

Studi Literatur: Perbandingan Algoritma K-Means dan Fuzzy C-Means dalam Analisis Clustering

Jaelani¹, Octaviana², Elkin Rilvani³

^{1,2,3} Universitas Pelita Bangsa, Bekasi,

jthebluess@gmail.com, rumapeaocta@gmail.com, elkin.rilvani@pelitabangsa.ac.id

Abstract

The increasing volume and complexity of data across sectors such as healthcare, business, education, and security necessitates analytical methods capable of handling uncertainty and structural variability. One widely used approach is clustering, which groups data based on similarity. Two commonly applied algorithms are K-Means and Fuzzy C-Means (FCM), each with distinct characteristics: K-Means applies hard clustering, while FCM adopts soft clustering using degrees of membership. This study presents a literature review of both national and international scientific journals published from 2020 to 2025. The selected studies apply clustering algorithms in contexts such as employee evaluation, spatial disease mapping, patient classification, and customer segmentation. In addition to 12 national journals, this review incorporates international perspectives (e.g., Shahapure & Nicholas, 2020; Abirami & Chitra, 2021; Ma, Zhang, & Li, 2021), which reinforce the theoretical foundation. The findings show that algorithm selection depends on data characteristics and analytical goals. K-Means is computationally efficient and easy to interpret, while FCM is more flexible for overlapping and complex data. Several recent studies also recommend hybrid approaches combining both algorithms to enhance accuracy and robustness. Therefore, this review aims to serve as a credible academic and practical reference for selecting appropriate clustering techniques.

Keywords: *Clustering, K-Means, Fuzzy C-Means, Literature Review, Soft Clustering, Data Analysis*

Abstrak

Peningkatan volume dan kompleksitas data di berbagai sektor seperti kesehatan, bisnis, pendidikan, dan keamanan mendorong kebutuhan akan metode analisis yang mampu menangani ketidakpastian dan variasi struktur data. Salah satu pendekatan yang umum digunakan adalah clustering, yaitu teknik pengelompokan data berdasarkan kemiripan. Dua algoritma yang sering diterapkan adalah K-Means dan Fuzzy C-Means (FCM), dengan karakteristik berbeda: K-Means menggunakan pendekatan hard clustering, sementara FCM menerapkan soft clustering dengan derajat keanggotaan. Penelitian ini merupakan studi literatur terhadap jurnal ilmiah nasional dan internasional yang diterbitkan antara tahun 2020 hingga 2025. Studi-studi tersebut membahas implementasi algoritma clustering dalam berbagai konteks seperti evaluasi kinerja, pemetaan penyakit secara spasial, klasifikasi pasien, dan segmentasi pelanggan. Selain merujuk pada 12 jurnal nasional, tinjauan ini diperkuat dengan perspektif internasional terkini (misalnya: Shahapure & Nicholas, 2020; Abirami & Chitra, 2021; Ma, Zhang, & Li, 2021) untuk memperkuat landasan teoritis dan aplikatif. Hasil studi menunjukkan bahwa pemilihan algoritma sangat dipengaruhi oleh karakteristik data dan tujuan analisis. K-Means unggul dalam efisiensi komputasi dan interpretasi yang sederhana, sementara FCM lebih fleksibel untuk data yang kompleks dan tumpang tindih. Sejumlah studi juga merekomendasikan pendekatan hibrida untuk meningkatkan akurasi dan ketahanan. Oleh karena itu, kajian ini diharapkan menjadi referensi akademik dan praktis dalam memilih metode clustering yang tepat.

Kata Kunci: *Clustering, K-Means, Fuzzy C-Means, Studi Literatur, Soft Clustering, Analisis Data*

1. Pendahuluan

Perkembangan pesat teknologi informasi menyebabkan peningkatan volume dan kompleksitas data di berbagai

bidang, seperti kesehatan, bisnis, pendidikan, hingga keamanan. Hal ini mendorong kebutuhan akan metode analisis yang efisien dan adaptif. Salah

satu pendekatan penting dalam unsupervised learning adalah clustering, yaitu teknik pengelompokan data berdasarkan kemiripan tanpa label kelas yang eksplisit (Ma, Zhang, & Li, 2021). Teknik ini memungkinkan identifikasi pola tersembunyi dalam data, yang sangat relevan dalam era big data dan kecerdasan buatan.

