

ANALISIS PENERAPAN ALGORITMA FUZZY C-MEANS DAN FUZZY K-MODES UNTUK SEGMENTASI PEMUSTAKA DI PERPUSTAKAAN UNIVERSITAS MARITIM RAJA ALI HAJI

Yuni Sunarsih, Nerfita Nikentari, Nolan Efranda

Teknik Informatika, Universitas Maritim Raja Ali Haji

Jl. Raya Dompok, Kota Tanjung Pinang, Kepulauan Riau, Indonesia

yuunipulut123@gmail.com

ABSTRAK

Perpustakaan perguruan tinggi memiliki peran strategis dalam mendukung kegiatan akademik melalui penyediaan dan pengelolaan sumber informasi yang relevan. Namun, keberagaman karakteristik pemustaka sering kali menyulitkan perpustakaan dalam penambahan koleksi yang tepat sasaran. Oleh karena itu, diperlukan suatu metode segmentasi pemustaka yang mampu mengelompokkan pemustaka berdasarkan kesamaan karakteristiknya. Penelitian ini bertujuan untuk menganalisis penerapan algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes* dalam melakukan segmentasi pemustaka. Data yang digunakan adalah data peminjaman buku perpustakaan Universitas Raja Ali Haji Tahun 2024 sebanyak 7.610 data dengan variabel tipe keanggotaan, kategori buku, dan durasi. Kedua algoritma menghasilkan tiga kelompok pemustaka yaitu aktif, moderat dan kasual. Hasil pengelompokkan pemustaka aktif didominasi oleh pemustaka dari FEBM dengan minat utama pada kategori buku ilmu-ilmu terapan dengan durasi pinjam standar dimana menggunakan *Fuzzy C-Means* menunjukkan pemustaka aktif memiliki nilai derajat keanggotaan yaitu 0,88 dan segmentasi menggunakan *Fuzzy K-Modes* menghasilkan derajat keanggotaan bernilai 0,910. Segmen pemustaka moderat dengan menggunakan *Fuzzy C-Means* memiliki derajat keanggotaan sebesar 0,798 dimana tipe keanggotaan terbanyak FKIP, kategori buku karya umum, dengan durasi pinjam standar dan menggunakan *Fuzzy K-Modes* memiliki derajat keanggotaan sebesar 0,657 dengan tipe keanggotaan terbanyak dari FISIP. Sementara itu, pemustaka kasual memiliki nilai rata-rata derajat keanggotaan terendah dengan dominasi pemustaka dari FISIP dan FKIP.

Kata kunci : Segmentasi Pemustaka, Perpustakaan, Data Mining, Fuzzy C-Means, Fuzzy K-Modes

1. PENDAHULUAN

Perpustakaan perguruan tinggi memiliki peran strategis sebagai penyedia sumber informasi sekaligus pendukung utama aktivitas akademik, penelitian, dan pengembangan ilmu pengetahuan [1]. Kondisi ini juga berlaku pada Perpustakaan Universitas Maritim Raja Ali Haji (UMRAH) yang melayani pemustaka dari berbagai fakultas dengan kebutuhan akses informasi yang beragam.

Seiring perkembangan teknologi informasi, perpustakaan dituntut tidak hanya menyediakan koleksi yang memadai, tetapi juga mampu menghadirkan layanan yang relevan, adaptif, dan berbasis pemahaman menyeluruh terhadap karakteristik pemustaka. Namun, perpustakaan sering menghadapi kendala dalam memahami perilaku pemustaka secara komprehensif sehingga kebijakan layanan masih bersifat umum dan kurang responsif terhadap kebutuhan spesifik kelompok pengguna tertentu [2][3].

Salah satu pendekatan yang dapat digunakan untuk memahami karakteristik pemustaka yaitu dengan segmentasi pemustaka. Dalam konteks perpustakaan, segmentasi pemustaka dapat diartikan sebagai proses mengelompokkan pengguna perpustakaan berdasarkan kesamaan pola perilaku dalam mengakses, meminjam, dan menggunakan sumber daya perpustakaan [4]. Dengan memanfaatkan data historis peminjaman buku, segmentasi dapat memberikan gambaran nyata mengenai tingkat

keaktifan pemustaka dan kecenderungan penggunaan koleksi.

Pengolahan data peminjaman buku dalam skala besar membutuhkan pendekatan komputasional yang mampu mengekstraksi informasi, koneksi, dan pola-pola tersembunyi dalam data. *Data mining* merupakan pendekatan yang akan digunakan untuk menggali informasi berharga dari kumpulan data [5]. Salah satu penerapan data mining adalah *clustering*, yaitu proses identifikasi data ke dalam sejumlah kelompok berdasarkan atribut yang memiliki kemiripan. *Clustering* bersifat *unsupervised learning*, sehingga tidak memerlukan label kelas awal dan sangat sesuai diterapkan untuk kasus segmentasi pemustaka, di mana kategori pemustaka tidak ditentukan sebelumnya [6].

