cluster* menggunakan K-Means dan Fuzzy C-Means dalam Pengelompokan Provinsi Menurut Data Intesitas Bencana Alam di Indonesia Tahun 2017-2021,” *J. Ilm. Mat.*, vol. 10, no. 01, hal. 50–60, 2022, doi: <https://doi.org/10.26740/mathunesa.v10n1.p50-60>.
- [7] E. Y. Ahmadv, “Comparative Analysis of K-Means and Fuzzy C-Means Algorithms on Demographic Data using the PCA Method,” *Probl. Inf. Technol.*, vol. 14, no. 1, hal. 15–22, 2023, doi: <http://doi.org/10.25045/jpit.v14.i1.03>.
- [8] A. A. Wani, “Comprehensive Analysis of *Clustering* Algorithms : Exploring Limitations and Innovative Solutions,” *PeerJ Comput. Sci.*, 2024, doi: [10.7717/peerj-cs.2286](https://doi.org/10.7717/peerj-cs.2286).
- [9] E. K. A. Mala, S. Rochman, dan H. Suprajitno, “Comparison of *Clustering* in Tuberculosis using Fuzzy C-Means and K-Means Methods,” *Commun. Math. Biol. Neurosci.*, hal. 1–20, 2022, doi: <https://doi.org/10.28919/cmbn/7335>.
- [10] D. Chicco, A. Campagner, A. Spagnolo, D. Ciucci, dan G. Jurman, “The Silhouette Coefficient and the Davies-Bouldin Index are more Informative than Dunn index, Calinski-Harabasz index, Shannon entropy, and Gap Statistic for Unsupervised *Clustering* Internal Evaluation of Two Convex Clusters,” 2025, doi: [10.7717/peerj-cs.3309](https://doi.org/10.7717/peerj-cs.3309).
- [11] M. K. S. Md Amran Hossen Pabel, Biswanath Bhattacharjee, Sonjoy Kumar Dey, Sakib Salam Jamee, Md Omar Obaid, Md Sakib Mia, Sajidul Islam Khan, “Business Analytics for Customer Segmentation: A Comparative Study of Machine Learning Algorithms in Personalized Banking Services,” *International J. Econ. Financ. Manag. Sci.*, hal. 1–13, doi: <https://doi.org/10.55640/ijefms/Volume10Issue03-01>.
- [12] Z. Parinzka, S. Surono, dan A. Thobirin, “Algoritma Support Vector Regression dan Analisis Long Short-Term Memory sebagai Penanganan Missing data,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 13, no. 1, hal. 73–82, 2026, doi: <https://doi.org/10.25126/jtiik.2026131>.
- [13] Fahrillah dan Z. Fatah, “Pengelompokan Data Nilai Siswa Madrasah Ta’hiliyah menggunakan Metode K-means Clustering,” *J. Ris. Sist. Inf.*, vol. 2, no. 1, hal. 53–59, 2025, doi: <https://doi.org/10.69714/0v1pkz05>.
- [14] A. Zhu, Z. Hua, Y. Shi, dan Y. Tang, “An Improved K-Means Algorithm Based on Evidence Distance,” *entropy*, 2021, doi: <https://doi.org/10.3390/e23111550>.
- [15] M. Sania Fitri Octavia, “Penerapan K-Means dan Fuzzy C-Means untuk Pengelompokan Data Kasus Covid-19 di Kabupaten Indragiri Hilir,” *Build. Informatics, Technol. Sci.*, vol. 3, no. 2, hal. 88–94, 2021, doi: [10.47065/bits.v3i2.1005](https://doi.org/10.47065/bits.v3i2.1005).
- [16] F. A. Totti dan N. Setiyawati, “Perbandingan Algoritma *Clustering* K-Means, Gaussian Mixture Model, dan Spectral *Clustering* untuk Facial Emotion Recognition A Comparative Study of K-Means, Gaussian Mixture Model, and,” *J. Sist. Inf.*, vol. 14, hal. 3007–3019, 2025, doi: <https://doi.org/10.32520/stmsi.v14i6.5668>.
- [17] S. H. Marwoto, “A Comparative Analysis of DBSCAN and Gaussian Mixture Model for *Clustering* Indonesian Provinces Based on Socioeconomic Welfare Indicators,” *J. Ilmu Mat. dan Terap.*, vol. 19, no. 3, hal. 2039–2056, 2025, doi: <https://doi.org/10.30598/barekengvol19iss3pp2039-2056>.
