

Implementasi K-Means Clustering dalam Pengelompokan Risiko Pinjaman Nasabah pada Koperasi Simpan Pinjam

¹Jihan Martha, ²Danar Putra Pamungkas, ³Made Ayu Dusea Widyadara

^{1,2,3}Teknik Informatika, Universitas Nusantara PGRI Kediri

E-mail: ¹jihannmarthaa@gmail.com, ²danar@unpkediri.ac.id, ³madedara@unpkediri.ac.id

Penulis Korespondens : Jihan Martha

Abstrak— Koperasi Simpan Pinjam Sri Hadi Dharma menghadapi tantangan dalam pengelolaan risiko pinjaman karena proses analisis masih dilakukan secara manual dan bergantung pada penilaian subjektif pengelola. Kondisi tersebut dapat memengaruhi objektivitas dalam pengambilan keputusan terkait pengelolaan pinjaman nasabah. Penelitian ini bertujuan mengelompokkan risiko pinjaman nasabah menggunakan algoritma K-Means *Clustering* berdasarkan karakteristik pinjaman. Data yang digunakan merupakan data riwayat pinjaman nasabah koperasi periode 2021-2025. Tahapan penelitian meliputi *preprocessing* data, proses *clustering* menggunakan K-Means yang dievaluasi melalui pengujian stabilitas *cluster*, serta interpretasi hasil pengelompokan. Hasil penelitian menunjukkan bahwa K-Means mampu membentuk tiga kelompok nasabah dengan karakteristik yang berbeda berdasarkan pola pinjaman dan pembayaran. Pengelompokan yang dihasilkan memberikan informasi yang lebih terstruktur sehingga dapat mendukung analisis risiko pinjaman dan membantu pengambilan keputusan dalam pengelolaan pinjaman nasabah.

Kata Kunci— koperasi, risiko pinjaman, prediksi risiko, k-means

Abstract— The Sri Hadi Dharma Savings and Loan Cooperative faces challenges in managing loan risk because the analysis process is still conducted manually and relies on the subjective judgment of managers. This situation can affect the objectivity of decision-making regarding the management of customer loans. This study aims to cluster customer loan risks using the K-Means Clustering algorithm based on loan characteristics. The data used consists of the cooperative's customer loan history for the period 2021–2025. The research stages include data preprocessing, the clustering process using K-Means evaluated through cluster stability testing, and the interpretation of the clustering results. The results show that K-Means is capable of forming three groups of customers with different characteristics based on loan and payment patterns. The resulting clustering provides more structured information, thereby supporting loan risk analysis and aiding decision-making in customer loan management.

Keywords— cooperative, loan risk, risk prediction, K-Means

This is an open access article under the CC BY-SA License.



I. PENDAHULUAN

Koperasi Simpan Pinjam memiliki peran penting dalam mendukung perekonomian masyarakat melalui penyediaan layanan pembiayaan bagi nasabah [1]. Kemudahan akses pinjaman yang diberikan koperasi membantu memenuhi kebutuhan modal usaha maupun kebutuhan ekonomi lainnya [2]. Namun, tingginya jumlah pinjaman yang dikelola juga diikuti oleh risiko pinjaman, seperti keterlambatan pembayaran angsuran dan potensi terjadinya

pinjaman bermasalah. Kondisi tersebut dapat memengaruhi kesehatan keuangan koperasi dan menghambat keberlanjutan operasional Lembaga [3].

Pengelolaan risiko pinjaman yang masih dilakukan secara manual sering kali bergantung pada penilaian subjektif pengelola koperasi sehingga kurang efektif [4]. Pemanfaatan teknik *data mining* dapat membantu proses analisis secara lebih objektif dan terukur. Salah satu teknik yang banyak digunakan adalah *clustering*, yaitu metode pengelompokan data berdasarkan kemiripan karakteristik untuk menghasilkan segmentasi yang lebih terstruktur [5].