Dalam konteks data mining, clustering merupakan salah satu teknik eksploratif yang banyak digunakan untuk segmentasi pasar, analisis spasial, klasifikasi citra, dan pengenalan pola (Shahapure & Nicholas, 2020). Dua algoritma clustering yang sering diaplikasikan adalah K-Means dan Fuzzy C-Means (FCM). K-Means menerapkan pendekatan hard clustering di mana setiap data hanya masuk ke satu kluster, sedangkan FCM memungkinkan soft clustering dengan memberikan derajat keanggotaan pada beberapa kluster, sehingga lebih fleksibel terhadap ambiguitas data (Abirami & Chitra, 2021).

Berbagai studi literatur telah membahas keunggulan dan kelemahan kedua algoritma ini. Misalnya, penelitian oleh Syaifudin et al. (2023) menunjukkan bahwa FCM lebih unggul dalam klasifikasi pelanggan berdasarkan pola perilaku yang kompleks. Sedangkan Rahmansyah et al. (2023) menekankan keunggulan FCM dalam mengklasifikasikan status gizi balita yang memiliki ambiguitas tinggi. Sebaliknya, Rosyada (2021) dan Saputri & Wulandari (2020) menggunakan K-Means untuk pemetaan spasial penyakit karena efisiensi komputasinya yang tinggi.

Meskipun keduanya populer, efektivitas dan efisiensinya sangat tergantung pada jenis data dan tujuan analisis. Dalam beberapa studi, K-Means terbukti unggul dalam kecepatan komputasi dan interpretasi hasil yang sederhana (Arora, Varshney, & Yadav, 2020), sementara FCM lebih akurat untuk data yang tumpang tindih atau kompleks (Abirami & Chitra, 2021).

Berdasarkan paparan di atas, dapat disimpulkan bahwa terdapat perbedaan karakteristik dan performa antara algoritma clustering K-Means dan FCM. Oleh karena itu, penelitian ini bertujuan untuk meninjau dan membandingkan efektivitas kedua algoritma melalui studi literatur terbaru yang relevan dengan konteks aplikasi nyata. Pertanyaan penelitian yang diajukan adalah: **“Bagaimana efektivitas dan kecocokan algoritma clustering K-**

Means dan FCM berdasarkan karakteristik data dan tujuan analisisnya?”

Dengan meninjau berbagai studi ilmiah nasional dan internasional terkini, tulisan ini diharapkan dapat memberikan kontribusi teoretis dan praktis, khususnya bagi peneliti dan praktisi yang ingin menerapkan metode clustering dalam analisis data kompleks. Selain itu, penulisan ini memperkuat posisi data mining sebagai pendekatan strategis dalam pengambilan keputusan berbasis data.

2. Landasan Teori

2.1. Teori Dasar Clustering

Clustering merupakan metode *unsupervised learning* yang bertujuan mengelompokkan data berdasarkan kemiripan tanpa menggunakan label kelas. Setiap kluster merepresentasikan data dengan karakteristik serupa, sedangkan antar kluster dijaga tetap berbeda secara signifikan. Metode ini banyak diterapkan dalam segmentasi pasar, klasifikasi citra, pengenalan pola, dan analisis data spasial. Tujuan utamanya adalah meminimalkan variasi dalam kluster (*intra-cluster distance*) dan memaksimalkan perbedaan antar kluster (*inter-cluster distance*).

