

Implementasi Paralel *K-Means Clustering* dan *Decision Tree* untuk Deteksi Risiko Kredit Macet

Halili Maar¹, Elin Nurjanah², Gina Suraya³, Dea Khoirunnisa⁴

^{1,2,3,4} Teknik Informatika, Ilmu Komputer, Universitas Pamulang

¹dosen02957@unpam.ac.id*, ²elinnurjanah65807@gmail.com, ³surayagina@gmail.com, ⁴deakhns101@email.com

Abstract

Non-Performing Loan (NPL) pose a major threat to the stability of financial institutions. This study implements parallel K-Means Clustering and Decision Tree methods in RapidMiner 2026.0.5 to identify potential credit default risks using the Credit Risk Dataset. The dataset consists of 32,576 borrower records with 12 attributes after filtering. Preprocessing includes handling missing values, dummy coding, converting loan status into a binary variable, and applying Z-Transformation normalization. The Multiply operator enables K-Means (k=3) and Decision Tree to run simultaneously. The Clustering results show that Cluster 0 (33.10%) is a low-risk group dominated by homeowners with mortgages and an A loan rating; Cluster 1 (32.07%) is a medium-risk group dominated by tenant borrowers with a loan rating of B, and Cluster 2 (34.83%) is a high-risk group dominated by a history of delinquency, loan ratings of C–G, and the highest interest rates. The Decision Tree model achieved an Accuracy of 89.26%, Precision of 76.30%, Recall of 73.05%, and an F-Measure of 74.58%, making it effective as an early detection system for nonperforming loans.

Keywords: k-means clustering, decision tree, non-performing loan, gain ratio, rapidminer

Abstrak

Kredit macet atau *Non-Performing Loan (NPL)* merupakan ancaman yang dapat mengganggu stabilitas institusi keuangan. Penelitian ini mengimplementasikan *K-Means Clustering* dan *Decision Tree* secara paralel pada *RapidMiner 2026.0.5* untuk mendeteksi risiko kredit macet dengan menggunakan *Credit Risk Dataset*. Dataset terdiri atas 32.576 data debitur dengan 12 atribut yang mencakup informasi demografis, karakteristik pinjaman, dan riwayat kredit. Tahap pra-pemrosesan meliputi penanganan *missing values*, *dummy coding*, konversi *loan status* menjadi variabel *binomial*, serta normalisasi *Z-Transformation*. Operator *Multiply* digunakan untuk menjalankan *K-Means (k=3)* dan *Decision Tree* secara bersamaan. Hasil *Clustering* menghasilkan Cluster 0 (33,10%) sebagai kelompok risiko rendah dengan debitur mapan berhipotek dan *loan grade* A dominan, Cluster 1 (32,07%) sebagai kelompok risiko menengah dengan debitur penyewa dan *loan grade* B dominan, dan Cluster 2 (34,83%) sebagai kelompok risiko tinggi yang didominasi riwayat gagal bayar, *loan grade* C–G, serta suku bunga tertinggi. *Decision Tree* memperoleh akurasi 89.26%, *Precision* 76.30%, *Recall* 73.05% dan *F-Measure* 74.58%, sehingga efektif sebagai sistem deteksi dini kredit macet.

Kata kunci: *k-means clustering, decision tree, kredit macet, gain ratio, rapidminer*

©This work is licensed under a Creative Commons Attribution -ShareAlike 4.0 International License

1. Pendahuluan

Perbankan berperan penting dalam perekonomian nasional, dengan kredit sebagai produk utama sumber pendapatan. Namun, risiko kredit macet atau *Non-Performing Loan (NPL)* tetap menjadi tantangan yang tidak dapat diabaikan. Tingginya NPL mengancam profitabilitas serta stabilitas sistem keuangan. Otoritas Jasa Keuangan (OJK) menetapkan batas maksimal rasio NPL sebesar 5% sebagai indikator kesehatan perbankan [1]. Deteksi risiko kredit sejak tahap awal sangat penting agar pihak manajemen dapat melakukan tindakan pencegahan sebelum kerugian besar terjadi.