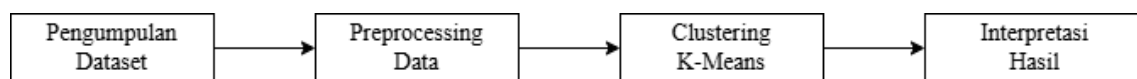
Algoritma K-Means merupakan metode *clustering* yang banyak diterapkan dalam bidang keuangan dan pembiayaan. Metode ini digunakan untuk mengelompokkan nasabah berdasarkan data pinjaman guna mendukung penentuan kelayakan pinjaman berikutnya, serta mengidentifikasi calon nasabah potensial untuk membantu proses persetujuan kredit [6], [7]. K-Means juga digunakan untuk melakukan segmentasi pengguna kartu kredit berdasarkan perilaku penggunaan dan menunjukkan performa yang baik dibandingkan beberapa metode *clustering* lainnya [8]. Selain itu, metode ini mampu menghasilkan segmentasi risiko kredit berdasarkan karakteristik pelanggan pembiayaan [9]. Dalam pengelompokan data penyaluran pinjaman, K-Means juga terbukti memberikan kualitas klaster yang lebih baik dibandingkan metode Complete Linkage [10].

Berdasarkan penelitian terdahulu, algoritma K-Means telah menunjukkan kemampuan yang baik dalam mengelompokkan data pada berbagai bidang keuangan dan pembiayaan. Namun, kajian yang secara khusus membahas pengelompokan risiko pinjaman pada nasabah koperasi simpan pinjam masih belum banyak dikaji secara mendalam [11]. Kondisi tersebut menunjukkan perlunya analisis pengelompokan yang dapat menghasilkan segmentasi risiko secara lebih terstruktur untuk mendukung pengambilan keputusan pada koperasi

Penelitian ini bertujuan menerapkan algoritma K-Means untuk mengelompokkan risiko pinjaman nasabah koperasi simpan pinjam berdasarkan karakteristik data riwayat pinjaman yang dimiliki. Hasil pengelompokan diharapkan dapat menghasilkan segmentasi risiko yang lebih objektif sehingga memberikan informasi mengenai profil risiko nasabah dan membantu koperasi dalam meningkatkan efektivitas pengelolaan pinjaman.

II. METODE

Penelitian ini menerapkan pendekatan *data mining* untuk mengelompokkan risiko pinjaman nasabah pada Koperasi Simpan Pinjam Sri Hadi Dharma menggunakan algoritma K-Means. Tahapan penelitian meliputi pengumpulan data, *preprocessing*, *clustering*, dan interpretasi hasil.



Gambar 1. Tahapan Penelitian

A. Pengumpulan Data

Data yang digunakan merupakan data sekunder berupa riwayat pinjaman nasabah yang diperoleh secara langsung oleh penulis dari pengurus Koperasi Simpan Pinjam Sri Hadi Dharma dalam bentuk dokumen digital. Data tersebut mencakup riwayat pinjaman nasabah periode tahun 2021 hingga 2025 dan digunakan sebagai dasar dalam proses pengolahan serta analisis pengelompokan risiko pinjaman.

B. Preprocessing Data

Preprocessing dilakukan untuk memastikan data yang digunakan memiliki kualitas yang baik sebelum proses *clustering* [12]. Tahapan yang dilakukan meliputi:

1. Pembersihan Data (*Data Cleaning*)
Menghilangkan data yang tidak lengkap atau tidak sesuai.
2. Perhitungan Tunggal Maksimum
Menentukan nilai tunggal tertinggi yang dimiliki setiap ID Nasabah.
3. Reduksi Data (*Data Reduction*)
Memilih satu data untuk setiap ID Nasabah berdasarkan Sisa Tenor terkecil.
4. Seleksi Atribut (*Feature Selection*)
Menentukan atribut yang digunakan dalam proses pengelompokan.
5. Transformasi Data (Normalisasi)
Menyamakan rentang nilai atribut menggunakan metode *Min-Max Normalization*.

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}} \quad (1)$$

Keterangan:

- a. x' adalah nilai hasil normalisasi
- b. x adalah nilai data awal
- c. $x_{\min}$ adalah nilai minimum atribut
- d. $x_{\max}$ adalah nilai maksimum atribut

Normalisasi dilakukan agar setiap atribut memiliki skala yang seragam sehingga proses *clustering* dapat berjalan lebih optimal.

C. Clustering Menggunakan Algoritma K-Means

Proses *clustering* dilakukan menggunakan algoritma K-Means untuk mengelompokkan nasabah berdasarkan kemiripan karakteristik pinjaman [13]. Tahapan yang dilakukan meliputi:

1. Menentukan jumlah *cluster* (k),
2. Menginisialisasi *centroid* awal.
3. Menghitung jarak data terhadap *centroid* menggunakan *Euclidean Distance*.

