

Penerapan Algoritma *K-Means* untuk Pengelompokan Koleksi Perpustakaan dengan Data Mining

Ahmad Fairus Zabidi

Program Studi Teknik Informatika, Fakultas Teknik, Universitas Mercu Buana
ahmadfairuszabidi@email.com

Abstract

This study aims to apply the Clustering data mining technique using the K-Means algorithm to library collections to improve information management and accessibility. By utilizing the type of borrowed collection data, this study analyzes the usage patterns and preferences of library visitors. The research method used is a quantitative approach with technical implementation steps, to ensure that the implementation of K-Means can run systematically and provide useful results for library management. Data collection methods used include direct observation of user interactions, interviews with library staff and users, and electronic data collection from the library management system to obtain information about collection usage. The results of the application of the K-Means algorithm show that library collections can be grouped based on similar characteristics, which allows library managers to better understand user needs and optimize existing collections. This research is expected to provide significant contributions such as assisting managers in managing collections or making data-based recommendations in developing library management systems, as well as improving user experience in accessing relevant information.

Keywords: *Clustering, Data Mining, Library Collections, K-Means, Information Management.*

Abstrak

Penelitian ini bertujuan untuk menerapkan teknik data mining Clustering menggunakan algoritma K-Means pada koleksi perpustakaan untuk meningkatkan pengelolaan dan aksesibilitas informasi. Dengan memanfaatkan jenis data koleksi yang dipinjam, penelitian ini melakukan analisis terhadap pola penggunaan dan preferensi pengunjung perpustakaan. Metode penelitian yang digunakan adalah pendekatan kuantitatif dengan langkah-langkah implementasi teknis, untuk memastikan implementasi K-Means dapat berjalan secara sistematis dan memberikan hasil yang bermanfaat bagi pengelolaan perpustakaan. Metode pengumpulan data yang digunakan meliputi observasi langsung terhadap interaksi pengguna, wawancara dengan staf dan pengguna perpustakaan, serta pengumpulan data elektronik dari sistem manajemen perpustakaan guna mendapatkan informasi seputar penggunaan koleksi. Hasil dari penerapan algoritma K-Means menunjukkan bahwa koleksi perpustakaan dapat dikelompokkan berdasarkan kesamaan karakteristik, yang memungkinkan pengelola perpustakaan untuk lebih memahami kebutuhan pengguna dan mengoptimalkan koleksi yang ada. Penelitian ini diharapkan dapat memberikan kontribusi signifikan seperti membantu pengelola dalam mengelola koleksi atau membuat rekomendasi berbasis data dalam pengembangan sistem pengelolaan perpustakaan, serta meningkatkan pengalaman pengguna dalam mengakses informasi yang relevan.

Kata Kunci: *Clustering, Data Mining, Koleksi Perpustakaan, K-Means, Pengelolaan Informasi.*

I. PENDAHULUAN

Perpustakaan sebagai pusat informasi memegang peranan yang sangat penting dalam mendukung proses pendidikan, penelitian, dan pengembangan pengetahuan di berbagai bidang. Selain berfungsi sebagai tempat penyimpanan informasi, perpustakaan juga berperan sebagai fasilitator bagi para pengguna dalam mengakses berbagai sumber daya yang dibutuhkan untuk mendalami berbagai topik. Dengan semakin berkembangnya teknologi informasi, perpustakaan modern kini dihadapkan pada tantangan besar dalam mengelola koleksi yang semakin besar dan beragam, baik dalam format cetak maupun digital. Salah satu tantangan utama

yang dihadapi adalah pengelolaan koleksi yang terus berkembang dengan jumlah item yang sangat banyak, sehingga menyulitkan pengelompokan informasi yang relevan untuk tiap pengguna. Selain itu, kebutuhan untuk meningkatkan pengalaman pengguna, mengoptimalkan pencarian informasi, serta menyediakan layanan yang lebih personal menjadi tantangan lain yang memerlukan perhatian serius.

Untuk menghadapi tantangan tersebut, perpustakaan perlu beradaptasi dengan perkembangan teknologi, salah satunya dengan menerapkan teknik data mining. Dalam konteks ini, teknik *Clustering* menjadi relevan karena memungkinkan perpustakaan untuk mengelompokkan koleksi berdasarkan pola dan kesamaan yang ada, sehingga dapat mempermudah pencarian dan

pengorganisasian koleksi. *Clustering* adalah proses yang digunakan untuk mengelompokkan data ke dalam grup-grup yang memiliki kesamaan tertentu. Dalam hal ini, teknik *Clustering* dapat digunakan untuk mengelompokkan koleksi berdasarkan berbagai atribut, seperti jenis bahan pustaka, topik, frekuensi peminjaman, atau minat pengguna [1]. Salah satu algoritma *Clustering* yang banyak digunakan dalam konteks ini adalah *K-Means*, yang dikenal karena kesederhanaan dan efisiensinya dalam menangani data besar [2]. Algoritma *K-Means* bekerja dengan cara mengarahkan data untuk dikelompokkan ke dalam sejumlah titik pusat atau centroid yang berfungsi sebagai acuan dalam pengelompokan, sehingga data yang serupa dapat dikelompokkan dalam grup yang sama [3].

Namun, meskipun algoritma *K-Means* dikenal efisien, algoritma ini memiliki beberapa keterbatasan yang perlu diperhatikan. Salah satunya adalah bahwa *K-Means* mengharuskan pengguna untuk menentukan jumlah kluster (*K*) terlebih dahulu, yang bisa menjadi masalah ketika jumlah kluster yang optimal tidak jelas atau sulit diprediksi. Selain itu, *K-Means* sangat sensitif terhadap outlier, yang dapat mempengaruhi hasil pengelompokan secara signifikan. Kelemahan lainnya adalah *K-Means* bekerja dengan asumsi bahwa data yang dikelompokkan memiliki bentuk yang relatif bulat dan simetris, sehingga kurang efektif untuk menangani data yang lebih kompleks atau memiliki distribusi yang tidak teratur [4]. Oleh karena itu, pemilihan algoritma yang tepat serta evaluasi hasil klusterisasi menjadi faktor penting dalam memastikan bahwa pengelompokkan koleksi perpustakaan dapat memberikan manfaat maksimal.

Penerapan algoritma *K-Means* dalam konteks pengelolaan koleksi perpustakaan dapat memberikan berbagai manfaat. Sebagai contoh, penelitian sebelumnya menunjukkan bahwa analisis *Clustering* dapat digunakan untuk memahami pola perilaku pengguna perpustakaan dan mengidentifikasi sumber daya yang paling sering digunakan, yang kemudian bisa menjadi dasar untuk pengambilan keputusan dalam mengoptimalkan koleksi [5]. Di sisi lain, penelitian lain juga menunjukkan bahwa analisis *Clustering* dapat membantu perpustakaan dalam menyediakan layanan yang lebih sesuai dengan kebutuhan pengguna, seperti mengelompokkan koleksi berdasarkan minat atau topik yang sering dicari, sehingga memudahkan pengguna dalam menemukan informasi yang relevan [6]. Dengan demikian, penerapan *K-Means* dapat meningkatkan pengalaman pengguna dan memudahkan pengelola perpustakaan dalam merespons kebutuhan informasi yang lebih spesifik.