Kemajuan dalam bidang *Machine Learning* membuka kemungkinan baru untuk menganalisis risiko kredit dengan akurasi dan efisiensi yang lebih tinggi. Dua teknik yang relevan meliputi *K-Means Clustering* sebagai metode segmentasi yang menggunakan pembelajaran tak terawasi, serta *Decision Tree* sebagai metode klasifikasi yang berbasis pembelajaran terawasi. Implementasi keduanya secara paralel dalam

satu alur proses memberikan analisis yang komprehensif, yaitu *insight* segmentasi debitur sekaligus prediksi status kredit individual yang akurat [2].

K-Means Clustering memungkinkan pengelompokan debitur berdasarkan kesamaan karakteristik tanpa memerlukan *label*, mengungkap segmen dengan profil risiko serupa [3]. Algoritma ini bekerja dengan cara memulai dari sejumlah *k* titik pusat (*centroid*), menempatkan setiap titik data pada *centroid* terdekat berdasarkan ukuran jarak, memperbarui posisi *centroid* dengan menghitung rata-rata anggota *Cluster*, dan mengulangi proses tersebut hingga tercapai konvergensi. Parameter *Mixed Euclidean Distance* digunakan untuk menangani data yang terdiri atas variabel numerik dan biner hasil *dummy encoding*. Pada penelitian ini, nilai *k=3* dipilih berdasarkan pertimbangan domain bisnis perbankan yang membagi portofolio kredit ke dalam tiga segmen risiko (rendah, menengah, dan tinggi), yang merupakan kategorisasi standar dalam manajemen portofolio kredit perbankan.

Untuk memastikan pemilihan $k=3$ secara kuantitatif, dilakukan evaluasi memakai *Elbow Method* serta *Silhouette Score* dalam rentang $k=2$ sampai $k=10$ dengan Python (*scikit-learn*). Hasil *Elbow Method* mengindikasikan bahwa kurva WCSS (*Within-Cluster Sum of Squares*) tidak menampilkan sudut tajam, melainkan menurun secara perlahan dari $k=2$ (WCSS sekitar 760.000) hingga $k=10$ (WCSS sekitar 520.000). Penurunan paling besar terjadi antara $k=2$ dan $k=3$ (sekitar 60.000), sedangkan penurunan setelah $k=3$ menjadi jauh lebih halus. Sementara itu, *Silhouette Score* memperlihatkan nilai tertinggi pada $k=2$ (0,158), diikuti $k=4$ (0,120), dan $k=3$ (0,115). Walaupun secara statistik $k=2$ dan $k=10$ memperoleh *Silhouette Score* yang lebih tinggi, $k=2$ tidak memberikan detail yang cukup untuk membedakan segmen risiko menengah dari rendah, dan $k=10$ tidak relevan secara bisnis karena menghasilkan terlalu banyak segmen yang sulit diinterpretasikan serta tidak sejalan dengan praktik pengelompokan risiko kredit perbankan yang biasanya memakai tiga hingga lima segmen. Oleh karena itu, $k=3$ dipilih sebagai kompromi optimal antara validitas statistik, kemudahan interpretasi bisnis, dan kesesuaian dengan standar industri keuangan. Parameter $\text{max runs}=10$ dipakai untuk menjamin sepuluh inisialisasi independen dijalankan guna memperoleh konfigurasi terbaik, serta $\text{max optimization steps}=100$ untuk membatasi jumlah iterasi per percobaan demi efisiensi komputasi.

Decision Tree membangun model klasifikasi sebagai pohon hierarkis yang *interpretable* dan mudah dipahami praktisi bisnis, dengan node internal merepresentasikan kondisi pengujian atribut, cabang merepresentasikan hasil pengujian, dan node daun merepresentasikan kelas prediksi. [4]. Kriteria *Gain Ratio* yang digunakan dalam penelitian ini merupakan penyempurnaan *Information Gain* yang mengoreksi bias terhadap atribut dengan banyak nilai unik, dengan cara membagi nilai *Information Gain* dengan *Split Information*. Mekanisme *pre-pruning* dengan parameter *minimal gain* sebesar 0,01, *minimal leaf* sebesar 2, *minimal size for split* sebesar 4, dan tiga alternatif percabangan menghentikan pertumbuhan pohon secara proaktif, sementara *post-pruning* berbasis tingkat keyakinan 0,1 menyederhanakan pohon setelah terbentuk. Kedua mekanisme tersebut bekerja secara sinergis untuk mencegah *overfitting* sekaligus mempertahankan akurasi generalisasi model. *Platform RapidMiner* menyertakan operator *Multiply* yang memungkinkan satu set data diproses secara paralel menjadi dua sub-alur terpisah, sehingga tahap praproses tetap konsisten pada kedua algoritma.