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (2)$$

Keterangan:

- a. (x, y) adalah jarak antara data dan *centroid*
 - b. x_i adalah nilai atribut ke- i pada data
 - c. y_i adalah nilai atribut ke- i pada *centroid*
 - d. n adalah jumlah atribut yang digunakan
4. Menempatkan data ke *cluster* dengan jarak terdekat.
 5. Memperbarui *centroid* hingga kondisi konvergen tercapai.

$$\mu_k = \frac{1}{N_k} = \sum_{i=1}^{N_k} x_i \quad (3)$$

Keterangan:

- a. μ_k Adalah *centroid* pada *cluster* ke- k
- b. N_k Adalah jumlah data pada *cluster* ke- k

- c. x_i Adalah data ke- i yang menjadi anggota *cluster* ke- k

Stabilitas hasil *clustering* dievaluasi menggunakan *Adjusted Rand Index* (ARI) melalui beberapa variasi *random state* untuk memastikan konsistensi pembentukan *cluster*. Selain itu, *Principal Component Analysis* (PCA) digunakan untuk mereduksi dimensi data sehingga pola persebaran *cluster* dapat divisualisasikan dengan lebih jelas [14].

D. Interpretasi Hasil Cluster

Cluster yang terbentuk kemudian diinterpretasikan menjadi kategori risiko pinjaman berdasarkan karakteristik *centroid* masing-masing *cluster*. Hasil pengelompokan digunakan untuk membedakan tingkat risiko pinjaman nasabah ke dalam kategori risiko rendah, risiko sedang, dan risiko tinggi sehingga lebih mudah digunakan sebagai informasi pendukung dalam pengambilan keputusan.

III. HASIL DAN PEMBAHASAN

Bagian ini memaparkan hasil pengolahan data menggunakan teknik *data mining* dengan algoritma K-Means *Clustering* pada data riwayat pinjaman nasabah Koperasi Simpan Pinjam Sri Hadi Dharma. Pembahasan difokuskan pada hasil pengelompokan data, karakteristik *cluster* yang terbentuk, serta interpretasi tingkat risiko pinjaman berdasarkan hasil *clustering*.

A. Dataset

Dataset yang digunakan merupakan riwayat pinjaman nasabah Koperasi Simpan Pinjam Sri Hadi Dharma pada periode 2021 hingga 2025. Hasil penyaringan diperoleh sebanyak 7.155 data dengan atribut ID Nasabah, Nama Nasabah, Tanggal Pencairan, Status Pekerjaan, Besar Pinjaman, Sisa Pinjaman, Angsuran Pokok, Jasa, Tunggalan dan Sisa Tenor yang digunakan sebagai dasar dalam analisis pengelompokan risiko pinjaman.

B. Hasil Preprocessing Data

Tahap awal penelitian dilakukan dengan melakukan *preprocessing* terhadap data riwayat pinjaman nasabah agar siap digunakan dalam proses *clustering*. Perbandingan data sebelum dan sesudah preprocessing ditunjukkan pada Tabel berikut.

Tabel 1. Data Sebelum *Preprocessing*

ID	Nama Nasabah	Status Pekerjaan	Besar Pinjaman	Sisa Pinjaman	Angsuran Pokok	Jasa	Tunggakan	Sisa Tenor
10	Soepandi/Hentri	Pensiunan	5400000	1800000	450000	108000	252000	4
26	Pitran Triyono	Umum	5040000	840000	420000	50400	67500	2
46	Soni	Umum	9000000	750000	750000	45000	270000	1
46	Soni	Umum	5040000	2100000	420000	126000	0	5
46	Soni	Umum	5040000	5040000	420000	302400	0	12

Tabel 2. Data Setelah *Preprocessing*

ID	Nama Nasabah	Status Pekerjaan	Besar Pinjaman	Sisa Pinjaman	Angsuran Pokok	Sisa Tenor	Tunggakan Maksimum
10	Soepandi/Hentri	Pensiunan	0.198	0.0842	0.2492	0.1765	0.0571
26	Pitran Triyono	Umum	0.1834	0.038	0.2308	0.0588	0.0153
46	Soni	Umum	0.3447	0.0338	0.4338	0	0.0612
59	Juari	Umum	0.1418	0.0139	0.1785	0	0.0336
60	Didik Subandi	Pensiunan	0.1002	0.0338	0.1262	0	0.1574