Penelitian ini bertujuan untuk menerapkan algoritma *K-Means* dalam mengelompokkan koleksi perpustakaan dengan tujuan yang lebih spesifik, yaitu untuk mengeksplorasi dan menganalisis data koleksi yang ada guna meningkatkan efektivitas pengelolaan koleksi perpustakaan. Selain itu, penelitian ini juga bertujuan untuk mengidentifikasi pola penggunaan koleksi serta preferensi pengunjung, sehingga dapat memberikan wawasan yang lebih dalam tentang bagaimana pengelolaan koleksi dapat dioptimalkan. [7] Lebih lanjut, penelitian ini akan menganalisis apakah metode *Clustering* dengan *K-Means* dapat memberikan gambaran yang lebih jelas mengenai sumber daya yang paling

banyak diakses dan bagaimana hal tersebut berhubungan dengan kebutuhan pengguna. Hasil dari penelitian ini diharapkan dapat memberikan solusi yang lebih efisien dalam mengelola koleksi perpustakaan, yang pada gilirannya akan meningkatkan layanan perpustakaan secara keseluruhan.

Kontribusi penelitian ini terletak pada penerapan metode data mining secara konkret dalam pengelolaan koleksi perpustakaan. Dengan menggunakan algoritma *K-Means* untuk mengelompokkan koleksi berdasarkan berbagai atribut yang relevan, penelitian ini akan memberikan panduan bagi pengelola perpustakaan dalam mengidentifikasi pola dan preferensi pengguna. Selain itu, temuan dari penelitian ini diharapkan dapat memberikan dasar untuk perencanaan koleksi yang lebih baik, yang mengarah pada penghematan sumber daya dan peningkatan kualitas layanan. Secara konkret, penelitian ini akan menghasilkan sebuah model pengelompokan koleksi yang dapat digunakan oleh perpustakaan lain untuk mengelola koleksi mereka dengan lebih efektif. Dengan demikian, penelitian ini tidak hanya berkontribusi pada pengembangan ilmu perpustakaan, tetapi juga memberikan solusi praktis yang dapat diterapkan untuk menghadapi tantangan modern dalam pengelolaan informasi dan koleksi perpustakaan. Selain itu, temuan-temuan ini dapat menjadi referensi untuk penelitian selanjutnya yang berfokus pada penerapan metode data mining dalam berbagai aspek pengelolaan perpustakaan dan layanan informasi [8][9].

Dengan demikian, penelitian ini memiliki potensi untuk memberikan wawasan yang signifikan dalam memahami penerapan algoritma *K-Means* dalam pengelolaan koleksi perpustakaan dan memberikan kontribusi yang lebih luas dalam bidang data mining serta pengelolaan informasi perpustakaan [10][11].

II. METODE PENELITIAN

Studi ini mengadopsi pendekatan kuantitatif yang bertujuan untuk menganalisis penerapan pengelompokan data mining pada koleksi perpustakaan dengan memanfaatkan algoritma *K-Means*. Metode kuantitatif dipilih karena memungkinkan pengumpulan dan analisis data numerik yang dapat diukur dan dianalisis secara statistik.

a. Proses Pengumpulan Data

Pengumpulan data dilakukan dengan beberapa teknik yang saling melengkapi, yaitu observasi langsung, wawancara, studi dokumen, dan pengumpulan data elektronik. Berikut adalah rincian dari masing-masing teknik pengumpulan data:

Observasi dilakukan untuk memperoleh pemahaman mendalam tentang proses pengelolaan dan penggunaan koleksi perpustakaan. Peneliti melakukan observasi selama dua minggu untuk mempelajari cara koleksi diorganisasi dan diakses oleh pengguna. Observasi ini mencakup pengamatan terhadap interaksi antara pengunjung dan staf, serta pengamatan terhadap proses pencarian informasi oleh pengguna. Melalui pengamatan ini, peneliti dapat mengidentifikasi pola penggunaan koleksi yang relevan untuk dianalisis lebih lanjut menggunakan teknik pengelompokan. Observasi ini juga membantu dalam memperoleh data kontekstual yang

mendalam mengenai kebutuhan pengguna yang kemudian dapat menjadi dasar dalam penerapan algoritma *K-Means* [12][13].



Gambar 1 : Koleksi Buku dan Ruang Redaksi

Wawancara dilakukan dengan 15 staf perpustakaan dan 20 pengguna untuk menggali pengalaman mereka terkait penggunaan koleksi. Pertanyaan wawancara dirancang untuk mengidentifikasi jenis koleksi yang paling sering digunakan, kebutuhan informasi pengguna, serta tantangan dalam pengelolaan koleksi. Data wawancara memberikan perspektif kontekstual yang lebih kaya, yang tidak selalu terlihat melalui observasi langsung atau dokumen. Hasil wawancara juga membantu dalam menilai relevansi kategori pengelompokan yang dihasilkan oleh *K-Means* dalam konteks kebutuhan pengguna perpustakaan [14][15].

Tabel 1 : Notulensi Wawancara Dengan Pengembang Perangkat Lunak

Wawancara dengan Bapak Farid (Pengembang Perangkat Lunak)	
Mohon izin bertanya Bapak, Saya mempunyai ide terkait dengan penerapan data mining <i>Clustering</i> data perpustakaan menggunakan algoritma <i>K-Means</i> di Perpustakaan Lemhannas RI, bagaimana tanggapan Bapak, terkait ide tersebut?	Sebagai pengembang perangkat lunak di perpustakaan Lemhannas RI. Saya setuju dalam pengembangan perangkat lunak yang mengimplementasikan algoritma <i>K-Means</i> untuk mengklasifikasikan koleksi buku di Perpustakaan Lemhannas RI. Tanggapan saya, mungkin nanti kamu bisa konsultasikan lebih lanjut kepada saya meliputi penulisan kode, pengujian sistem, serta memastikan perangkat lunak berjalan dengan baik yang tentunya harus sesuai dengan kebutuhan pengguna.
Menurut Bapak, apa saja tantangan utama yang mungkin akan dihadapi terkait dengan ide ini?	Menurut saya, tantangan utama yang mungkin akan dihadapi adalah kamu akan mengolah data yang cukup besar dan kompleks dari koleksi perpustakaan. Saran dari saya, data tersebut terlebih dahulu harus dibersihkan dan dipersiapkan dengan baik sebelum dapat digunakan dalam algoritma <i>K-Means</i> . Selain itu, kamu juga harus memastikan bahwa hasil <i>Clustering</i> relevan dan bermanfaat bagi perpustakaan.
Menurut Bapak, Bagaimana cara mengatasi tantangan tersebut?	Kamu mungkin harus menggunakan berbagai teknik pra-pemrosesan data, seperti normalisasi dan pembersihan data, untuk memastikan data yang dimasukkan ke dalam algoritma sudah siap. Selain itu, kamu juga harus

Wawancara dengan Bapak Farid (Pengembang Perangkat Lunak)	
	melakukan beberapa iterasi pengujian dan penyempurnaan algoritma untuk memastikan hasil yang optimal.