Penelitian terdahulu menunjukkan bahwa integrasi metode *Clustering* sebelum atau secara paralel dengan klasifikasi dapat meningkatkan kualitas analisis risiko [5]. Beberapa peneliti membuktikan bahwa *K-Means* efektif dalam segmentasi debitur [6], sementara peneliti lain mengonfirmasi bahwa *Decision Tree* merupakan *classifier* yang andal untuk prediksi status kredit [7].

Penelitian terdahulu juga membandingkan algoritma *machine learning* untuk *credit scoring* namun tidak menggunakan pendekatan paralel [8], penelitian lain mengimplementasikan *Decision Tree* C4.5 untuk kredit macet perbankan Indonesia tanpa integrasi *Clustering* [9], dan penelitian lain melakukan *benchmarking* komprehensif algoritma klasifikasi kredit namun terbatas pada pembelajaran terawasi saja [10]. Kajian komprehensif terhadap sepuluh algoritma *machine learning* menggunakan *dataset* kredit nyata berukuran lebih dari 2,5 juta observasi juga menemukan bahwa model-model *ensemble*, khususnya XGBoost, secara konsisten mengungguli algoritma tunggal seperti regresi logistik dan *Decision Tree* dalam klasifikasi risiko kredit [11], namun menegaskan bahwa tidak ada satu pun algoritma yang secara universal optimal karena performa terbaik bergantung pada karakteristik dataset, teknik penanganan ketidakseimbangan kelas, dan pemilihan metrik evaluasi yang digunakan. Namun, keberhasilan pemodelan risiko kredit tidak hanya ditentukan oleh pemilihan algoritma. Kajian literatur terbaru menunjukkan bahwa tahapan praproses data, pemilihan fitur, penanganan ketidakseimbangan kelas, serta optimasi parameter turut berperan penting dalam meningkatkan performa model *machine learning* [12]. Temuan-temuan tersebut memperkuat argumen bahwa pendekatan kombinasi atau paralel lebih adaptif dibandingkan algoritma tunggal, sekaligus menjadi dasar penelitian ini untuk mengeksplorasi pendekatan paralel *K-Means Clustering* dan *Decision Tree* yang saling melengkapi, yaitu *K-Means Clustering* memberikan segmentasi populasi berbasis karakteristik data, sementara *Decision Tree* tetap dipilih karena strukturnya yang mudah diinterpretasikan dibandingkan algoritma *black-box* lainnya sehingga sesuai untuk kebutuhan transparansi pengambilan keputusan kredit.

Meskipun demikian, kajian terhadap literatur yang ada mengungkap tiga kesenjangan penelitian yang belum terjawab. Pertama, sebagian besar penelitian menerapkan *K-Means Clustering* dan *Decision Tree* secara terpisah atau sekuensial, di mana hasil *Clustering* dijadikan fitur tambahan untuk klasifikasi sehingga kedua algoritma tidak dapat memanfaatkan data sumber yang identik secara bersamaan, dan bias dari tahap *preprocessing* yang berbeda berpotensi memengaruhi konsistensi hasil. Kedua, analisis *centroid K-Means* sebagai instrumen interpretasi profil risiko debitur secara kuantitatif umumnya disajikan secara terpisah dari evaluasi model klasifikasi padahal penyajian keduanya dalam satu kerangka analisis yang terpadu berpotensi menghasilkan informasi yang lebih lengkap bagi pengambil keputusan kredit, yaitu tidak hanya mengetahui prediksi status kredit individual tetapi juga memahami karakteristik kelompok risiko di mana debitur tersebut berada. Ketiga, implementasi pendekatan paralel menggunakan *platform visual workflow* yang *reproducible* seperti RapidMiner belum banyak didokumentasikan dalam literatur nasional, sehingga membatasi kemampuan peneliti lain untuk

mereplikasi dan mengembangkan eksperimen serupa [13]. Ketiga kesenjangan tersebut menjadi landasan penelitian ini untuk merancang alur proses paralel menggunakan operator *Multiply* di RapidMiner, yang memungkinkan *K-Means Clustering* dan *Decision Tree* berjalan secara simultan dari satu sumber data yang identik, sehingga konsistensi *preprocessing* terjamin dan hasil kedua algoritma dapat dibandingkan secara adil.