- [18] V. C. Pezoulas *et al.*, “Bayesian Inference-Based Gaussian Mixture Estimation Towards Large-Scale Synthetic Data Generation for In Silico Clinical Trials,” *Eng. Med. Biol.*, vol. 3, hal. 108–114, 2022, doi: <https://doi.org/10.1109/OJEMB.2022.3181796>.
- [19] D. I. Yulianti, T. I. Hermanto, dan M. Defriani, “Analisis *Clustering* Donor Darah dengan Metode Agglomerative Hierarchical Clustering,” *J. Rekayasa Tek. Inform. dan Inf.*, vol. 3, no. 6, hal. 303–308, 2023, doi: <https://djournals.com/resolusi.2111>.
- [20] S. S. Brigitta Melati Kumarahadi, Hasih Pratiwi, “Penerapan Metode Hierarchical *Clustering* untuk Pengelompokan Kota/Kabupaten di Indonesia berdasarkan Indikator Kemiskinan,” vol. 11, no. 2, 2023, doi: <https://doi.org/10.30646/tikomsin.v11i2.754>.
- [21] Z. Alamtaha, I. Djakaria, dan N. I. Yahya, “Implementasi Algoritma Hierarchical *Clustering* dan Non-Hierarchical *Clustering* untuk Pengelompokan Pengguna Media Sosial,” *J. Stat. an Its Appl.*, vol. 4, no. 1, hal. 33–43, 2023, doi: [10.20956/ejsa.vi.24830](https://doi.org/10.20956/ejsa.vi.24830).
- [22] F. Rasidia, R. Goejantoro, M. Fathurahman, dan L. S. Komputasi, “Analisis Klaster Menggunakan Metode Average Linkage dengan Validasi Multiscale Bootstrap (Studi Kasus: Indikator Pendidikan di Indonesia Tahun 2021),” *J. Ekspansional*, vol. 16, no. April, hal. 23–31, 2025, doi: [10.30872/ekspansional.v16i1.1392](https://doi.org/10.30872/ekspansional.v16i1.1392).
- [23] R. Gustriansyah, J. Alie, dan N. Suhandi, “Hierarchical *Clustering* for Crime Rate Mapping in Indonesia,” *J. Ilm.*, vol. 14, no. 3, hal. 275–283, 2022.
- [24] A. Septianingsih, “Pemetaan Kabupaten Kota di Provinsi Jawa Timur berdasarkan Tingkat Kasus Penyakit menggunakan Pendekatan Agglomeratif Hierarchical Clustering,” *J. Ilm. Pendidik. Mat. Mat. dan Stat.*, vol. 3, no. 2, 2022.
- [25] A. M. Sarusu, D. Akmila, M. Wijana, dan M. E. Habiby, “Sistem Informasi Manajemen Data Penduduk Berbasis Website,” *Intern. Inf. Syst. J.*, vol. 6, no. 2, hal. 127–136, 2024, doi: <https://doi.org/10.32627/internal.v6i2.859>.
- [26] H. Ramadhan, M. Rizal, A. Kamaludin, M. A. Nasrullah, dan D. Rolliawati, “Comparison of Hierarchical, K-Means and DBSCAN *Clustering* Methods for Credit Card Customer Segmentation Analysis Based on Expenditure Level,” *J. Appl. Informatics Comput.*, vol. 7, no. 2, hal. 246–251, 2023, doi: <https://doi.org/10.30871/jaic.v7i2.5790>.
- [27] D. A. Tarigan, “Optimization of the K-Means *Clustering* Algorithm Using Davies Bouldin Index in Iris Data Classification,” vol. 4, no. 1, hal. 545–552, 2023, doi: [10.30865/klik.v4i1.964](https://doi.org/10.30865/klik.v4i1.964).
- [28] Y. Lin dan S. Chen, “A Centroid Auto-Fused Hierarchical Fuzzy c-Means Clustering,” *IEEE Trans. Fuzzy Syst.*, vol. 29, no. 7, hal. 2006–2017, 2021, doi: [10.1109/TFUZZ.2020.2988848](https://doi.org/10.1109/TFUZZ.2020.2988848).
- [29] M. S. Hasibuan, R. Rizal Nul Fikri, dan D. A. Dewi, “Job Clustering Based on AI Adoption and Automation Risk Levels: An Analysis Using the K-Means Algorithm in the Technology and Entertainment Industries,” *Int. J. Appl. Inf. Manag.*, vol. 4, no. 2, hal. 54–69, 2024, doi: [10.47738/ijaim.v4i2.83](https://doi.org/10.47738/ijaim.v4i2.83).