Tabel 2 : Notulensi Wawancara Dengan Sistem Analis

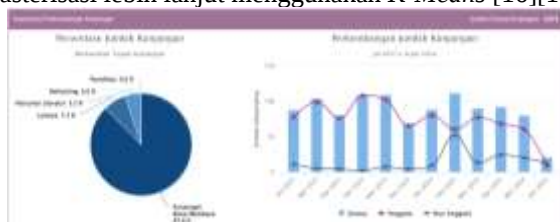
Wawancara dengan Ibu Ita Puspitasari (Analis Sistem)	
Mohon izin bertanya, Ibu Ita. Saya mempunyai ide terkait dengan penerapan data mining <i>Clustering</i> data perpustakaan menggunakan algoritma <i>K-Means</i> di Perpustakaan Lemhannas RI, bagaimana tanggapan Ibu, terkait ide tersebut?	Sebagai analis sistem di perpustakaan Lemhannas RI. Saya setuju dalam pengembangan perangkat lunak yang mengimplementasikan algoritma <i>K-Means</i> untuk mengklasifikasikan koleksi buku di Perpustakaan Lemhannas RI. Tanggapan saya, mungkin nanti kamu bisa konsultasikan lebih lanjut kepada saya terkait dengan pengumpulan dan penganalisan kebutuhan dari pengguna, serta perancangan sistem yang dapat memenuhi kebutuhan tersebut. Selain itu, nanti kamu juga bisa bekerja sama dengan tim pengembang untuk memastikan sistem yang dikembangkan sesuai dengan spesifikasi yang telah ditetapkan di perpustakaan Lemhannas RI.
Menurut Ibu, apa saja tantangan utama yang mungkin akan dihadapi terkait dengan ide ini?	Menurut saya, tantangan utama yang mungkin akan dihadapi adalah Bagaimana kedepannya cara kamu memahami secara mendalam kebutuhan pengguna dan bagaimana mereka berinteraksi dengan sistem perpustakaan saat ini. Hal ini tentu penting agar solusi yang kamu rancang benar-benar membantu mereka. Selain itu, kamu juga harus memastikan bahwa data yang digunakan dalam <i>Clustering</i> akurat dan representatif.
Menurut Ibu, Bagaimana cara mengatasi tantangan tersebut?	Kamu harus melakukan uji coba sistem dengan pengguna untuk mendapatkan umpan balik dan melakukan penyesuaian yang diperlukan.

Tabel 3 : Notulensi Wawancara Dengan Kepala Perpustakaan

Wawancara dengan Bapak Suparmo (Kepala Perpustakaan)	
Mohon izin bertanya Bapak, Saya mempunyai ide terkait dengan penerapan data mining <i>Clustering</i> data perpustakaan menggunakan algoritma <i>K-Means</i> di Perpustakaan Lemhannas RI, bagaimana tanggapan Bapak, terkait ide tersebut?	Saya setuju dalam pengembangan perangkat lunak yang mengimplementasikan algoritma <i>K-Means</i> untuk mengklasifikasikan koleksi buku di Perpustakaan Lemhannas RI. Saya melihat ide ini sangat bermanfaat bagi pengelolaan koleksi perpustakaan. Dengan adanya <i>Clustering</i> data, akan dapat mengelompokkan buku-buku berdasarkan kesamaan tertentu, yang memudahkan dalam penataan dan pencarian buku oleh para pengguna.

Wawancara dengan Bapak Suparmo (Kepala Perpustakaan)	
Menurut Bapak, apa saja tantangan utama yang mungkin akan dihadapi terkait dengan ide ini?	Tantangannya adalah penyesuaian dengan sistem baru, terutama bagi staf yang sudah terbiasa dengan cara kerja lama. Namun, dengan penjelasan yang rinci dan dukungan dari tim pengembang, mungkin kami akan dapat beradaptasi dengan baik.
Menurut Bapak, Bagaimana cara mengatasi tantangan tersebut?	Setelah ide ini terealisasi tim pengembang akan menyediakan pelatihan dan dukungan teknis yang cukup, serta melakukan penyesuaian pada sistem berdasarkan umpan balik dari kami. Hal ini tentunya akan sangat membantu dalam mempercepat proses adaptasi dan memastikan sistem berjalan sesuai dengan kebutuhan kami.

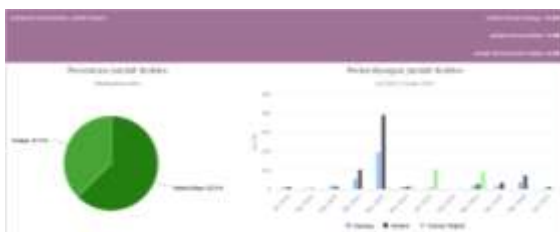
Studi dokumen dilakukan dengan menganalisis berbagai dokumen yang terkait dengan koleksi perpustakaan, seperti katalog buku, laporan peminjaman, dan data statistik penggunaan koleksi selama 6 bulan terakhir. Dokumen ini memberikan informasi mengenai jenis koleksi yang tersedia, frekuensi peminjaman, dan tren penggunaan dari waktu ke waktu. Dengan analisis ini, peneliti dapat mengidentifikasi pola penggunaan koleksi yang berulang, yang kemudian akan menjadi data untuk klusterisasi lebih lanjut menggunakan *K-Means* [16][17].



Gambar 2 : Grafik Persentase Jumlah Kunjungan Perpustakaan Lemhannas RI



Gambar 3 : Grafik Persentase Jumlah Anggota Perpustakaan Lemhannas RI



Gambar 4 : Grafik Persentase Jumlah Koleksi Perpustakaan Lemhannas RI

Pengumpulan data elektronik mencakup data dari sistem manajemen perpustakaan, seperti database koleksi, data peminjaman, dan data pengguna. Sebelum dianalisis dengan *K-Means*, data ini diproses terlebih dahulu untuk

memastikan kualitasnya. Proses pembersihan data dilakukan untuk menghilangkan data duplikat dan menangani missing values. Data yang relevan seperti frekuensi peminjaman per koleksi dan jenis koleksi yang dipinjam akan diekstraksi dan dikonversi menjadi format numerik agar dapat dianalisis menggunakan algoritma *K-Means* [18][19]. Data ini kemudian diimputasi dan dipersiapkan untuk *Clustering* berdasarkan pola penggunaan yang terdeteksi.