Berdasarkan kesenjangan tersebut, penelitian ini memiliki tiga tujuan utama: (1) merancang dan mengimplementasikan alur kerja paralel *K-Means Clustering* dan *Decision Tree* pada RapidMiner menggunakan operator *Multiply*, (2) menganalisis profil risiko debitur secara kuantitatif melalui interpretasi tabel *centroid K-Means*, dan (3) mengevaluasi kinerja prediksi model *Decision Tree* sebagai instrumen deteksi dini kredit macet. Kebaruan (*novelty*) utama penelitian ini terletak pada pengembangan kerangka *profiling* risiko kredit berbasis interpretasi kuantitatif tabel *centroid K-Means*, yang menghubungkan nilai Z-score atribut diskriminatif khususnya *default*, *loan_grade*, dan *interest_rate* dengan segmentasi risiko debitur secara terukur dan dapat diaudit oleh analisis kredit. Berbeda dari penelitian terdahulu yang umumnya menyajikan hasil *Clustering* hanya sebagai pengelompokan statistik tanpa interpretasi profil risiko yang mendalam [6], penelitian ini menghasilkan tiga segmen debitur dengan karakteristik kuantitatif yang spesifik dan dapat langsung diterjemahkan menjadi kebijakan kredit yang terdiferensiasi. Kontribusi pendukung dari penelitian ini adalah desain *framework* paralel menggunakan operator *Multiply* di RapidMiner yang memastikan kedua algoritma *K-Means* untuk segmentasi dan *Decision Tree* untuk prediksi individual memproses representasi data yang benar-benar identik dalam satu alur terintegrasi, sehingga konsistensi *preprocessing* terjamin dan hasil kedua algoritma dapat diinterpretasikan secara bersama-sama secara koheren.

2. Metode Penelitian

2.1. Pengumpulan Data

Penelitian menggunakan *Credit Risk Dataset* dari repositori OpenML. *Dataset* original 32.581 *record*, setelah filter *person_age* ≤ 100 diperoleh 32.576 *record* valid (5 outlier usia dihapus) dengan 12 atribut (Tabel 1).

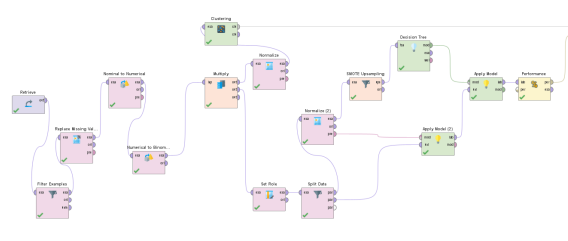
Tabel 1. Deskripsi Atribut Dataset

No	Atribut	Tipe	Keterangan
1	<i>person_age</i>	Integer	Usia debitur (tahun)
2	<i>income</i>	Integer	Pendapatan tahunan (USD)
3	<i>home_ownership</i>	Nominal	Status kepemilikan rumah
4	<i>employment_length</i>	Integer	Lama bekerja (tahun)
5	<i>loan_intent</i>	Nominal	Tujuan pinjaman

6	<i>loan_grade</i>	Nominal	Peringkat pinjaman (A–G)
7	<i>loan_amount</i>	Integer	Jumlah pinjaman (USD)
8	<i>interest_rate</i>	Real	Suku bunga (%)
9	<i>loan_status</i>	Binominal	Target: 0=Lancar, 1=Macet
10	<i>loan_percent_income</i>	Real	Rasio cicilan/pendapatan
11	<i>default</i>	Nominal	Riwayat gagal bayar (Y/N)
12	<i>credit_history_length</i>	Integer	Panjang riwayat kredit

2.2. Desain Alur Proses Paralel

Alur proses pada RapidMiner 2026.0.5 dirancang dalam lima fase: (1) *Input & Filtering*, (2) *Praproses Sekuensial*, (3) *Percabangan Paralel via Multiply*, (4) *Penyeimbangan Kelas via SMOTE Upsampling* pada data latih, dan (5) *Analisis Ganda*. Gambar 1 menunjukkan visualisasi alur proses lengkap sebagaimana diimplementasikan.