Algoritma *Fuzzy C-Means* efektif digunakan pada data numerik, sedangkan *Fuzzy K-Modes* sesuai diterapkan untuk data kategorik. Kedua algoritma ini termasuk dalam *fuzzy clustering* yang dimana memungkinkan data masuk ke dalam dua atau lebih cluster dimana titik dalam suatu cluster ditentukan dengan derajat keanggotaan [7]. Segmentasi dilakukan untuk mengelompokkan pemustaka dalam tiga kelompok yaitu pemustaka aktif, pemustaka moderat, dan pemustaka kasual dengan variabel yang digunakan yaitu tipe keanggotaan, kategori buku, dan durasi. Pemilihan ketiga variabel tersebut didasarkan pada masukan dari pustakawan di perpustakaan Universitas Raja Ali Haji.

Fokus penelitian ini lebih diarahkan pada pemanfaatan hasil segmentasi sebagai dasar untuk memahami karakteristik pemustaka, bukan pada perbandingan hasil pengelompokan menggunakan metode *Fuzzy C-Means* dan *Fuzzy K-Modes*. Hasil penelitian diharapkan dapat memberikan gambaran karakteristik pemustaka lebih jelas serta menjadi dasar rekomendasi bagi perpustakaan dalam pengembangan layanan dan pengelolaan koleksi yang lebih adaptif terhadap kebutuhan pemustaka.

2. TINJAUAN PUSTAKA

2.1. Clustering

Clustering, dapat disebut juga segmentasi (*segmentation*), merupakan metode untuk mengidentifikasi kelompok atau cluster dalam sebuah kasus didasarkan pada atribut yang memiliki kemiripan. Cara kerja clustering adalah membagi sejumlah data menjadi beberapa kelompok berdasarkan ciri yang dimiliki, di mana objek yang dimaksud bisa berupa orang, peristiwa dan lainnya yang dikelompokkan sedemikian rupa sehingga muncul beberapa tingkatan saling berhubungan antar anggota cluster [8].

Clustering merupakan salah satu metode dalam penambangan data yang bersifat *unsupervised*, karena data yang digunakan tidak memiliki label atau kategori untuk memandu proses pembelajaran. Hal ini mencakup penggunaan algoritma untuk mengidentifikasi pola dalam kumpulan data. Pola-pola ini dapat ditemukan melalui pengamatan langsung atau melalui pembuatan data simulasi. Proses pengamatan data melibatkan upaya klasifikasi variabel independen tanpa mengetahui variabel target sebelumnya [9].

2.2. Fuzzy C-Means

Fuzzy C-Means (FCM) merupakan salah satu algoritma dalam data mining clustering berbasis logika fuzzy yang cukup umum digunakan. Dasar algoritma yang diterapkan dalam FCM adalah mengelompokkan data berdasarkan seberapa besar derajat keanggotaan data terhadap cluster. Derajat keanggotaan memiliki nilai 0 hingga 1. Semakin tinggi derajat keanggotaan, maka semakin akurat pengelompokan data tersebut [10]. FCM bekerja membagi data-data dalam dataset ke dalam beragam *cluster* sehingga data yang ada tidak hanya tergabung mutlak pada satu *cluster* melainkan bergabung ke banyak *cluster*. Data yang berada di batas antara dua atau lebih *cluster* tidak dapat dianggap sepenuhnya menjadi anggota dari satu *cluster*. Sebaliknya, data tersebut dapat tergabung dalam beberapa *cluster* sekaligus dengan derajat keanggotaan sebagai indikator utama yang menunjukkan seberapa kuat data itu terpengaruh oleh karakteristik *cluster*.

Konsep utama dari *Fuzzy C-Means* adalah mengidentifikasi terlebih dahulu pusat dari *cluster* yang merepresentasikan pusat *cluster*. Pada tahap awal, pusat *cluster* masih belum akurat. Setiap data memiliki derajat keanggotaan dalam *cluster*, sehingga

harus berulang kali memperbaiki pusat *cluster* dan derajat keanggotaan setiap titik data [11].

Adapun tahapan dalam melakukan *clustering* dengan *Fuzzy C-Means* sebagai berikut [12]:

- Input data yang akan dikelompokkan, X_{ij} data sampel ke $-i$ ($i = 1, 2, \dots, n$), atribut ke $-j$ ($j = 1, 2, \dots, m$).
- Tentukan parameter yang dibutuhkan seperti banyak *cluster* (c), parameter fuzziness (m), eror terkecil yang diharapkan (ϵ), maksimum iterasi (MaxIter), fungsi objektif awal = $P_0 = 0$, dan iterasi awal = $t = 0$.
- Bangkitkan bilangan random μ_{ik} , dimana $i = 1, 2, \dots, n$, $k = 1, 2, \dots, c$; sebagai elemen matriks partisi awal U . Jumlah nilai untuk setiap kolom harus sama dengan 1.
- Hitung pusat *cluster* (*centroid*) ke $-k$ untuk atribut j , dengan $k = 1, 2, \dots, c$; dan $j = 1, 2, \dots, m$.

$$V_{kj} = \frac{\sum_{i=1}^n ((\mu_{ik})^m \cdot X_{ij})}{\sum_{i=1}^n (\mu_{ik})^m} \quad (1)$$

Keterangan :