Gambar 5 : Portal Utama



Gambar 6 : Kebijakan tentang Keanggotaan



Gambar 7 : Prosedur Sistem dan Jam Layanan

b. Proses Analisis dengan *K-Means*

Pada tahap analisis, algoritma *K-Means* digunakan untuk mengelompokkan koleksi perpustakaan berdasarkan kesamaan karakteristik dan pola penggunaan. Untuk menentukan jumlah kluster yang optimal, peneliti menggunakan metode elbow, yaitu dengan menghitung nilai *within-cluster sum of squares* (WCSS) untuk berbagai nilai K (jumlah kluster) dan memilih K yang menghasilkan penurunan terbesar dalam WCSS. Metode ini membantu dalam menentukan jumlah kluster yang paling sesuai dengan distribusi data koleksi perpustakaan. Setelah jumlah kluster ditentukan, algoritma *K-Means* dijalankan untuk mengelompokkan koleksi berdasarkan kesamaan karakteristik yang terdeteksi dalam data [20].

c. Penggunaan *Progressive Web App* (PWA)

Untuk mendukung tujuan penelitian, sistem yang dikembangkan berbasis *Progressive Web App* (PWA), yang memungkinkan pengelolaan koleksi perpustakaan secara lebih efisien dan interaktif. PWA dipilih karena

kemampuannya untuk diakses melalui perangkat apa pun, baik desktop maupun mobile, tanpa memerlukan aplikasi terpisah yang rumit. PWA memungkinkan pengguna untuk mengakses informasi mengenai koleksi perpustakaan dan hasil analisis *Clustering* secara langsung dan responsif, sehingga meningkatkan pengalaman pengguna dalam menemukan koleksi yang relevan. Selain itu, PWA juga memungkinkan pengelola perpustakaan untuk memonitor penggunaan koleksi secara real-time, yang membantu dalam pengelolaan koleksi yang lebih dinamis dan berbasis data [21].

Dengan menggunakan pendekatan kuantitatif yang sistematis, penelitian ini bertujuan untuk mengembangkan dan mengimplementasikan sistem data mining berbasis *K-Means* yang dapat mengelompokkan koleksi perpustakaan di Lembaga Ketahanan Nasional RI (Lemhannas RI). Dengan memanfaatkan berbagai teknik pengumpulan data yang melibatkan observasi langsung, wawancara, studi dokumen, dan data elektronik, diharapkan penelitian ini dapat memberikan wawasan yang lebih mendalam mengenai penggunaan koleksi perpustakaan dan mendukung pengelolaan koleksi yang lebih efisien di masa depan. Penerapan sistem berbasis *Progressive Web App* (PWA) juga diharapkan dapat meningkatkan pengalaman pengguna dengan memudahkan akses informasi dan hasil analisis secara lebih cepat dan mudah diakses.

III. HASIL PENELITIAN

Penelitian ini bertujuan untuk mengidentifikasi pola peminjaman buku di Perpustakaan Lemhannas RI melalui penerapan algoritma *K-Means*, sebuah metode dalam data mining yang digunakan untuk mengelompokkan koleksi buku berdasarkan pola peminjaman oleh pengguna. Pengelompokan ini diharapkan dapat memberikan wawasan yang lebih baik mengenai preferensi pengguna, tren peminjaman, serta kebutuhan pengelolaan koleksi yang lebih efisien. Dengan memanfaatkan teknik *Clustering*, diharapkan pengelola perpustakaan dapat mengoptimalkan pengadaan dan penyusunan koleksi buku sesuai dengan kebutuhan pengguna.

a. Dataset

Dataset yang digunakan dalam penelitian ini mencakup lebih dari 1.500 transaksi peminjaman buku yang dilakukan oleh lebih dari 500 pengguna dalam periode tertentu. Setiap transaksi mencatat informasi terkait peminjam dan buku yang dipinjam, termasuk ID Pengguna, ID Buku, Judul Buku, Kategori Buku, Frekuensi Peminjaman, Durasi Peminjaman, dan Tanggal Peminjaman. Data ini memberikan gambaran lengkap tentang interaksi pengguna dengan koleksi buku, memungkinkan untuk analisis lebih lanjut mengenai pola peminjaman dan preferensi pengguna terhadap buku tertentu.

Beberapa kolom utama dalam dataset ini adalah sebagai berikut:

1. ID Pengguna: Setiap peminjam memiliki ID unik yang digunakan untuk mengidentifikasi siapa yang meminjam buku.
2. ID Buku: Kode unik yang menunjukkan buku yang dipinjam.
3. Judul Buku: Nama atau judul buku yang dipinjam.

4. Kategori Buku: Kategori buku yang dipinjam, yang dapat mencakup berbagai topik seperti keamanan negara, pertahanan nasional, strategi pertahanan, kebijakan luar negeri, dan lainnya.
5. Frekuensi Peminjaman: Banyaknya peminjaman buku dalam periode yang dianalisis.
6. Durasi Peminjaman: Lama waktu buku dipinjam, dihitung dalam satuan hari.
7. Tanggal Peminjaman: Tanggal peminjaman dilakukan.

Dataset ini memiliki potensi besar untuk memberikan wawasan terkait penggunaan koleksi perpustakaan, serta membantu pengelola dalam mengambil keputusan terkait pengadaan dan pemeliharaan koleksi buku. Namun, untuk memastikan bahwa hasil analisis yang diperoleh dari dataset ini akurat dan bermanfaat, diperlukan serangkaian langkah pra-pemrosesan data yang menyeluruh.

b. Langkah-langkah Pra-pemrosesan Data

Sebelum penerapan algoritma *K-Means*, serangkaian tahap pra-pemrosesan dilakukan untuk memastikan data yang digunakan berkualitas dan siap untuk dianalisis. Langkah-langkah pra-pemrosesan ini mencakup pembersihan data (*data cleaning*), normalisasi, penghapusan outlier, dan transformasi variabel kategorikal.

1. Pembersihan Data (*Data Cleaning*)

Langkah pertama dalam pra-pemrosesan adalah pembersihan data, yang bertujuan untuk mengidentifikasi dan menangani masalah dalam dataset seperti nilai yang hilang (*missing values*), duplikasi, atau data yang tidak relevan. Proses ini dilakukan dengan langkah-langkah sebagai berikut:

- Identifikasi dan Penanganan *Missing Values*: Beberapa data, seperti durasi peminjaman atau frekuensi peminjaman, mungkin memiliki nilai yang hilang. Untuk menangani masalah ini, data yang hilang pada atribut numerik seperti durasi peminjaman dan frekuensi peminjaman diimputasi dengan nilai rata-rata dari masing-masing kolom. Sedangkan untuk data kategorikal seperti kategori buku, nilai yang hilang diimputasi dengan modus kategori yang paling sering muncul.
- Penghapusan Data Duplikat: Beberapa transaksi mungkin tercatat lebih dari sekali dalam dataset. Penghapusan duplikat dilakukan untuk memastikan bahwa setiap transaksi hanya dihitung satu kali, menghindari bias dalam analisis yang disebabkan oleh data yang ganda.
- Penyaringan Data Tidak Relevan: Jika terdapat transaksi yang tidak relevan atau anomali dalam dataset (seperti transaksi yang tidak sesuai dengan konteks penelitian), baris data tersebut akan dihapus untuk menjaga kualitas analisis.

2. Normalisasi Data

Algoritma *K-Means* sangat bergantung pada jarak antara titik data dalam ruang fitur untuk menentukan kluster. Oleh karena itu, penting untuk memastikan bahwa semua variabel dalam dataset memiliki skala yang seragam. Misalnya, durasi peminjaman yang berkisar antara beberapa hari hingga 30 hari, dan frekuensi peminjaman yang berkisar antara 1 hingga 10 kali, keduanya memiliki rentang yang sangat berbeda. Jika tidak dinormalisasi, variabel dengan rentang lebih besar,

seperti durasi peminjaman, akan mendominasi proses *Clustering*, sementara frekuensi peminjaman yang lebih kecil akan kurang diperhatikan.