2.2. Algoritma K-Means

K-Means adalah algoritma hard clustering yang membagi data ke dalam k cluster berdasarkan jarak terdekat ke pusat cluster (centroid). Langkah-langkah dalam algoritma ini meliputi:

1. Menentukan jumlah cluster (k)
2. Menginisialisasi centroid awal secara acak
3. Mengelompokkan data ke cluster terdekat berdasarkan jarak Euclidean
4. Menghitung ulang centroid berdasarkan rata-rata data dalam masing-masing cluster
5. Mengulangi proses hingga centroid tidak berubah lagi (konvergen)

Keunggulan K-Means adalah kemudahan implementasi dan efisiensi waktu komputasi. Namun, algoritma ini sensitif terhadap inisialisasi centroid dan kurang efektif untuk data dengan

bentuk non-linier atau distribusi tidak seimbang.

2.3. Algoritma Fuzzy C-Means

Fuzzy C-Means (FCM) merupakan algoritma soft clustering yang memungkinkan satu data memiliki derajat keanggotaan terhadap beberapa cluster sekaligus, dengan nilai keanggotaan antara 0 dan 1. Prosedur FCM mencakup:

1. Inisialisasi jumlah cluster (c) dan matriks derajat keanggotaan
2. Menghitung pusat cluster dengan mempertimbangkan derajat keanggotaan
3. Memperbarui derajat keanggotaan berdasarkan jarak data ke masing-masing pusat cluster
4. Melakukan iterasi hingga nilai keanggotaan stabil

Kelebihan FCM terletak pada kemampuannya menangani ambiguitas dalam data dan mendeteksi struktur cluster yang tumpang tindih. Kekurangannya adalah lebih kompleks secara komputasi dan sensitif terhadap nilai parameter awal.

2.4. Studi-Studi Terdahulu yang Relevan

Berikut adalah beberapa studi terdahulu yang menjadi dasar perbandingan algoritma K-Means dan FCM:

- 1) Hamdi Syaifudin et al. (2023): Menganalisis perilaku pelanggan dan menyimpulkan FCM memberikan segmentasi yang lebih fleksibel.
- 2) Akim Kurnia Rahmansyah et al. (2023): Mengelompokkan status gizi balita; FCM menghasilkan pemisahan yang lebih baik untuk data tidak tegas.
- 3) Pradana Putra et al. (2021): Mengembangkan sistem pendukung keputusan investasi dengan metode clustering untuk mengelompokkan profil risiko.
- 4) Anissa Enggar Pramitasari (2021) & Andika Dwi Saputra (2022): Mengkaji kinerja pegawai dengan hasil bahwa FCM lebih adaptif pada data evaluasi kerja yang kompleks.

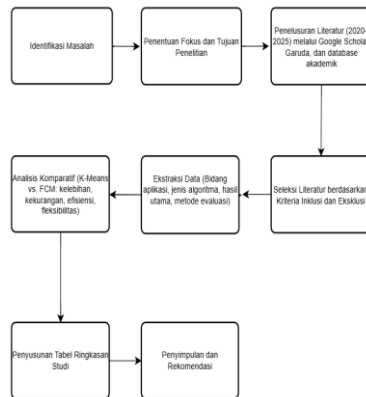
- 5) Franklyn & Nataliani (2022): Menganalisis performa atlet basket dan menunjukkan bahwa kombinasi K-Means dan FCM memberi hasil yang lebih informatif.
- 6) Amalia Rosyada (2021) & Mega Saputri & Niken Wulandari (2020): Menggunakan K-Means untuk memetakan persebaran penyakit, menekankan efisiensi algoritma terhadap data spasial.
- 7) Aina Sugiarta Firdusie et al. (2021): Menunjukkan efektivitas FCM dalam pengelompokan wilayah rawan kriminalitas.
- 8) Adelia Indira Maharani (2023): Menerapkan FCM untuk clustering data imunisasi, menunjukkan akurasi dan konsistensi hasil yang tinggi.
- 9) Teguh Triyanto & Yunita Eka Putri (2021): Menggabungkan metode fuzzy dan algoritma genetika untuk deteksi penyakit mata, menyoroti potensi kombinasi metode.