V_{kj} = Pusat *cluster* ke- k untuk variabel ke- j

μ_{ik} = Derajat keanggotaan untuk data sampel ke- i pada *cluster* ke- k

X_{ij} = Data ke i , variabel ke j

- Hitung fungsi objektif pada iterasi ke- t , P_t

$$P_t = \sum_{i=1}^n \sum_{k=1}^c ([\sum_{j=1}^m (X_{ij} - V_{kj})^2] (\mu_{ik})^m) \quad (2)$$

Keterangan :

P_t = Fungsi Objektif

n = Jumlah data

c = Jumlah *cluster*

- Hitung perubahan matriks partisi

$$\mu_{ik} = \frac{[\sum_{j=1}^m (X_{ij} - V_{kj})^2]^{-\frac{1}{m-1}}}{\sum_{k=1}^c [\sum_{j=1}^m (X_{ij} - V_{kj})^2]^{-\frac{1}{m-1}}} \quad (3)$$

Keterangan :

μ_{ik} = derajat keanggotaan data ke i pada *cluster* ke- k

- Periksa kondisi berhenti:

- Jika : $(|P_t - P_{t-1}| < \epsilon)$ atau $(t > \text{iterasi maksimal})$ maka berhenti;
- Jika tidak: lanjutkan langkah 4 dengan $t = t + 1$.

2.3. Fuzzy K-Modes

Fuzzy K-Modes (FKMO) adalah pengembangan dari algoritma *K-Modes* yang mengintegrasikan konsep logika fuzzy. *K-Modes* sendiri merupakan modifikasi dari *K-Means* dan menjadi salah satu algoritma clustering pertama yang dikembangkan untuk analisis data kategorik skala besar. Beberapa perbedaan antara algoritma *K-Means* dan *K-Modes* adalah sebagai berikut:

- Menggunakan ukuran ketidaksamaan (*dissimilarity measure*) untuk data kategorik
- Penggunaan modus sebagai ganti rata-rata, dan
- Menggunakan metode berbasis frekuensi untuk mencari modus [13].

Fuzzy K-Modes mempunyai konsep yang mirip dengan FCM dimana setiap data tidak hanya dimiliki

satu *cluster*, tetapi dapat memiliki derajat keanggotaan terhadap beberapa *cluster* secara bersamaan. Pada metode ini, pusat *cluster* diperbarui pada setiap iterasi dengan mengukur jarak antar masing-masing pusat *cluster* dan masing-masing objek [14]. Tahapan *clustering* dengan *Fuzzy K-Modes* dilakukan dengan algoritma dasar sebagai berikut [15]:

- Input data yang akan dikelompokkan
- Tentukan parameter yang dibutuhkan seperti jumlah *cluster* (c), pangkat derajat keanggotaan (α), dan maksimum iterasi (MaxIter).
- Inisialisasi pusat *cluster* atau *centroid* awal dengan memilih modus dari data secara acak.
- Hitung jarak data poin untuk setiap *cluster* berdasarkan derajat keanggotaan *fuzzy*

$$d(X, Y) = \sum_{j=1}^m \delta(x_j, y_j) \quad (4)$$

Keterangan :

- $d(X, Y)$ = perhitungan ketidaksamaan antar dua variabel
- x_j, y_j = nilai variabel ke j dari objek x dan y
- $\delta(x_i, y_i)$ = fungsi ketidakmiripan per atribut, didefinisikan sebagai :

$$\delta(x_i, y_i) = \begin{cases} 0 & \text{jika } x_i = y_i \\ 1 & \text{jika } x_i \neq y_i \end{cases}$$

- Hitung modus setiap *cluster* dengan derajat keanggotaan *fuzzy*.

$$u_{ij} = \frac{1}{\sum_{h=1}^k \left(\frac{d(x_i, z_h)}{d(x_i, z_j)} \right)^{\frac{1}{\alpha-1}}} \quad (5)$$

Keterangan :

- u_{ij} = derajat keanggotaan *fuzzy* data ke- i ke *cluster* ke- j
 - x_i, z_i = jarak antara data ke- i ke himpunan pusat *cluster* ke- i
- Update *centroid* dengan mengambil semua kemungkinan nilai kategori setiap variabel. Kategori dengan jumlah nilai anggota terbesar akan menjadi *centroid* baru variabel X . Teoremanya sebagai berikut.

$$r = \arg \max \sum_{i, x_{ij}=a_{jr}} w_{il}^{\alpha} \quad (6)$$

Dimana

$$\sum_{i, x_{ij}=a_{jr}^{(r)}} w_{il}^{\alpha} \geq \sum_{i, x_{ij}=a_{jt}^{(t)}} w_{il}^{\alpha}$$

Untuk $1 \leq j \leq d$ dan $1 \leq t \leq k$

Keterangan:

- w_{il} = derajat keanggotaan data ke *cluster* i
 - $x_{i,j}$ = nilai variabel ke- j dari data ke- i
 - $a_j^{(r)}$ = salah satu kategori dari variabel ke- j
- Iterasi dilanjutkan hingga pusat *cluster* sudah stabil atau maksimum iterasi sudah tercapai.