Untuk menangani masalah ini, dilakukan *Min-Max Scaling* pada atribut numerik. Dengan teknik ini, nilai-nilai dalam setiap kolom data akan dipetakan dalam rentang [0, 1], sesuai dengan rumus berikut:

$$X_{\text{norm}} = \frac{X - X_{\text{min}}}{X_{\text{maks}} - X_{\text{min}}}$$

Metode ini memastikan bahwa setiap variabel memiliki skala yang setara dan tidak ada satu pun variabel yang mendominasi dalam proses *Clustering*.

3. Penghapusan Outlier

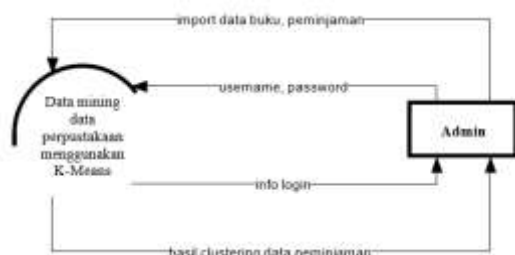
Outlier atau nilai yang sangat jauh dari distribusi data lainnya dapat mempengaruhi hasil *Clustering* dengan membuat beberapa kluster terlihat tidak representatif. Oleh karena itu, penting untuk mengidentifikasi dan menghapus outlier yang ada dalam dataset. Untuk mendeteksi outlier, digunakan *Z-score* dengan nilai ambang batas lebih besar dari 3 atau kurang dari -3. Data yang memiliki *Z-score* lebih besar dari 3 atau kurang dari -3 dianggap sebagai outlier dan dihapus dari dataset. Hal ini bertujuan untuk menjaga agar hasil *Clustering* tetap akurat dan menggambarkan pola peminjaman yang sebenarnya.

4. Transformasi Variabel Kategorikal

Variabel kategori buku, yang berupa data kategorikal, harus diubah menjadi format numerik agar bisa diproses oleh algoritma *K-Means*. Untuk itu, digunakan teknik *One-Hot Encoding*, di mana setiap kategori buku diubah menjadi kolom baru dengan nilai 0 atau 1. Nilai 1 menandakan bahwa buku dalam transaksi tersebut termasuk dalam kategori tertentu, sementara 0 menandakan sebaliknya. Dengan cara ini, kategori buku dapat dimasukkan dalam proses analisis tanpa mengubah makna kategorikalnya.

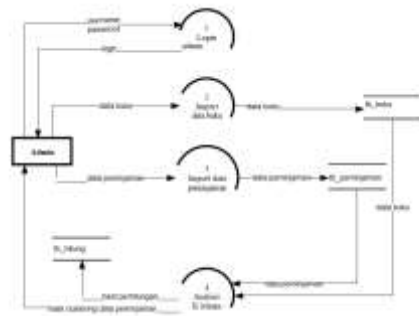
Diagram Konteks

Diagram konteks merupakan representasi yang mencakup input-input utama, sistem secara keseluruhan, dan output yang dihasilkan. Diagram ini berada pada tingkat paling atas dalam hierarki diagram aliran data dan hanya menampilkan satu proses saja.



Gambar 8 : Diagram Konteks

Proses secara keseluruhan yang diilustrasikan dalam diagram konteks akan dijelaskan dengan lebih rinci pada DFD level 1, guna memberikan kejelasan mengenai fungsi-fungsi serta aliran data yang ada dalam sistem.



Gambar 9 : DFD Level 1

Tabel peminjaman menyajikan rincian transaksi peminjaman yang berlangsung di perpustakaan Lemhannas RI.

Tabel I. TB_PEMINJAMAN

No	Nama Field	Type	Size
1	id peminjaman	int	11
2	id buku	int	11
3	bulan	Varchar	11
4	jumlah pinjam	varchar	11

Gambar 3 : Tabel Peminjaman

Tabel buku menyimpan informasi dasar mengenai buku-buku yang terlibat dalam transaksi.

Tabel II. LABEL TB_BUKU

No	Nama Field	Type	Size
1	id buku	int	11
2	kategori	varchar	20
3	judul buku	varchar	100
4	pengarang	varchar	50
5	penerbit	varchar	20
6	tahun terbit	varchar	5
7	jumlah buku	varchar	5

Gambar 4 : Tabel Buku

Tabel hitung menyimpan hasil dari perhitungan yang dilakukan menggunakan metode *K-Means*.

Tabel III. STRUKTUR TB_HITUNG

No	Nama Field	Type	Size
1	id hitung	int	11
2	ulangan	int	11
3	ses hitung	varchar	128

Gambar 5 : Tabel Hitung

Tabel kategori terdiri dari dua kolom, yaitu *id_kategori* dan *nama_kategori*. Tabel ini menyimpan informasi mengenai kategori buku.

Tabel IV. TABEL TB_KATEGORI

No	Nama field	Type	Size
1	id_kategori	varchar	11
2	nama_kategori	varchar	50

Gambar 10 : Tabel Kategori

c. Hasil Pengelompokan dengan Algoritma *K-Means*

Setelah data melalui serangkaian tahapan pra-pemrosesan yang menyeluruh, algoritma *K-Means* diterapkan untuk mengelompokkan koleksi buku berdasarkan pola peminjaman. Salah satu tantangan utama dalam penerapan *K-Means* adalah menentukan jumlah kluster yang optimal. Untuk itu, digunakan metode Elbow untuk mengevaluasi *Within-Cluster Sum of Squares* (WCSS), yang mengukur total jarak antara setiap titik data dan pusat kluster. Dengan memplot WCSS terhadap jumlah kluster, penurunan yang tajam menunjukkan jumlah kluster yang optimal.

Berdasarkan hasil metode Elbow, ditemukan bahwa 4 kluster adalah jumlah kluster optimal. Setelah mencapai 4

klaster, penurunan WCSS melambat, yang menandakan bahwa menambah jumlah klaster lebih lanjut tidak akan memberikan hasil yang lebih baik.

Hasil pengelompokan menunjukkan empat klaster utama dengan karakteristik yang berbeda:

1. Klaster 1: Buku yang sering dipinjam dengan durasi peminjaman yang pendek, seperti buku mengenai strategi pertahanan dan keamanan negara. Buku-buku dalam klaster ini cenderung dicari untuk referensi cepat dan mendalam mengenai isu-isu strategis.
2. Klaster 2: Buku dengan peminjaman yang jarang namun durasi peminjaman panjang, seperti buku tentang analisis kebijakan pertahanan yang lebih banyak digunakan untuk riset mendalam.
3. Klaster 3: Buku yang memiliki frekuensi peminjaman sedang, mencakup topik-topik seperti keamanan siber dan teknologi pertahanan yang banyak diminati oleh pengguna yang tertarik pada perkembangan teknologi terbaru dalam pertahanan.
4. Klaster 4: Buku yang sangat jarang dipinjam, tetapi memiliki nilai referensial tinggi dalam jangka panjang, seperti dokumen kebijakan ketahanan nasional yang digunakan dalam konteks riset tingkat tinggi.

d. Evaluasi Hasil Pengelompokan

Untuk mengevaluasi kualitas hasil *Clustering*, digunakan beberapa metrik yang umum digunakan dalam analisis *Clustering*, antara lain *Silhouette Score*, WCSS, dan *Purity*.