Studi-studi di atas menunjukkan bahwa konteks aplikasi, jenis data, serta tujuan analisis sangat memengaruhi performa masing-masing algoritma. Oleh karena itu, pemilihan metode clustering harus mempertimbangkan karakteristik data dan tujuan akhir dari pengelompokan.

3. Metodologi Penelitian

3.1. Jenis Penelitian

Penelitian ini menggunakan pendekatan studi literatur (literature review) yang bertujuan untuk menganalisis dan membandingkan algoritma K-Means dan Fuzzy C-Means berdasarkan temuan dari berbagai sumber ilmiah. Kajian literatur ini bersifat deskriptif-kualitatif, di mana data diperoleh dari jurnal-jurnal yang telah dipublikasikan dan dianalisis secara sistematis.



Gambar 3.1 Alur Tahapan Penelitian Studi Literatur

3.2. Sumber dan Teknik Pengumpulan Data

Sumber data dalam penelitian ini diperoleh dari 12 jurnal ilmiah nasional yang relevan, diterbitkan dalam rentang waktu 2020 hingga 2023. Jurnal-jurnal tersebut diakses melalui basis data daring seperti Garuda, Google Scholar, dan portal akademik lainnya. Teknik pengumpulan data dilakukan dengan menelusuri kata kunci seperti “K-Means,” “Fuzzy C-Means,” dan “clustering.”

3.3. Kriteria Seleksi Literatur

Dalam penyaringan literatur, digunakan beberapa kriteria inklusi dan eksklusi sebagai berikut:

- 1) Kriteria inklusi:
 - a) Jurnal ilmiah yang membahas implementasi algoritma K-Means dan/atau FCM
 - b) Studi dengan konteks aplikasi nyata (data pelanggan, kesehatan, kinerja pegawai, dll.)
 - c) Diterbitkan antara tahun 2020–2023
- 2) Kriteria eksklusi:
 - a) Jurnal yang hanya membahas teori tanpa implementasi algoritma
 - b) Artikel non-peer-reviewed atau opini

3.4. Teknik Analisis Data

Data yang terkumpul dari jurnal-jurnal tersebut dianalisis dengan pendekatan komparatif. Analisis dilakukan terhadap aspek-aspek berikut:

- 1) Tujuan dan konteks penelitian

- 2) Algoritma yang digunakan
- 3) Jenis data dan variabel
- 4) Hasil analisis dan evaluasi performa
- 5) Kelebihan dan kekurangan masing-masing algoritma

Analisis akan disajikan secara naratif dan dilengkapi dengan tabel ringkasan guna memperjelas perbandingan antar studi yang dikaji.

4. Hasil dan Pembahasan

Penelitian ini menganalisis 12 jurnal ilmiah yang mengimplementasikan algoritma K-Means dan Fuzzy C-Means (FCM) dalam berbagai konteks aplikasi. Berdasarkan hasil telaah, ditemukan bahwa pemilihan algoritma sangat dipengaruhi oleh jenis data, tujuan analisis, dan kompleksitas permasalahan.

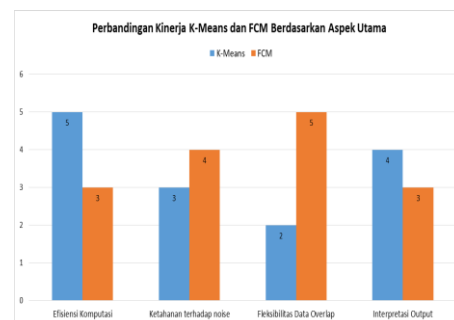
4.1. Penggunaan Algoritma dalam Studi-Studi Terdahulu

Dari jurnal-jurnal yang ditelaah:

1. 4 studi menggunakan algoritma K-Means secara tunggal.
2. 3 studi menggunakan FCM sebagai metode utama.
3. 5 studi membandingkan atau mengombinasikan kedua algoritma tersebut.