2.4. Segmentasi Pemustaka

Segmentasi adalah upaya mengelompokkan populasi atau kumpulan data menjadi beberapa kelompok (segmen) yang memiliki karakteristik serupa. Segmentasi merupakan tindakan membagi pasar yang beragam menjadi bagian-bagian yang memiliki kesamaan. Segmen-segmen ini dapat memiliki kesamaan dalam hal kebutuhan, keinginan,

tingkah laku, atau tanggapan terhadap pemasaran yang spesifik [16].

Jika penjelasan ini diadaptasi ke dalam konteks perpustakaan, maka segmentasi pemustaka merupakan proses pengelompokan pemustaka berdasarkan pola perilaku dan karakteristik dalam memanfaatkan layanan perpustakaan. Tujuannya adalah memahami kebutuhan pemustaka lebih mendalam sehingga perpustakaan dapat menyediakan layanan tepat sasaran serta meningkatkan kepuasan pemustaka.

Pemustaka dalam penelitian ini adalah seluruh pengguna layanan perpustakaan UMRAH, baik untuk kepentingan akademik, penelitian, ataupun sekedar rekreasi. Pada penelitian ini, terdapat 3 kelompok pemustaka yang akan disegmentasi sebagai berikut:

- Pemustaka Aktif, merupakan kelompok pengguna dengan pola peminjaman konsisten. Koleksi yang dipinjam umumnya bersifat akademik atau ilmiah, yang digunakan untuk menunjang kegiatan perkuliahan dan penelitian. Derajat keanggotaan *fuzzy* kelompok pemustaka aktif paling tinggi dan mendekati 1, yang menunjukkan bahwa perilaku pemustaka berada sangat dekat dengan pusat *cluster* dan memiliki tingkat konsistensi kuat.
- Pemustaka Moderat, merupakan kelompok pengguna yang tidak menunjukkan kecenderungan spesifik terhadap satu jenis koleksi tertentu, melainkan memanfaatkan berbagai koleksi buku.
- Pemustaka Kasual, merupakan kelompok pengguna yang memiliki tingkat peminjaman relatif rendah atau tidak secara rutin. Koleksi yang dipinjam bersifat umum atau non-spesifik. Derajat keanggotaan *fuzzy* pada jenis pemustaka ini paling jauh dari nilai 1, yang mencerminkan penggunaan layanan perpustakaan tidak konsisten.

3. METODE PENELITIAN

3.1. Teknik Pengumpulan Data

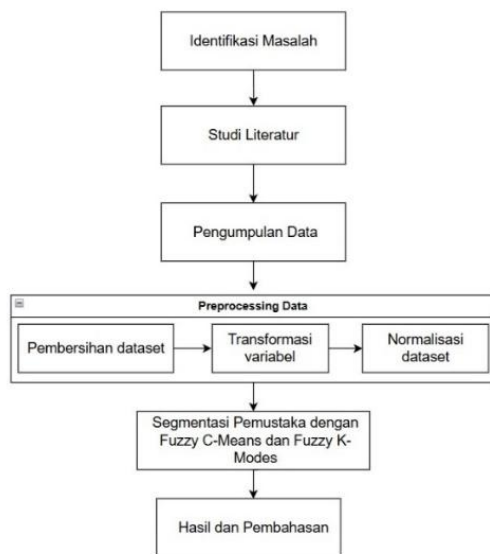
Teknik pengumpulan data pada penelitian “Analisis Penerapan Algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes* Untuk Segmentasi Pemustaka di Perpustakaan Universitas Maritim Raja Ali Haji” melibatkan beberapa metode sebagai berikut:

- Dokumentasi
Data yang digunakan dalam penelitian ini berupa data historis peminjaman buku yang tercatat pada sistem informasi perpustakaan Universitas Maritim Raja Ali Haji.
- Studi Pustaka
Studi Pustaka yaitu mempelajari teori-teori baik dari Jurnal, buku cetak maupun non cetak yang berhubungan dengan objek penelitian.

3.2. Prosedur Pelaksanaan Penelitian

Penelitian ini bertujuan menganalisis penerapan algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes* untuk segmentasi pemustaka di Perpustakaan Universitas

Maritim Raja Ali Haji (UMRAH). Prosedur pelaksanaan penelitian dapat dilihat alurnya pada Gambar 1.

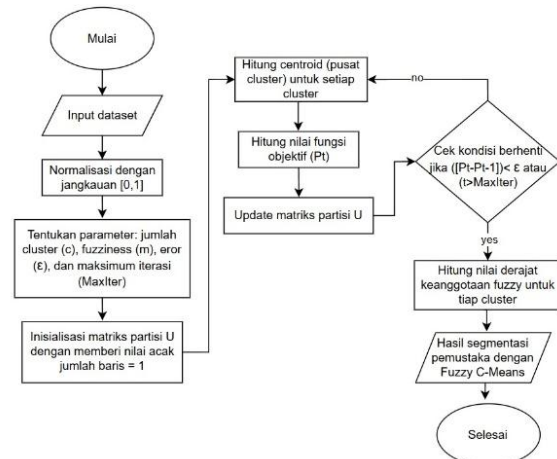


Gambar 1. Tahapan Penelitian

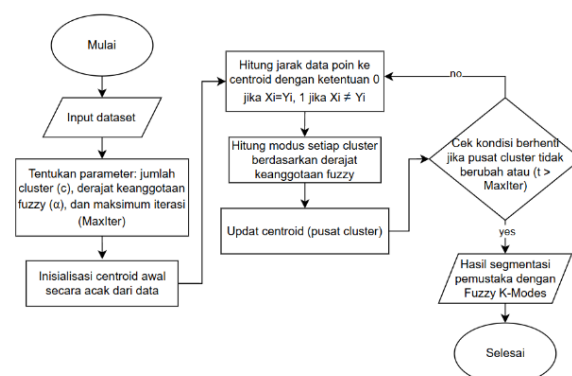
Penjelasan *flowchart* tahapan penelitian diatas dapat diuraikan sebagai berikut:

- Identifikasi masalah, Langkah pertama dalam melakukan penelitian ini penulis memahami terlebih dahulu masalah yang akan dikaji,
- Studi Literatur, penulis mengumpulkan informasi sebanyak mungkin mengenai penelitian yang pernah dilakukan sebelumnya beserta landasan teori relevan dengan masalah yang akan diteliti. Penelitian berfokus pada segmentasi pemustaka dengan Algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes*,
- Pengumpulan Data, pada tahap ini data peminjaman buku dikumpulkan. Data ini berupa daftar historis peminjaman buku perpustakaan UMRAH tahun 2024,
- Preprocessing* data, pada tahap ini data yang sudah dikumpulkan akan di bersihkan dan di transformasi.
- Clustering* dengan algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes*, pada tahap ini dilakukan pengelompokan data dengan kedua metode menggunakan data yang telah di praproses sebelumnya,
- Hasil dan Pembahasan, Tahapan terakhir yaitu menganalisis hasil segmentasi dan dikaitkan berdasarkan tujuan penelitian.

Pemrosesan data dilakukan dengan algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes* yang dihitung melalui sistem untuk segmentasi pemustaka UMRAH. Adapun tahapan dalam algoritma *Fuzzy C-Means* dan *Fuzzy K-Modes* dapat dilihat dalam flowchart berikut.



Gambar 2. Flowchart algoritma *Fuzzy C-Means*



Gambar 3. Flowchart algoritma *Fuzzy K-Modes*

4. HASIL DAN PEMBAHASAN

Hasil penelitian yang dijelaskan pada pembahasan ini menjelaskan tentang proses clustering dataset menggunakan *Fuzzy C-Means* dan *Fuzzy K-Modes* segmentasi pemustaka di perpustakaan UMRAH.

4.1. Data

Data yang digunakan dalam penelitian ini berjumlah sebanyak 7.610 data yang telah melalui proses pembersihan dataset. Dataset tersebut bersumber Perpustakaan UMRAH yang beralamat di kampus UMRAH Dompok. Dataset dapat dilihat pada gambar berikut.

ID	Tipe Keanggotaan	Judul	Kategori Buku	Tanggal Pinjam	Tanggal Kembali
0	FKIP	Harga sebuah percaya	Kesusastraan (800)	31/01/2024	07/02/2024
1	FKIP	Gadis kretek	Kesusastraan (800)	31/01/2024	07/02/2024
2	FEBM	Manajemen Sumber Daya	Ilmu-ilmu Terapan (600)	30/01/2024	06/02/2024
3	FEBM	Manajemen sumber daya	Ilmu-ilmu Terapan (600)	30/01/2024	06/02/2024
4	FISIP	Metodologi penelitian ki	Ilmu-ilmu Terapan (600)	30/01/2024	06/02/2024
...	...	...	...	...	...
7605	FEBM	Manajemen keuanganpe	Ilmu-ilmu Terapan (600)	02/12/2024	09/12/2024
7606	FEBM	Manajemen keuanganpe	Ilmu-ilmu Terapan (600)	02/12/2024	09/12/2024
7607	FEBM	AKUNTANSI BIAYA EDISI	Ilmu-ilmu Terapan (600)	02/12/2024	09/12/2024
7608	FISIP	Methods in Behavioral R	Karya Umum (000)	02/12/2024	09/12/2024
7609	FISIP	Methods in Behavioral R	Karya Umum (000)	02/12/2024	09/12/2024

Gambar 4. Dataset peminjaman buku perpustakaan

Pada penelitian ini variabel yang digunakan adalah tipe keanggotaan, kategori buku, dan durasi peminjaman.

4.2. Proses Fuzzy C-Means

Sebelum dilakukan segmentasi, terlebih dahulu dilakukan transformasi data dengan insialisasi variabel kategorik yaitu tipe keanggotaan, kategori buku, dan durasi. Penentuan nilai transformasi variabel diurutkan berdasarkan frekuensi pada data peminjaman, frekuensi data tinggi memiliki nilai variabel yang besar pula [17]. Transformasi data tipe keanggotaan dapat dilihat pada tabel 1, transformasi variabel kategori buku pada tabel 2 dan tabel 3 menampilkan transformasi variabel durasi.