1. *Silhouette Score*: Nilai *Silhouette Score* yang diperoleh dalam penelitian ini adalah 0,68, yang mengindikasikan kualitas *Clustering* yang baik. Nilai ini menunjukkan bahwa data dalam setiap klaster memiliki kesamaan yang lebih besar dibandingkan dengan data di klaster lainnya, menandakan pengelompokan yang efektif.
2. *Within-Cluster Sum of Squares (WCSS)*: Penurunan WCSS yang tajam setelah penambahan jumlah klaster keempat menunjukkan bahwa klaster yang terbentuk sudah cukup homogen dan mewakili kelompok dengan karakteristik peminjaman yang serupa.
3. *Purity*: Nilai *Purity* yang dihitung sebesar 0,85 menunjukkan bahwa sebagian besar data dalam setiap klaster sesuai dengan kategori yang diharapkan. Hal ini menunjukkan bahwa algoritma *K-Means* dapat mengelompokkan buku dengan baik, sesuai dengan preferensi pengguna yang tercermin dalam pola peminjaman.
- e. Pentingnya Hasil untuk Pengelolaan Koleksi Perpustakaan

Hasil dari pengelompokan ini memberikan informasi yang sangat berguna bagi pengelola perpustakaan dalam menyusun koleksi. Dengan memahami pola peminjaman yang ada, pengelola perpustakaan dapat menyesuaikan pengadaan buku baru dengan preferensi pengguna. Sebagai contoh, klaster yang berisi buku dengan peminjaman sering dan durasi pendek akan membutuhkan buku-buku tambahan dengan topik yang serupa, sementara klaster dengan buku jarang dipinjam mungkin menunjukkan kebutuhan untuk penyegaran atau perbaikan dalam pengelolaan koleksi tersebut.

Selain itu, penerapan hasil *Clustering* ini juga dapat meningkatkan pengalaman pengguna, dengan memberikan rekomendasi buku yang relevan dengan

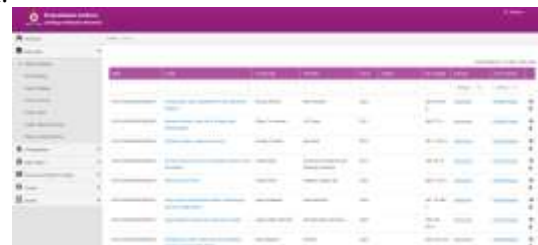
minat mereka. Misalnya, pengguna yang sering meminjam buku dalam klaster keamanan siber dapat diarahkan untuk meminjam buku dalam klaster yang serupa, meningkatkan keterlibatan dan kepuasan mereka terhadap layanan perpustakaan.

Penerapan algoritma *K-Means* dalam penelitian ini berhasil mengelompokkan koleksi buku di Perpustakaan Lemhannas RI berdasarkan pola peminjaman dengan baik. Evaluasi hasil *Clustering* menggunakan berbagai metrik menunjukkan kualitas pengelompokan yang baik, yang dapat diterjemahkan dalam peningkatan pengelolaan koleksi perpustakaan yang lebih efisien dan sesuai dengan kebutuhan pengguna. Hasil ini menunjukkan bahwa data mining dan *Clustering* dapat digunakan untuk membuat keputusan yang lebih tepat dalam pengelolaan koleksi perpustakaan dan meningkatkan pengalaman pengguna secara keseluruhan.

IV. PEMBAHASAN

Temuan yang signifikan dan analisis terkait penelitian ini adalah sistem yang diimplementasikan untuk data mining *Clustering* ini bertujuan untuk mengidentifikasi penambahan koleksi buku di perpustakaan Lemhannas RI, dan dapat menghasilkan *Clustering* dari data transaksi peminjaman buku. Website ini dirancang untuk mempermudah perpustakaan Lemhannas RI dalam menentukan penambahan koleksi buku melalui analisis data transaksi peminjaman yang dilakukan.

Halaman buku adalah halaman yang berfungsi untuk menampilkan informasi dasar mengenai buku-buku yang tercantum dalam tabel buku di perpustakaan Lemhannas RI.



Gambar 11 : Halaman Buku

Halaman peminjaman berfungsi menampilkan seluruh data transaksi peminjaman buku yang terdapat dalam tabel transaksi peminjaman buku di perpustakaan Lemhannas RI.



Gambar 12 : Halaman Peminjaman

Halaman *cluster* memiliki fungsi untuk melaksanakan proses *Clustering*. Di bagian atas, terdapat opsi untuk memilih bulan dan tahun yang digunakan untuk menentukan rentang waktu data transaksi peminjaman yang akan dianalisis. Sedangkan tombol hitung berfungsi untuk menjalankan perhitungan dengan algoritma *K-Means*.



Gambar 13 : Halaman Cluster

f. Interpretasi Hasil Clustering

Hasil dari penerapan algoritma *K-Means* pada data transaksi peminjaman buku menghasilkan beberapa kluster yang memiliki karakteristik dan pola yang berbeda. Berikut adalah penjelasan lebih lanjut mengenai interpretasi masing-masing kluster yang terbentuk:

1. Kluster 1: Buku dengan Peminjaman Tinggi dalam Waktu Singkat

- Karakteristik: Buku yang termasuk dalam kluster ini dipinjam dengan frekuensi yang sangat tinggi namun dengan durasi peminjaman yang relatif singkat. Buku-buku ini cenderung berfokus pada topik-topik praktis dan informatif terkait dengan strategi keamanan nasional, kebijakan pertahanan, dan peraturan terbaru yang sering digunakan oleh pengguna dalam pengambilan keputusan yang cepat.
- Interpretasi: Kluster ini mencerminkan kelompok buku yang memiliki relevansi yang tinggi dalam konteks kebutuhan pengguna yang mendesak, seperti profesional atau pejabat yang membutuhkan informasi praktis dalam waktu singkat.

2. Kluster 2: Buku dengan Peminjaman Lama namun Jarang

- Karakteristik: Buku-buku dalam kluster ini dipinjam dengan durasi yang lebih lama, tetapi dengan frekuensi peminjaman yang lebih rendah. Buku ini cenderung membahas topik yang lebih mendalam, seperti analisis kebijakan strategis, teori pertahanan, dan sejarah pertahanan negara.
- Interpretasi: Kluster ini menggambarkan minat pengguna yang terlibat dalam riset mendalam atau studi lanjutan, yang memerlukan waktu lebih lama untuk memahami materi yang kompleks. Pengguna yang meminjam buku dalam kluster ini biasanya terlibat dalam penelitian atau analisis kebijakan jangka panjang.