4.2. Perbandingan Karakteristik K-Means dan FCM

Dari segi efisiensi, K-Means unggul karena memiliki waktu komputasi yang relatif lebih cepat dan struktur algoritma yang sederhana. Namun, dalam hal fleksibilitas, terutama pada data yang mengandung ambiguitas atau cluster yang saling tumpang tindih, FCM lebih mampu merepresentasikan keanggotaan data secara bertingkat.



Gambar 4.2 Perbandingan Kinerja K-Means dan Fuzzy C-Means Berdasarkan Aspek Utama.

4.3. Kesesuaian Algoritma terhadap Bidang Aplikasi

Berdasarkan analisis bidang, K-Means banyak digunakan dalam studi yang bersifat spasial dan dengan batas cluster yang jelas, seperti penyebaran COVID-19 (Rosyada, 2021) dan wilayah penyakit TBC (Saputri & Wulandari, 2020). Sebaliknya, FCM digunakan pada bidang yang memiliki kompleksitas variabel tinggi, seperti evaluasi kinerja pegawai (Pramitasari, 2021), atau klasifikasi status gizi (Rahmansyah, 2023).

Distribusi bidang studi adalah sebagai berikut:

1. Kesehatan: 5 studi
2. Kinerja/Evaluasi Pegawai: 3 studi
3. Kriminalitas & Keamanan: 2 studi
4. Pelanggan/Bisnis: 2 studi

4.4. Temuan Utama

Berdasarkan hasil analisis terhadap 12 jurnal yang dikaji, terdapat beberapa temuan penting yang menunjukkan karakteristik, kelebihan, dan kekurangan algoritma K-Means dan Fuzzy C-Means (FCM) dalam berbagai aplikasi:

1. Efisiensi Proses Komputasi K-Means unggul dari sisi kecepatan dan efisiensi proses, cocok untuk dataset besar dan aplikasi spasial seperti pemetaan penyakit (Rosyada, 2021; Saputri & Wulandari, 2020).
2. Fleksibilitas dan Kejelasan Pengelompokan FCM lebih efektif dalam menangani data dengan batas kluster yang tidak tegas atau tumpang tindih, seperti pada evaluasi kinerja dan status gizi (Pramitasari, 2021; Rahmansyah et al., 2023).
3. Tingkat Akurasi dan Interpretabilitas FCM menghasilkan akurasi lebih tinggi untuk data kompleks, meskipun interpretasinya lebih menantang. Sebaliknya, K-Means lebih mudah dipahami, namun kurang akurat untuk data yang overlap.
4. Keselarasan terhadap Bidang Aplikasi K-Means sering

digunakan dalam studi spasial dan klasifikasi cepat, sedangkan FCM lebih sesuai untuk analisis evaluatif dan multidimensional.

5. Penggunaan Gabungan Algoritma -Means sering digunakan dalam studi spasial dan klasifikasi cepat, sedangkan FCM lebih sesuai untuk analisis evaluatif dan multidimensional.

Untuk memperjelas perbandingan antara K-Means dan FCM, berikut disajikan tabel yang merangkum karakteristik teknis masing-masing algoritma berdasarkan aspek analisis utama:

Aspek	K-Means	Fuzzy C-Means (FCM)
Tipe Klastering	Hard clustering (tegas satu kluster)	Soft clustering (keanggotaan ganda)
Efisiensi Komputasi	Tinggi (cepat & ringan)	Lebih kompleks, lambat pada data besar
Toleransi terhadap Overlapping	Rendah	Tinggi
Interpretasi Hasil	Sederhana dan langsung	Lebih fleksibel namun interpretatif
Ketahanan terhadap Noise	Rentan terhadap outlier	Relatif lebih tahan tergantung nilai fuzzy
Cocok untuk	Segmentasi sederhana, spasial	Evaluasi kompleks, multidimensi, sosial

Tabel 1. Perbandingan Karakteristik K-Means dan Fuzzy C-Means

Selanjutnya, untuk memperkuat temuan dari studi literatur, Tabel di bawah ini menyajikan ringkasan dari jurnal yang ditelaah dalam penelitian ini. Tabel mencantumkan nama penulis, bidang kajian, algoritma yang digunakan, dan hasil utama dari masing-masing studi.