Tabel 1. Transformasi variabel tipe keanggotaan

No	Variabel	Nilai
1	STAF	1
2	FIKP	2
3	FTTK	3
4	DOSEN	4
5	FISIP	5
6	FKIP	6
7	FEBM	7

Tabel 2. Transformasi variabel kategori buku

No	Variabel Kategori Buku	Nilai
1	Kesenian, Hiburan dan Olahraga (700)	1
2	Geografi dan Sejarah (900)	2
3	Agama (200)	3
4	Filsafat (100)	4
5	Ilmu-ilmu Murni (500)	5
6	Karya Umum (000)	6
7	Bahasa (400)	7
8	Kesusastraan (800)	8
9	Ilmu-ilmu Terapan (600)	9
10	Ilmu-ilmu Sosial (300)	10

Tabel 3. Transformasi variabel durasi

No	Variabel Durasi	Nilai
1	0-6 Hari	1
2	7-14 hari	2
3	>14 hari	3

Langkah awal proses perhitungan dengan *Fuzzy C-Means* yaitu menormalisasi dataset untuk menghindari dominasi salah satu variabel dalam proses clustering. Selanjutnya menentukan parameter untuk menyelesaikan permasalahan. Parameter-parameter tersebut adalah banyak cluster ($c = 3$), parameter fuzzy ($m = 2$), eror terkecil yang diharapkan ($\epsilon = 0,01$), dan iterasi maksimum ($\text{MaxIter} = 100$).

Pada langkah selanjutnya perhitungan menggunakan sistem menghasilkan output sebagai berikut: Hasil segmentasi menggunakan *Fuzzy C-Means* dapat dilihat pada table 4.

Tabel 4. Jumlah data segmentasi *Fuzzy C-Means*

Cluster	Jumlah Data	Membership
C 1	2001 Data	0,798
C 2	3578 Data	0,880
C 3	2031 Data	0,765

Berdasarkan hasil segmentasi menggunakan algoritma *Fuzzy C-Means*, diperoleh tiga segmen

utama pemustaka. Pada segmen pemustaka aktif yaitu berada di cluster 2 dengan rata-rata derajat keanggotaan tertinggi dari tiga cluster sebesar 0,880. data yang mendominasi pemustaka aktif adalah tipe keanggotaan FEBM, kategori buku Ilmu-ilmu terapan, dengan durasi pinjam standar. Nilai derajat keanggotaan tersebut mengindikasikan bahwa karakteristik pemustaka aktif memiliki pola peminjaman yang konsisten. Segmen pemustaka moderat berada di cluster 1 dengan rata-rata derajat keanggotaan sebesar 0,798. Pada segmen pemustaka moderat, tipe keanggotaan terbanyak FKIP, kategori buku karya umum, dengan durasi pinjam standar. Hal ini menunjukkan variasi karakteristik lebih besar dibanding pemustaka aktif. Sementara itu segmen pemustaka kasual berada di cluster 3 dengan rata-rata derajat keanggotaan sebesar 0,765. yang merupakan nilai terendah diantara ketiga segmen *Fuzzy C-Means*. Tipe keanggotaan terbanyak FISIP, kategori buku ilmu-ilmu sosial, dengan durasi pinjam standar.

4.3. Proses Fuzzy K-Modes

Proses segmentasi pemustaka dengan *Fuzzy K-Modes* menggunakan data kategorik sehingga terlebih dahulu dilakukan transformasi data untuk variabel durasi menjadi bentuk kategorik. tabel 5 menampilkan transformasi variabel durasi kategorik.

Tabel 5. Transformasi variabel durasi kategorik

No	Variabel Durasi	Nilai
1	0-6 Hari	Pendek
2	7-14 hari	Standar
3	>14 hari	Panjang

Langkah awal proses perhitungan dengan *Fuzzy K-Modes* yaitu menentukan parameter yang dibutuhkan dalam menyelesaikan permasalahan. Pada penelitian ini ditentukan jumlah cluster ($c = 3$), pangkat derajat keanggotaan ($\alpha = 2$), dan iterasi maksimal ($\text{MaxIter} = 100$).

Pada langkah selanjutnya perhitungan menggunakan sistem menghasilkan output yang dapat dilihat pada tabel 6 sebagai berikut:

Tabel 6. Jumlah data segmentasi *Fuzzy K-Modes*

Cluster	Jumlah Data	Membership
C 1	3655 Data	0,567
C 2	1917 Data	0,556
C 3	2038 Data	0,910

Berdasarkan hasil segmentasi menggunakan algoritma *Fuzzy K-Modes* juga menghasilkan tiga segmen pemustaka. Pada segmen pemustaka aktif berada di cluster 3 dengan rata-rata derajat keanggotaan paling besar yaitu 0,910 yang menunjukkan tingkat kedekatan sangat tinggi antara data dengan pusat cluster. Data yang mendominasi pemustaka aktif adalah tipe keanggotaan FEBM, kategori buku ilmu-ilmu terapan, dengan durasi pinjam standar. Segmen pemustaka moderat berada di

cluster 1 dengan rata-rata derajat keanggotaan 0,567, menunjukkan adanya variasi karakteristik pemustaka yang tumpang tindih dengan segmen lain. Tipe keanggotaan terbanyak FISIP, kategori buku ilmu-ilmu sosial, dengan durasi pinjam standar. Sementara itu pada segmen pemustaka kasual berada di cluster 2 memiliki rata-rata derajat keanggotaan 0,556, yang merupakan nilai terendah pada hasil *Fuzzy K-Modes*. Nilai tersebut menunjukkan bahwa pemustaka kasual memiliki pola perilaku yang paling beragam dan tingkat keterikatan yang relative rendah terhadap satu cluster tertentu. Tipe keanggotaan terbanyak FKIP, kategori buku ilmu-ilmu murni, dengan durasi pinjam standar.