3. Kluster 3: Buku dengan Minat Sedang pada Topik Teknologi

- Karakteristik: Kluster ini berisi buku yang terkait dengan keamanan siber, teknologi pertahanan, dan inovasi dalam teknologi keamanan. Buku-buku dalam kluster ini dipinjam secara stabil dengan durasi peminjaman yang sedang.
- Interpretasi: Kluster ini menunjukkan bahwa ada minat yang berkembang terhadap teknologi dalam bidang pertahanan, dengan banyak pengguna yang mencari informasi terkait dengan teknologi baru yang dapat diterapkan dalam kebijakan keamanan nasional.

4. Kluster 4: Buku dengan Peminjaman Langka dan Relevansi Jangka Panjang

- Karakteristik: Buku dalam kluster ini jarang dipinjam, tetapi memiliki nilai referensi yang sangat

tinggi dan cenderung digunakan untuk riset jangka panjang. Buku ini sering kali berkaitan dengan sejarah pertahanan atau arsip kebijakan pertahanan yang lebih bersifat arsitektur atau fondasi kebijakan.

- Interpretasi: Kluster ini berisi buku yang jarang digunakan tetapi memiliki relevansi yang besar dalam penelitian atau pengembangan kebijakan jangka panjang. Pengguna yang meminjam buku dalam kluster ini kemungkinan besar terlibat dalam riset yang membutuhkan analisis sejarah atau teori dasar dalam bidang pertahanan.

g. Manfaat Pengelompokan dengan Algoritma *K-Means*

Implementasi algoritma *K-Means* dalam pengelompokan data peminjaman buku memberikan sejumlah manfaat yang signifikan dalam pengelolaan koleksi perpustakaan di Lemhannas RI:

1. Optimasi Pengelolaan Koleksi: Pengelompokan buku berdasarkan pola peminjaman memungkinkan pengelola untuk menyusun koleksi lebih efisien dan terstruktur. Buku dapat dikelompokkan ke dalam kategori-kategori yang relevan, sehingga memudahkan pengguna untuk menemukan buku yang mereka butuhkan.
2. Personalisasi Layanan Pengguna: Dengan memahami lebih baik preferensi dan minat pengguna, perpustakaan dapat memberikan layanan yang lebih tepat sasaran. Hal ini termasuk memberikan rekomendasi buku berdasarkan pola peminjaman yang teridentifikasi melalui *Clustering*.
3. Efisiensi Pengadaan Buku: Berdasarkan hasil *Clustering*, pengelola perpustakaan dapat mengidentifikasi kebutuhan koleksi yang belum terpenuhi, sehingga memudahkan dalam pengadaan buku baru yang lebih sesuai dengan minat pengguna dan menghindari pemborosan sumber daya.
4. Pengambilan Keputusan Berdasarkan Data: Analisis *Clustering* berbasis data memberikan pengelola perpustakaan dengan wawasan yang lebih mendalam mengenai tren peminjaman dan kebutuhan pengguna. Dengan demikian, keputusan yang diambil lebih didasarkan pada informasi yang akurat dan relevan.
5. Peningkatan Pengalaman Pengguna: Dengan sistem rekomendasi buku yang didukung oleh hasil *Clustering*, pengguna dapat menemukan buku yang lebih relevan dengan cepat. Hal ini meningkatkan kenyamanan dan efisiensi penggunaan perpustakaan.

g. Keunggulan dan Keterbatasan Algoritma *K-Means*

Berdasarkan hasil penelitian, berikut adalah analisis keunggulan dan keterbatasan algoritma *K-Means* yang diterapkan pada data transaksi peminjaman buku di perpustakaan Lemhannas RI:

Keunggulan:

1. Sederhana dan Efisien: *K-Means* merupakan algoritma yang relatif sederhana dan cepat diterapkan, membuatnya ideal untuk pengelolaan koleksi buku dalam jumlah besar.
2. Skalabilitas: *K-Means* mampu menangani dataset yang besar dengan cepat, sesuai dengan jumlah data peminjaman buku yang ada di perpustakaan Lemhannas RI.
3. Fleksibilitas: *K-Means* dapat diterapkan pada berbagai jenis data, baik numerik maupun kategori, setelah dilakukan pra-pemrosesan yang tepat.

Keterbatasan:

1. Kebutuhan Penentuan K: Salah satu keterbatasan *K-Means* adalah kebutuhan untuk menentukan jumlah kluster (*K*) sebelumnya. Jika *K* yang ditentukan tidak tepat, hasil *Clustering* dapat menjadi kurang relevan.
2. Sensitif terhadap Outlier: *K-Means* sangat sensitif terhadap data outlier, yang dapat mengganggu proses *Clustering* dan menghasilkan kluster yang tidak representatif.
3. Asumsi Bentuk Kluster yang Terbentuk Sederhana: *K-Means* mengasumsikan bahwa kluster berbentuk bulat atau sferis. Jika distribusi data tidak sesuai dengan asumsi ini, hasil *Clustering* yang dihasilkan bisa kurang optimal.

h. Perbandingan dengan Metode Lain

Untuk memperkaya analisis, perbandingan dilakukan antara *K-Means* dengan metode *Hierarchical Clustering* dan *DBSCAN*:

- *Hierarchical Clustering* lebih fleksibel dalam hal menentukan jumlah kluster dan tidak memerlukan pemilihan *K* di awal. Namun, metode ini lebih lambat dibandingkan dengan *K-Means*, terutama untuk dataset yang besar.
- *DBSCAN* lebih unggul dalam menangani kluster dengan bentuk yang lebih kompleks dan mampu mendeteksi outlier dengan lebih baik. Meskipun demikian, *DBSCAN* memerlukan pemilihan parameter seperti *eps* dan *min_samples*, yang dapat mempengaruhi hasil *Clustering* secara signifikan.

i. Validasi Hasil *Clustering*

Untuk memvalidasi keakuratan dan relevansi hasil *Clustering*, sejumlah metrik evaluasi digunakan:

1. *Silhouette Score* menunjukkan bahwa hasil *Clustering* dengan *K-Means* memiliki skor yang cukup tinggi, menandakan bahwa kluster yang terbentuk memiliki pemisahan yang jelas antara satu kluster dengan kluster lainnya.
2. *Within-Cluster Sum of Squares (WCSS)* menunjukkan bahwa kluster yang terbentuk memiliki keseragaman yang baik, dengan jarak antar data dalam kluster yang relatif kecil.
3. Visualisasi dengan PCA dan t-SNE mendukung temuan bahwa kluster yang terbentuk jelas terlihat terpisah dalam ruang multidimensi, memberikan indikasi bahwa *K-Means* berhasil mengelompokkan data dengan baik.

j. Pengumpulan Feedback Pengguna

Untuk mengevaluasi keefektifan sistem *Clustering* dalam memenuhi kebutuhan pengguna, feedback pengguna dikumpulkan melalui survei dan wawancara. Hasil feedback menunjukkan bahwa pengguna merasa terbantu dengan adanya sistem rekomendasi buku yang lebih relevan berdasarkan pola peminjaman. Sebagian besar pengguna juga mengapresiasi kemudahan dalam mencari buku yang sesuai dengan minat mereka melalui sistem yang lebih terstruktur.