Tabel 1. Data Sebelum *Preprocessing* dan Tabel 2. Data Setelah *Preprocessing* menunjukkan adanya perubahan struktur dan nilai data setelah dilakukan tahapan preprocessing secara bertahap. Proses diawali dengan *data cleaning* untuk menghapus data yang tidak lengkap, duplikat, serta nilai yang tidak relevan dari total 7.155 data awal. Selanjutnya dilakukan perhitungan tunggakan maksimum pada setiap ID nasabah dengan mengambil nilai tunggakan tertinggi dari seluruh riwayat pinjaman sebagai representasi tingkat risiko terburuk. Tahap berikutnya adalah *data reduction*, di mana data disederhanakan menjadi 471 data unik, sehingga setiap nasabah hanya diwakili oleh satu *record* berdasarkan tenor terkecil untuk menghindari pengulangan data yang dapat memengaruhi hasil *clustering*. Setelah itu, dilakukan normalisasi menggunakan metode *Min-Max Normalization* untuk menyamakan skala atribut seperti besar pinjaman, sisa pinjaman, angsuran pokok, sisa tenor, dan tunggakan maksimum ke rentang 0-1. Proses ini membuat data lebih seimbang sehingga tidak ada atribut yang mendominasi dalam proses pengelompokan.

C. Hasil *Clustering* Menggunakan K-Means

Proses *clustering* dilakukan dengan menetapkan jumlah *cluster* sebesar $K = 3$ untuk membentuk tiga kategori utama yang merepresentasikan tingkat risiko pinjaman nasabah. Pada proses ini digunakan *random state* 42 untuk menjaga konsistensi hasil *clustering* dan memastikan hasil tetap stabil meskipun inisialisasi *centroid* dilakukan secara acak.

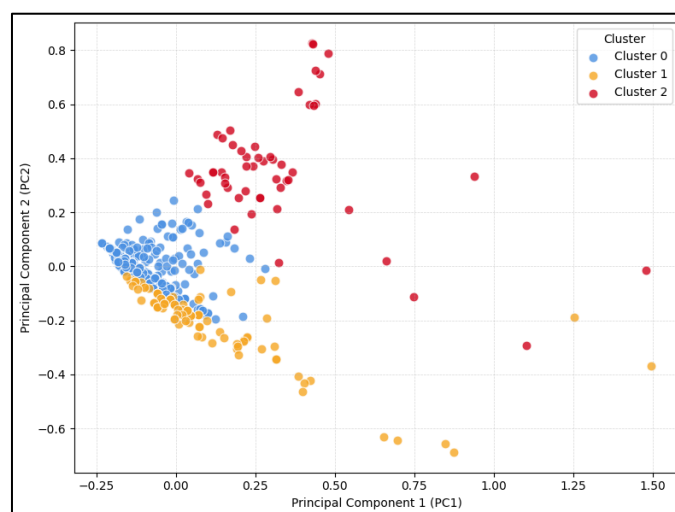
1. Distribusi Jumlah Nasabah Setiap *Cluster*

Tabel 3. Distribusi Jumlah Nasabah per *Cluster*

<i>Cluster</i>	Jumlah Nasabah
<i>Cluster</i> 0	323
<i>Cluster</i> 1	95
<i>Cluster</i> 2	53

Tabel 3. Distribusi Jumlah Nasabah per *Cluster* menunjukkan hasil pengelompokan 471 data nasabah ke dalam tiga *cluster* berbeda. Hasil tersebut memperlihatkan bahwa sebagian besar nasabah berada pada *Cluster* 0 dengan jumlah 323 orang, sedangkan sisanya terbagi ke dalam *Cluster* 1 sebanyak 95 orang dan *Cluster* 2 sebanyak 53 orang, yang masing-masing merepresentasikan karakteristik pinjaman yang berbeda antar kelompok.

2. Sebaran Data Berdasarkan *Cluster*



Gambar 2. Sebaran Data Berdasarkan *Cluster*

Gambar 2. Sebaran Data Berdasarkan *Cluster* memperlihatkan hasil visualisasi pengelompokan data yang dipetakan menggunakan *Principal Component Analysis* (PCA). Visualisasi ini menunjukkan adanya pemisahan antar *cluster* yang cukup jelas, di mana setiap titik merepresentasikan satu nasabah dengan karakteristik pinjaman yang berbeda. Komponen PC1 digunakan untuk menangkap variasi utama data, sedangkan PC2 merepresentasikan variasi tambahan sehingga pola sebaran antar *cluster* dapat terlihat lebih terstruktur dan mudah diamati.