Gambar 1. Alur Proses Paralel *K-Means* dan *Decision Tree* di RapidMiner

Tabel 2. Konfigurasi Operator RapidMiner

Operator	Parameter	Nilai
<i>Filter</i>	Kondisi	<i>person_age</i> ≤ 100
<i>Examples Replace Missing Val.</i>	Default	Average
<i>Nominal to Numerical</i>	Atribut	<i>default</i> , <i>home_ownership</i> , <i>loan_intent</i> , <i>loan_grade</i>
<i>Nominal to Numerical</i>	Coding Type	Dummy coding
<i>Nominal to Numerical</i>	Atribut	<i>loan_status</i> (single)
<i>Multiply</i>	Output	Out1→ <i>Clustering</i> ; Out2→ <i>Set Role</i>
<i>Normalize</i>	Method	<i>Z-Transformation</i> (Cabang <i>K-Means</i>)
<i>K-Means (Clustering)</i>	<i>k / Max runs</i>	3 / 10
<i>K-Means (Clustering)</i>	<i>Mixed Measure</i>	<i>MixedEuclideanDistance</i>
<i>K-Means (Clustering)</i>	<i>Max Opt. Steps</i>	100
<i>Set Role</i>	<i>loan_status</i>	label
<i>Split Data</i>	<i>Ratio / Sampling</i>	0,7:0,3 / <i>Shuffled</i>
<i>Normalize</i>	Method	<i>Z-Transformation</i> (Data latih)
<i>SMOTE Upsampling</i>	<i>number_of_neighbours</i>	5

SMOTE Upsamplin g	equalize_classes	true
SMOTE Upsamplin g	auto_detect_minority_class	true
SMOTE Upsamplin g	append_to_original	true
Decision Tree	Criterion	Gain Ratio
Decision Tree	Max depth	10
Decision Tree	Pruning/Confidence	True / 0,1
Decision Tree	Min Gain/Leaf/Split	0,01 / 2 / 4
Decision Tree	Prepruning Alt.	3
Apply Model	Default	Preprocessing Model
Apply Model	Default	Decision Tree Model
Performan ce	Main Criterion	First (Positive: true)
Performan ce	Metrik	Accuracy, Precision, Recall, F-Measure

2.3. Praproses Data

Proses praproses terdiri dari enam langkah: (1) *Replace Missing values* mengganti nilai yang hilang dengan nilai rata-rata; (2) *Filter Examples* menghapus lima data outlier dengan usia lebih dari 100 tahun; (3) *Nominal to Numerical* menggunakan *dummy coding* mengubah empat variabel kategorikal yaitu *default* (2 kelas), *home_ownership* (4 kelas), *loan_intent* (6 kelas), dan *loan_grade* (7 kelas) sehingga menjadi 19 kolom biner, bersama delapan atribut numerik yang tidak mengalami transformasi (*person_age*, *income*, *employment_length*, *loan_amount*, *interest_rate*, *loan_status*, *loan_percent income*, dan *credit history_length*), dataset kemudian menghasilkan total 27 atribut; (4) *Numerical to Binominal* mengubah *loan_status* menjadi variabel binomial (*false/true*) yang berfungsi sebagai target klasifikasi; (5) Setelah proses *data preprocessing* yang meliputi *Filter Examples*, *Replace Missing Values*, *Nominal to Numerical*, dan *Numerical to Binominal*, *workflow* dipisahkan menjadi dua cabang menggunakan operator *Multiply*. (6) Pada cabang pertama, data dinormalisasi menggunakan metode *Z-Transformation* sebelum dilakukan proses *K-Means Clustering*. Normalisasi pada cabang ini diterapkan terhadap seluruh dataset karena *K-Means* merupakan metode *unsupervised learning* yang tidak menggunakan pembagian data latih dan data uji, sehingga tidak menimbulkan risiko data *leakage*. (7) Pada cabang kedua, data terlebih dahulu diberi label menggunakan operator *Set Role*, (8) kemudian dibagi menjadi 70% data *training* dan 30% data uji menggunakan operator *Split Data*. (9) Selanjutnya, proses *Normalize* di-fit hanya pada data latih untuk memperoleh parameter normalisasi berupa *mean* dan *standard deviation*. (10) Parameter normalisasi tersebut diterapkan pada data uji menggunakan operator *Apply Model* yang memanfaatkan