4.4. Analisa

Berdasarkan hasil segmentasi menggunakan algoritma *Fuzzy C-Means*, diperoleh tiga segmen utama pemustaka. Pada segmen pemustaka aktif yaitu berada di *cluster 2* dengan rata-rata derajat keanggotaan tertinggi dari tiga *cluster* sebesar 0,880. Nilai derajat keanggotaan tersebut mengindikasikan bahwa karakteristik pemustaka aktif memiliki pola peminjaman yang konsisten. Segmen pemustaka moderat berada di *cluster 1* dengan rata-rata derajat keanggotaan sebesar 0,798, yang menunjukkan variasi karakteristik lebih besar dibanding pemustaka aktif. Sementara itu segmen pemustaka kasual berada di *cluster 3* dengan rata-rata derajat keanggotaan sebesar 0,765, yang merupakan nilai terendah diantara ketiga segmen *Fuzzy C-Means*.

Berdasarkan hasil segmentasi menggunakan algoritma *Fuzzy K-Modes* juga menghasilkan tiga segmen pemustaka. Pada segmen pemustaka aktif berada di *cluster 3* dengan rata-rata derajat keanggotaan paling besar yaitu 0,910 yang menunjukkan tingkat kedekatan sangat tinggi antara data dengan pusat *cluster*. Segmen pemustaka moderat berada di *cluster 1* dengan rata-rata derajat keanggotaan 0,567, menunjukkan adanya variasi karakteristik pemustaka yang tumpang tindih dengan segmen lain. Sementara itu pada segmen pemustaka kasual berada di *cluster 2* memiliki rata-rata derajat keanggotaan 0,556, yang merupakan nilai terendah pada hasil *Fuzzy K-Modes*. Nilai tersebut menunjukkan bahwa pemustaka kasual memiliki pola perilaku yang paling beragam dan tingkat keterikatan yang relative rendah terhadap satu *cluster* tertentu.

Perbedaan hasil segmentasi yang diperoleh dari kedua algoritma merupakan konsekuensi logis dari perbedaan prinsip dasar kedua algoritma dalam memproses data. *Fuzzy C-Means* melakukan pengelompokan berdasarkan jarak antar data numerik dan pusat *cluster*. Sebaliknya, *Fuzzy K-Modes* dirancang khusus untuk data kategorik dan mengukur kemiripan berdasarkan kesamaan atribut non-numerik.

Perbedaan pendekatan ini menyebabkan masing-masing algoritma menonjolkan karakteristik pemustaka yang berbeda, sehingga distribusi segmen yang dihasilkan tidak sepenuhnya identik. Penggunaan

kedua algoritma dalam penelitian ini menjadi relevan karena data peminjaman buku memiliki karakteristik campuran, yang tidak sepenuhnya dapat direpresentasikan secara optimal oleh satu pendekatan *clustering* saja. Dengan mengombinasikan hasil dari kedua algoritma, segmentasi pemustaka yang dihasilkan menjadi lebih komprehensif dan interpretatif, tanpa bertujuan untuk menentukan algoritma yang paling unggul.

Pendekatan ini memungkinkan pihak perpustakaan untuk melihat perilaku pemustaka dari dua sudut pandang yang saling melengkapi, yaitu kuantitatif dan kategorik. Oleh karena itu, perbedaan hasil segmentasi yang muncul tidak dipandang sebagai inkonsistensi, melainkan sebagai informasi tambahan yang memperkaya pemahaman terhadap karakteristik pemustaka.

5. KESIMPULAN DAN SARAN

Fuzzy C-Means menghasilkan tiga segmen pemustaka. Pemustaka aktif berada pada cluster 2 dengan derajat keanggotaan tertinggi (0,880), didominasi FEBM, kategori ilmu-ilmu terapan, dan durasi pinjam standar, yang menunjukkan pola konsisten. Pemustaka moderat berada pada cluster 1 (0,798), didominasi FKIP dan karya umum, dengan variasi karakteristik lebih besar. Pemustaka kasual berada pada cluster 3 (0,765), didominasi FISIP dan ilmu-ilmu sosial, dengan pola lebih beragam.

Fuzzy K-Modes juga menghasilkan tiga segmen. Pemustaka aktif berada pada cluster 3 dengan nilai tertinggi (0,910), didominasi FEBM dan ilmu-ilmu terapan. Pemustaka moderat berada pada cluster 1 (0,567), didominasi FISIP dan ilmu-ilmu sosial, menunjukkan tumpang tindih karakteristik. Pemustaka kasual berada pada cluster 2 (0,556), didominasi FKIP dan ilmu-ilmu murni, dengan perilaku paling beragam dan keterikatan rendah.

Perbedaan hasil disebabkan oleh pendekatan algoritma: *Fuzzy C-Means* berbasis jarak numerik, sedangkan *Fuzzy K-Modes* berbasis kesamaan kategorik. Keduanya saling melengkapi dan memberikan pemahaman komprehensif terhadap karakteristik pemustaka.