3. Nilai Maksimum Fitur Setiap *Cluster*

Tabel 4. Nilai Maksimum Fitur pada Setiap *Cluster*

<i>Cluster</i>	Besar Pinjaman	Sisa Pinjaman	Angsuran Pokok	Sisa Tenor	Tunggakan Maksimum
0	9000000	1500000	750000	6	1275000
1	25080000	20900000	1670000	18	3660700
2	20040000	11690000	1670000	10	4414800

Tabel 4. Nilai Maksimum Fitur pada Setiap *Cluster* menggambarkan karakteristik utama yang membedakan setiap kelompok berdasarkan pola pinjaman nasabah. Karakteristik dari masing-masing *cluster* dapat dijelaskan sebagai berikut:

- Cluster 0*: Menunjukkan kelompok dengan kondisi keuangan paling stabil, ditandai oleh besar pinjaman mencapai 9.000.000 dan sisa pinjaman sebesar 1.500.000 yang merupakan nilai terendah. Nasabah pada kelompok ini umumnya berada pada tahap akhir pelunasan dengan sisa tenor hingga 6, sehingga beban angsuran pokok yang relatif kecil sebanyak 750.000 sejalan dengan rendahnya tunggakan maksimum sebesar 1.275.000.
- Cluster 1*: Merepresentasikan kelompok dengan tingkat paparan pinjaman paling tinggi, yaitu besar pinjaman mencapai 25.080.000 dan sisa pinjaman sebesar 20.900.000. Pemberian sisa tenor hingga 18 bulan mencerminkan adanya penyesuaian jangka waktu pembayaran terhadap besarnya pinjaman yang diterima. Kelompok ini memiliki angsuran pokok terbesar sebanyak 1.670.000 serta mencatat tunggakan maksimum mencapai 3.660.700.
- Cluster 2*: Menunjukkan kelompok dengan kondisi pembayaran yang paling berat, ditandai oleh tunggakan maksimum tertinggi yaitu mencapai 4.414.800. Hal ini terjadi meskipun besar pinjaman hanya sebesar 20.040.000 dan sisa pinjaman 11.690.000 yang lebih rendah dibanding *cluster 1*. Nasabah pada kelompok ini memiliki sisa tenor hingga 10 dengan angsuran pokok sebanyak 1.670.000, sehingga memperbesar tekanan pembayaran.

Pemetaan ini memberikan gambaran awkarakteristik masing-masing kelompok yang terbentuk dari proses *clustering* menggunakan K-Means.

4. Hasil Uji Stabilitas Menggunakan *Adjusted Rand Index* (ARI)

Tabel 5. Hasil Uji Stabilitas Menggunakan *Adjusted Rand Index* (ARI)

Perbandingan <i>Random State</i>	Nilai ARI
<i>Random State 1 vs 10</i>	1.0
<i>Random State 10 vs 20</i>	0.99
<i>Random State 20 vs 30</i>	1.0
<i>Random State 30 vs 42</i>	0.92

Tabel 5 menunjukkan hasil pengujian stabilitas *cluster* melalui beberapa variasi *random state* yang digunakan untuk melihat konsistensi hasil pengelompokan. Nilai ARI yang diperoleh

berada pada rentang sangat tinggi dan mendekati 1.0, yang menandakan bahwa hasil *clustering* tidak mengalami perubahan signifikan meskipun inisialisasi *centroid* dilakukan secara acak. Kondisi ini menunjukkan bahwa pola pengelompokan yang dihasilkan oleh K-Means bersifat konsisten, sehingga hasil *cluster* dapat digunakan sebagai dasar interpretasi karakteristik masing-masing kelompok.

D. Interpretasi Hasil *Cluster*

Tahap ini berfungsi untuk mengubah hasil *clustering* menjadi data kategorikal yang lebih mudah diinterpretasikan. Penentuan label setiap *cluster* didasarkan pada karakteristik dominan masing-masing kelompok agar perbedaan pola data lebih terstruktur.