No	Penulis & Tahun	Bidang Aplikasi	Algoritma	Hasil & Temuan
1	Hamdi Syaifudin et al. (2023)	Segmentasi Pelanggan	K-Means, FCM	FCM memberikan segmentasi perilaku konsumen yang lebih detail.
2	Akim K. Rahmansyah et al. (2023)	Gizi Balita	K-Means, FCM	FCM lebih akurat dalam menangani data status gizi yang tidak tegas.
3	Pradana Putra et al. (2021)	Investasi	K-Means	K-Means efektif dalam mendukung pengambilan keputusan investasi.
4	Nirmalasari et al. (2023)	Klasifikasi Obat	K-Means	K-Means berhasil mengelompokkan obat berdasarkan fitur komposisi.
5	Anissa E. Pramitasari (2021)	Kinerja Pegawai	FCM	FCM akurat dalam menangani penilaian kerja multidimensi.
6	Franklyn & Nataliani (2022)	Performa Atlet	K-Means, FCM	Kombinasi dua algoritma memberi wawasan performa yang lebih mendalam.
7	Andika Dwi Saputra (2022)	Evaluasi Kinerja Swasta	K-Means	K-Means efektif untuk penilaian kerja dengan parameter sederhana.
8	Aina S. Firdusie et al. (2021)	Wilayah Kriminalitas	FCM	FCM mampu mengelompokkan area rawan berdasarkan tingkat kejahatan.
9	Teguh T. & Yunita E. Putri (2021)	Medis (Retinopati)	Fuzzy, Genetik	Fuzzy logic efektif dalam deteksi penyakit, khususnya fitur gejala kompleks.
10	Amalia Rosyada (2021)	COVID-19 (Spasial)	K-Means	K-Means cepat dan efisien dalam pemetaan kasus secara spasial.
11	Mega Saputri & Niken Wulandari (2020)	Penyakit TBC	K-Means	K-Means membantu mengidentifikasi pola persebaran penyakit.
12	Adelia I. Maharani (2023)	Data Imunisasi	FCM	FCM unggul dalam klasifikasi data imunisasi yang tumpang tindih antar variabel.

Tabel 2. Ringkasan Studi Terdahulu Mengenai Implementasi K-Means dan FCM5.

Meskipun hasil analisis ini memberikan gambaran yang cukup komprehensif mengenai perbandingan K-Means dan Fuzzy C-Means berdasarkan 12

jurnal nasional, validitas dan generalisasi temuan masih terbatas. Hal ini disebabkan karena sebagian besar sumber literatur berasal dari jurnal lokal yang tidak seluruhnya terindeks pada basis data internasional seperti Scopus atau Web of Science. Oleh karena itu, untuk meningkatkan bobot akademik dan kekuatan kesimpulan, disarankan agar penelitian selanjutnya mengembangkan kajian literatur dengan melibatkan lebih banyak jurnal internasional bereputasi yang relevan dan terbaru (2020–2025). Dengan demikian, hasil analisis akan lebih representatif terhadap praktik global dan dapat dijadikan acuan dalam studi lanjutan atau implementasi sistem cerdas berbasis clustering.

5. Kesimpulan dan Saran

Berdasarkan kajian terhadap berbagai jurnal ilmiah nasional dan internasional yang diterbitkan antara tahun 2020 hingga 2025, disimpulkan bahwa algoritma K-Means dan Fuzzy C-Means (FCM) memiliki keunggulan yang saling melengkapi dan tidak bersifat superior secara mutlak. Pemilihan algoritma clustering sangat tergantung pada karakteristik data, tingkat kompleksitas, keberadaan noise, dan tujuan analisis.