Penerapan algoritma *K-Means* dalam pengelompokan koleksi buku di Perpustakaan Lemhannas RI telah memberikan dampak yang signifikan dalam pengelolaan koleksi, layanan pengguna, serta pengambilan keputusan berbasis data. Meskipun demikian, penting untuk terus melakukan evaluasi terhadap algoritma yang digunakan dan mengembangkan

sistem yang dapat beradaptasi dengan kebutuhan pengguna yang dinamis.

V. KESIMPULAN

Implementasi algoritma *K-Means Clustering* dalam pengelolaan koleksi Perpustakaan Lemhannas RI terbukti efektif untuk meningkatkan pengelolaan koleksi dan pengalaman pengguna. Dengan menganalisis pola peminjaman, sistem ini membantu pengelola perpustakaan dalam menata koleksi secara efisien, seperti menempatkan buku yang sering dipinjam pada lokasi strategis dan mengurangi koleksi yang jarang digunakan. Selain itu, sistem ini juga memfasilitasi pengadaan koleksi yang lebih relevan dan memberikan rekomendasi buku yang lebih tepat kepada pengguna berdasarkan minat mereka.

Namun, tantangan utama yang dihadapi adalah kualitas data yang harus diperbaiki untuk memastikan hasil yang akurat, serta kebutuhan untuk pengetahuan teknis dalam implementasi. Oleh karena itu, evaluasi berkelanjutan dan pelatihan pengguna menjadi penting untuk memaksimalkan manfaat dari sistem ini. Dengan demikian, penerapan teknik *K-Means Clustering* memberikan dampak positif dalam pengelolaan koleksi dan peningkatan layanan di perpustakaan, meskipun membutuhkan upaya berkelanjutan dalam pemeliharaan dan pengembangan sistem.

UCAPAN TERIMA KASIH

Penulis ingin menyampaikan rasa terima kasih yang mendalam kepada pihak Lemhannas RI, terutama kepada Perpustakaan Lemhannas RI, atas izin dan akses yang diberikan untuk melaksanakan penelitian ini. Dukungan serta kerjasama dari staf perpustakaan sangat berarti dalam proses pengumpulan data dan pengembangan sistem *Clustering* data perpustakaan, khususnya terkait koleksi yang ada. Penulis berharap hasil studi ini memberikan kontribusi positif untuk pengelolaan koleksi perpustakaan di Lemhannas RI serta perpustakaan lainnya di masa yang akan datang.

VI. REFERENSI

- [1] Lund, B. and Ma, J. A review of cluster analysis techniques and their uses in library and information science research: *K-Means and k-medoids Clustering*. Performance Measurement and Metrics, 22(3), 161-173. 2021.
- [2] Qin, X. Resource classification and knowledge aggregation of library and information based on data mining. *Ingénierie Des Systèmes D Information*, 25(5), 645-653. 2020.
- [3] Fammaldo, E., & Hakim, L. (2018). Penerapan Algoritma *K-Means Clustering* Untuk Pengelompokan Tingkat Kesejahteraan Keluarga Untuk Program Kartu Indonesia Pintar. *Jurnal Ilmiah Teknologi Infomasi Terapan*, 5(1), 23-31.
- [4] Larasati, A., Hajji, A., Handayani, A., Azzahra, N., Farhan, M., & Rahmawati, P. Profiling academic library patrons using *K-Means* and *x-means Clustering*. "International Journal of Technology", 10(8), 1567. 2019.

- [5] Kovačević, A., Devedžić, V., & Pocajt, V. *Using data mining to improve digital library services*. The Electronic Library, 28(6), 829-843. 2020.
- [6] Ulya, L. F. *Pengelompokan Data Dengan Menggunakan Algoritma K- Means Clustering Pada Perpustakaan Kendal*. "Jurnal Universitas Dian Nuswantoro".
- [7] Alfariqi, M. *Implementasi Data Mining Metode Clustering Pada Toko Buku Bi-Obses Bandung Penerapan Algoritma KMeans*. "Jurnal Ilmiah Komputer dan Informatika (KOMPUTA)".
- [8] Ulfah, M. and Irtawaty, A. *Implementation of data mining Clustering using the k-medoids method in grouping library books politeknik negeri balikpapan*. "Jurnal Ecotipe (Electronic Control Telecommunication Information and Power Engineering)", 9(2), 183-191. 2022.
- [9] Shieh, J. *The integration system for librarians' bibliomining*. The Electronic Library, 28(5), 709-721. 2020.
- [10] Cox, A. and Corral, S. *Evolving academic library specialties*. "Journal of the American Society for Information Science and Technology", 64(8), 1526-1542. 2013.
- [11] Burhanudin, J. *Studi Kinerja Pegawai Layanan Sirkulasi dan Referens di Perpustakaan UIN Sunan Gunung Djati Bandung*. Thesis FIB UI.
- [12] Akhyani, I. *Integrated marketing communication "gadis modis" sebagai usaha mikro kecil dan menengah dalam meningkatkan loyalitas konsumen*. Commicast, 1(1), 10. 2020.
- [13] Pinasthika, R. and Kaltsum, H. *Analisis penggunaan metode eksperimen pada pembelajaran ipa di sekolah dasar*. "Jurnal Basicedu", 6(4), 6558-6566. 2022.
- [14] Sulfikar, S. and Fawzani, N. *Pemanfaatan instagram dalam meningkatkan penguasaan mufradat mahasiswa*. "Jurnal Tahsinia", 4(1), 19-27. 2023.
- [15] Ummah, D., Muis, M., & Lailiyah, S. *Evaluasi program kewirausahaan di madrasah ibtidaiyah persmin wonokromo surabaya*. "At-Thullab Jurnal Pendidikan Guru Madrasah Ibtidaiyah", 7(1), 34. 2023.
- [16] Matje, I. *Hubungan pemberian reward (hadiah) terhadap minat belajar pada siswa sekolah dasar*. taksonomi, 2(2), 122-128. 2022.
- [17] Hakim, R. *Tren dan strategi pengumpulan dana zakat, infak dan sedekah (zis) di masa pandemi covid-19: studi multisitius pada badan amal zakat nasional (baznas) kota malang, kabupaten jombang dan kabupaten tanah laut, kalimantan selatan*. "Jurnal Ilmiah Ekonomi Islam", 9(2), 2431. 2023.
- [18] Oktaviani, E. and Husin, H. *Implementasi pembelajaran tahsin al-qur'an dan amaliyah keagamaan di sekolah dasar*. "Jurnal Basicedu", 6(3), 5063-5075. 2020.
- [19] Fanny, N. and Megawati, S. *Implementasi program bantuan pangan non tunai (bpnt) di kecamatan bancar kabupaten tuban*. Publika, 407-418. 2022.
- [20] Syah, H., Hasan, M., Kamaruddin, C., Nurdiana, N., & Nurjannah, N. *Strategi ketahanan pangan dalam program urban farming dalam menunjang keberlanjutan usaha keluarga di masa pandemi covid-19*. "Ideas Jurnal Pendidikan Sosial Dan Budaya", 8(3), 1093. 2022.
- [21] Nadiva, F. and Cahyadi, N. *Konflik peran ganda dan burnout terhadap kinerja karyawan wanita*. "Jurnal Informatika Ekonomi Bisnis", 221-226. 2022.