1. Korelasi Fitur Terhadap *Cluster*



Gambar 3. Korelasi Fitur Terhadap *Cluster*

Gambar 3. Korelasi Fitur Terhadap *Cluster* menunjukkan tingkat kontribusi masing-masing atribut dalam pembentukan kelompok hasil *clustering*. Berdasarkan hasil visualisasi, Tunggalan Maksimum memiliki nilai korelasi tertinggi sebesar 0.71 yang menunjukkan bahwa atribut tersebut paling dominan dalam membedakan karakteristik antar *cluster*. Selanjutnya, Sisa Tenor dengan nilai 0.35 dan Sisa Pinjaman sebesar 0.33 juga memberikan pengaruh yang cukup signifikan dalam memperjelas pemisahan kelompok. Sementara itu, Besar Pinjaman sebesar 0.26 dan Angsuran Pokok sebesar 0.25 memiliki kontribusi yang lebih rendah namun tetap berperan dalam melengkapi pola pembentukan *cluster* secara keseluruhan.

2. Pelabelan *Cluster*

Tabel 6. Pelabelan *Cluster*

<i>Cluster</i>	Label Risiko
<i>Cluster</i> 0	Rendah
<i>Cluster</i> 1	Sedang
<i>Cluster</i> 2	Tinggi

Tabel 6. Pelabelan *Cluster* menunjukkan hasil pemetaan *cluster* ke dalam bentuk kategori yang lebih mudah dipahami. Penentuan label dilakukan dengan mempertimbangkan keterkaitan seluruh fitur pada masing-masing kelompok, dengan Tunggalan Maksimum sebagai indikator paling berpengaruh dalam membedakan pola antar *cluster*. Hasil ini membuat setiap *cluster* memiliki representasi kategori yang lebih jelas dan terstruktur.

IV. KESIMPULAN

Penelitian ini berhasil mencapai tujuan untuk mengelompokkan risiko pinjaman nasabah pada Koperasi Simpan Pinjam Sri Hadi Dharma menggunakan algoritma *K-Means Clustering*. Metode yang diterapkan mampu menghasilkan segmentasi nasabah berdasarkan karakteristik pinjaman yang dimiliki sehingga memberikan gambaran risiko yang lebih objektif, terstruktur, dan mudah diinterpretasikan dibandingkan penilaian secara manual. Hasil penelitian ini menunjukkan bahwa penerapan teknik *data mining* melalui *K-Means Clustering* dapat mendukung proses analisis risiko pinjaman serta memberikan informasi yang bermanfaat sebagai dasar pengambilan keputusan dalam pengelolaan pinjaman nasabah. Meskipun penelitian ini masih terbatas pada data dan atribut yang digunakan, hasil pengelompokan yang diperoleh berpotensi untuk dikembangkan ke dalam sistem operasional koperasi agar dapat dimanfaatkan secara lebih optimal dalam kegiatan pengelolaan pinjaman.

DAFTAR PUSTAKA

- [1] L. Hasan and R. Delzy Perkasa, "PERAN KOPERASI SIMPAN PINJAM DALAM MEMBERDAYAKAN EKONOMI MASYARAKAT," *Jurnal Genta Mulia*, vol. 14, no. 1, Jan. 2023, doi: 10.61290/GM.V14I1.687.
- [2] D. P. Ompusunggu, D. R. I. Sutrisno, and A. Hukom, "KONSISTENSI DAN EFEKTIVITAS PERAN LEMBAGA KEUANGAN NON BANK (KOPERASI SIMPAN PINJAM) SEBAGAI PENGGERAK PEREKONOMIAN INDONESIA," *Jurnal Cahaya Mandalika ISSN 2721-4796 (online)*, vol. 4, no. 1, pp. 689–696, Feb. 2023, doi: 10.36312/JCM.V4I1.1449.
- [3] W. Wahyudin and R. Purnamasari, "Analisis Risiko Kredit Bermasalah terhadap Return On Equity (ROE)," *Coopetition : Jurnal Ilmiah Manajemen*, vol. 15, no. 2, pp. 305–316, Jun. 2024, doi: 10.32670/COOPETITION.V15I2.4392.
- [4] Roviani, D. Supriadi, and I. Dzulfiqar Iskandar, "PREDICTION OF COOPERATIVE LOAN FEASIBILITY USING THE K-NEAREST NEIGHBOR ALGORITHM," *Jurnal Pilar Nusa Mandiri*, vol. 17, no. 1, pp. 39–46, Mar. 2021, doi: 10.33480/PILAR.V17I1.2183.
- [5] A. Rifqi and R. T. Aldisa, "Penerapan Data Mining dalam Implementasi Algoritma K-Means Clustering untuk Pelanggan Potensial pada Koperasi Simpan Pinjam," *Building of Informatics, Technology and Science (BITS)*, vol. 5, no. 2, pp. 486–497, Sep. 2023, doi: 10.47065/BITS.V5I2.4278.
- [6] Y. O. Tamba and N. Ekawati, "KLAUSTERISASI PINJAMAN NASABAH KOPERASI MENGGUNAKAN ALGORITMA K-MEANS," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 6, pp. 11107–11114, Nov. 2024, doi: 10.36040/JATI.V8I6.11195.
- [7] M. Alfarisi and M. S. Ramadhan, "IMPLEMENTASI K-MEANS CLUSTERING UNTUK MENGIDENTIFIKASI CALON NASABAH POTENSIAL PADA LEMBAGA KEUANGAN KECIL DAN MENENGAH," *JOURNAL OF SCIENCE AND SOCIAL RESEARCH*, vol. 8, no. 3, pp. 3807–3813, Aug. 2025, doi: 10.54314/JSSR.V8I3.4094.
- [8] F. D. S. Alhamdani, A. A. Dianti, and Y. Azhar, "Segmentasi Pelanggan Berdasarkan Perilaku Penggunaan Kartu Kredit Menggunakan Metode K-Means Clustering," *JISKA*