K-Means unggul dari sisi efisiensi komputasi dan kesederhanaan implementasi, sehingga lebih cocok untuk data yang homogen, linear, dan memiliki batas klaster yang jelas, seperti pada pemetaan spasial atau segmentasi lokasi. Sebaliknya, FCM lebih sesuai untuk data yang kompleks dan tumpang tindih, seperti evaluasi kinerja multidimensi atau analisis perilaku pelanggan, karena kemampuannya dalam mengakomodasi keanggotaan ganda.

Sejumlah studi juga menunjukkan bahwa pendekatan gabungan antara K-Means dan FCM, atau integrasinya dengan metode lain seperti fuzzy-genetik, dapat meningkatkan akurasi dan ketahanan pengelompokan.

Untuk penelitian lanjutan, disarankan dilakukan studi kuantitatif eksperimental guna mengukur akurasi, kecepatan, dan stabilitas algoritma dalam berbagai jenis dataset. Selain itu, eksplorasi terhadap algoritma alternatif seperti DBSCAN, Mean Shift, dan pendekatan hybrid juga penting sebagai opsi pengembangan lebih lanjut dalam membangun sistem cerdas berbasis clustering.

6. Daftar Pustaka

- Adelia, I. M. (2023). *Clustering data imunisasi dengan FCM*. Jurnal Sistem Informasi Kesehatan.
- Aina, S. F., & rekan. (2021). *Clustering wilayah rawan kriminalitas dengan FCM*. Jurnal Ilmu Komputer.
- Amalia, R. (2021). *Identifikasi persebaran kasus COVID-19 menggunakan K-Means*. Jurnal Kesehatan Digital.
- Andika, D. S. (2022). *Evaluasi kinerja pegawai swasta menggunakan clustering*. Jurnal Teknologi Informasi.
- Anissa, E. P. (2021). *Clustering kinerja pegawai menggunakan Fuzzy C-Means*. Jurnal Sains dan Teknologi.
- Akim, K. R., & rekan. (2023). *Clustering status gizi balita menggunakan K-Means dan FCM*. Jurnal Ilmiah Komputer.
- Franklyn, & Nataliani. (2022). *Analisis performa pemain basket dengan K-Means dan FCM*. Jurnal Statistika.
- Hamdi, S., & rekan. (2023). *Perbandingan K-Means dan Fuzzy C-Means pada data pelanggan*. Jurnal Teknologi Informasi.
- Mega, S., & Wulandari, N. (2020). *Clustering wilayah penyakit TBC*. Jurnal Kesehatan Masyarakat.
- Nirmalasari, & rekan. (2023). *Klasifikasi jenis obat dengan K-Means*. Jurnal Informatika dan Komputer.
- Pradana, P., & rekan. (2021). *Sistem pendukung keputusan investasi dengan clustering*. Journal of Information System.
- Teguh, T., & Putri, Y. E. (2021). *Deteksi diabetic retinopathy dengan fuzzy dan algoritma genetik*. Jurnal Diagnostik Digital.
- Ma, H., Zhang, X., & Li, J. (2021). *A review on clustering methods and applications*. Journal of Big Data.
- Shahapure, K. R., & Nicholas, C. (2020). *Cluster quality analysis using silhouette score*.
- Abirami, S., & Chitra, R. (2021). *A review on fuzzy c-means clustering algorithm for medical image segmentation*.
- Arora, S., Varshney, S., & Yadav, P. (2020). *Performance evaluation of K-means and fuzzy C-means clustering*.
- Shahapure, K. R., & Nicholas, C. (2020). *Cluster quality analysis using silhouette score*.
- Abirami, S., & Chitra, R. (2021). *A review on fuzzy c-means clustering algorithm for medical image segmentation*.
- Ma, H., Zhang, X., & Li, J. (2021). *A review on clustering methods and applications*. Journal of Big Data.