Saran untuk peneliti selanjutnya dapat menambahkan lebih banyak variabel ke dalam analisis untuk memperoleh pemahaman yang lebih mendalam tentang perilaku setiap kelompok pemustaka serta mendapatkan hasil segmentasi yang lebih spesifik.

DAFTAR PUSTAKA

- [1] I. Sopwandin, *Manajemen Perpustakaan Perguruan Tinggi*. Bogor: Guepedia Group, 2021.
- [2] I. S. Ritonga, "Analisis data peminjaman perpustakaan untuk meningkatkan layanan dan efisiensi pengelolaan UPT Perpustakaan UIN Syekh Ali Hasan Ahmad Addary Padangsidempuan," *Al-Kuttab J. Kaji. Perpustakaan, Inf. dan Kearsipan*, vol. 6, no. 1,

- pp. 41–51, 2024.
- [3] S. Dewi, A. Kresnawati, S. Sopyanti, and A. Sulmainah, “Mapping Library User Behavior Base On K-Means Clustering Of Ma’soem University Student Pemetaan Perilaku Pengguna Perpustakaan Berbasis K-Means Pada Mahasiswa Universitas Ma’soem,” *Indones. J. Inform. Res. Softw. Eng.*, vol. 5, no. 2, pp. 130–136, 2025.
 - [4] A. Zohriah, A. Qurtubi, and A. Fatoni, “Segmentasi Pasar Jasa Pendidikan Menuju Era Society 5.0,” *J. Ilm. Wahana Pendidik.*, vol. 9, no. 10, pp. 66–79, 2023.
 - [5] P. W. Rahayu, *Buku Ajar Data Mining*. Jambi: PT. Sonpedia Publishing Indonesia, 2024.
 - [6] J. Fadlina, R. Utami, and N. M. Sitinjak, “Pengelompokan Buku dan Rekomendasi Buku Menggunakan K-Means Clustering Pada Dinas Perpustakaan dan Kearsipan Kota Medan,” *J. Widya*, vol. 5, no. 2, pp. 1045–1058, 2024.
 - [7] A. Nurzida, I. T. Utami, and M. Rochayani, “Perbandingan Metode Fuzzy C-Means Dan Gustafson-Kessel Dalam Penentuan Cluster Tingkat Risiko Penularan Tuberculosis Terhadap Penyakit di Jawa Timur,” *J. Gaussian*, vol. 13, no. 2, pp. 373–382, 2024.
 - [8] J. Nasir, “Penerapan Data Mining Clustering dalam Mengelompokkan Buku dengan Metode K-Means,” *SIMETRIS*, vol. 11, no. 2, pp. 690–703, 2021.
 - [9] T. Alasali and Y. Ortakci, “Clustering Techniques in Data Mining: A Survey of Methods, Challenges, and Applications,” *J. Comput. Sci.*, vol. 9, no. 1, pp. 32–50, 2024.
 - [10] G. . Nugraha, R. Dwiyanaputra, F. Bimantoro, and A. Aranta, “Implementasi Fuzzy C-Means untuk Pengelompokan Daerah Berdasarkan Penularan COVID-19,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 1, pp. 97–10, 2023.
 - [11] W. Monika, A. H. Nasution, F. A. Syam, and C. Wijesundara, “Clustering Of Library’s Patron Behavior Using Machine Learning,” *Digital Zone: Jurnal Teknologi Informasi dan Komunikasi*, vol. 16, no. 1, pp. 1–11, 2025.
 - [12] B. Destia and M. . Kartikasari, “Comparison Of Fuzzy C-Means And Fuzzy Gustafson-Kessel Clustering Methods In Provincial Grouping In Indonesia Based On Criminality-Related Factors,” *Jounal Math. Its Appl.*, vol. 17, no. 2, pp. 1093–1102, 2023.
 - [13] R. J. Kuo, Y. R. Zheng, and T. P. Q. Nguyen, “Metaheuristic-Based Possibilistic Fuzzy K-Modes Algorithms for Categorical Data Clustering,” *Inf. Sci. (Ny).*, vol. 557, no. 1, pp. 1–15, 2021.
 - [14] L. Qadrini, *Algoritma Pengelompokan Ensemble Fuzzy, K-Prototypes dan Density Peaks Clustering Mixed (DPC-M)*. Tasikmalaya: Perkumpulan Rumah Cemerlang Indonesia, 2021.
 - [15] P. D’Urso, L. De Giovanni, L. Federico, and V. Vitale, “Fuzzy C-modes clustering with spatial regularization and noise cluster,” *ASIA Adv. Stat. Anal.*, 2025.
 - [16] Hartini, A. Sudirman, and C. . Wardhana, *Manajemen Pemasaran (Era Revolusi industri 4.0)*. Bandung: CV. Media Sains Indonesia, 2022.
 - [17] Y. Syahra, A. Fadlil, and H. Yuliansyah, “Customer Segmentation Using RFM and K-Means Clustering to Support CRM in Retail Industry,” *Sinkron*, vol. 9, no. 3, pp. 1120–1131, Jul. 2025.