- (Jurnal Informatika Sunan Kalijaga), vol. 6, no. 2, pp. 70–77, May 2021, doi: 10.14421/jiska.2021.6.2.70-77.
- [9] A. Fikri, B. F. Hutabarat, and U. Khaira, “Komparasi K-Means Clustering Dan Complete Linkage Dalam Pengelompokan Penyaluran Pinjaman Financial Technology,” Jurnal Ilmiah Media Sisfo, vol. 17, no. 2, pp. 228–239, Oct. 2023, doi: 10.33998/MEDIASISFO.2023.17.2.1373.
- [10] A. Fikri, B. F. Hutabarat, and U. Khaira, “Komparasi K-Means Clustering Dan Complete Linkage Dalam Pengelompokan Penyaluran Pinjaman Financial Technology,” Jurnal Ilmiah Media Sisfo, vol. 17, no. 2, pp. 228–239, Oct. 2023, doi: 10.33998/MEDIASISFO.2023.17.2.1373.
- [11] A. Ramadhanu, S. Defit, and S. W. Wahhab Kareem, “Hybrid Data Mining with the Combination of K-Means Algorithm and C4.5 to Predict Student Achievement,” International Journal Of Artificial Intelligence Research, vol. 5, 2021, Accessed: Apr. 29, 2026. [Online]. Available: https://ijair.id/index.php/ijair/article/view/225?utm_source=chatgpt.com
- [12] H. Manurung, “PENGELOMPOKAN UMKM KOTA BINJAI MENGGUNAKAN METODE CLUSTERING K-MEANS UNTUK MENGIDENTIFIKASI POLA PERKEMBANGAN BISNIS,” Jurnal Informatika Kaputama (JIK), vol. 8, no. 2, pp. 93–99, Jul. 2024, doi: 10.59697/JIK.V8I2.844.
- [13] S. Ocktavia and W. T. Atmojo, “ANALISIS SEGMENTASI PELANGGAN PEMBIAYAAN BERDASARKAN DEMOGRAFI UNTUK MEMPREDIKSI TINGKAT KREDIT MENGGUNAKAN ALGORITMA K-MEANS,” Jurnal Informatika Teknologi dan Sains (Jinteks), vol. 7, no. 1, pp. 394–402, Mar. 2025, doi: 10.51401/JINTEKS.V7I1.5582.
- [14] M. Rifqi Mubarak, A. Tri, J. Harjanto, and R. Renaldy, “PENINGKATAN PERFORMA DBSCAN DENGAN REDUKSI DIMENSI PRINCIPAL COMPONENT ANALYSIS (PCA) DALAM KLASTERISASI TINGKAT KEMISKINAN DI INDONESIA,” Jurnal Informatika Teknologi dan Sains (Jinteks), vol. 7, no. 3, pp. 1176–1184, Aug. 2025, doi: 10.51401/JINTEKS.V7I3.6129.