preprocessing model hasil *Normalize* sehingga data uji mengalami transformasi menggunakan parameter yang diperoleh dari data latih. Pendekatan ini dilakukan untuk mencegah terjadinya *data leakage*, yaitu kondisi ketika informasi statistik dari data uji memengaruhi proses pelatihan model. (11) Mengingat distribusi kelas *loan_status* yang tidak seimbang (78,18% lancar vs 21,82% macet) sebagaimana teridentifikasi pada statistik deskriptif, penelitian ini menerapkan teknik *Synthetic Minority Oversampling Technique* (SMOTE) pada partisi data latih menggunakan operator *SMOTE Upsampling* di RapidMiner (*number_of_neighbours=5*, *equalize_classes=true*, *auto_detect_minority_class=true*). SMOTE diterapkan setelah proses *Split Data* sehingga hanya data latih (70%) yang mengalami penyeimbangan, sedangkan data uji (30%) tetap mencerminkan distribusi asli untuk memastikan evaluasi performa dilakukan pada kondisi nyata tanpa kontaminasi data sintetis sebelum dilakukan proses pembentukan model *Decision Tree*.

2.4. Metrik Evaluasi

Evaluasi *Decision Tree* menggunakan:

$$Accuracy = \frac{(TP+TN)}{(TP+TN+FP+FN)} \quad (1)$$

$$Precision = \frac{TP}{(TP+FP)} \quad (2)$$

$$Recall = \frac{TP}{(TP+FN)} \quad (3)$$

$$F - Measure = 2 \frac{(P \times R)}{(P+R)} \quad (4)$$

Positive class ditetapkan 'true' (kredit macet) karena merupakan kelas yang paling kritis dideteksi dalam konteks manajemen risiko [9]. Sejalan dengan penetapan ini, *Recall* kelas macet diprioritaskan sebagai metrik evaluasi utama dalam penelitian ini, mengingat biaya kegagalan mendeteksi debitur yang berpotensi macet (*false negative*) berupa potensi kerugian kredit tak tertagih dinilai lebih besar dibandingkan biaya verifikasi tambahan terhadap debitur yang sebenarnya lancar (*false positive*).

3. Hasil dan Pembahasan

3.1. Statistik Deskriptif

Tabel 3. Statistik Deskriptif Dataset (n=32.576)

Atribut	Mean	Std Dev	Min	Max
<i>person_age</i>	27,72	6,20	20	94
<i>income</i>	65.882	52.535	4.000	2.039.784
<i>employment_length</i>	4,79	4,14	0	123
<i>loan_amount</i>	9.589	6.322	500	35.000
<i>interest_rate (%)</i>	11,01	3,24	5,42	23,22
<i>loan_percent income</i>	0,17	0,11	0,00	0,83
<i>credit history_length</i>	5,80	4,05	2	30

Sebaran *loan_status* mengungkap adanya 25.468 pinjaman yang masih berjalan (78,18%) dan 7.108 pinjaman yang macet (21,82%), yang menunjukkan adanya ketidakseimbangan kelas yang harus

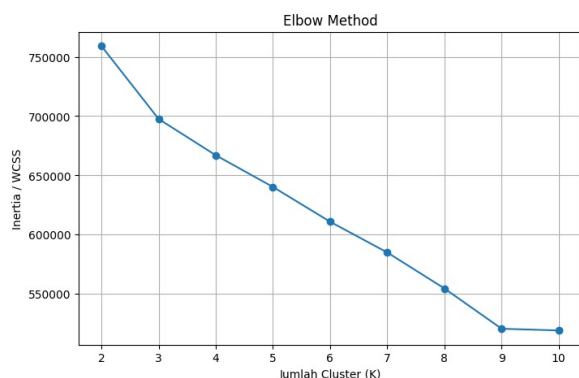
dipertimbangkan dalam penafsiran *Recall* model. Sebaran kepemilikan rumah didominasi oleh kategori RENT (50,48%) dan MORTGAGE (41,26%), sementara catatan gagal bayar (*default=Y*) tercatat sebanyak 5.745 *record* (17,64%).

Rata-rata nilai *loan_percent_income* sebesar 0,17 menunjukkan bahwa nilai pinjaman yang diajukan setara dengan sekitar 17% dari pendapatan debitur. Berdasarkan statistik deskriptif tersebut, dapat diinterpretasikan bahwa proporsi pinjaman terhadap pendapatan pada sebagian besar debitur dalam dataset relatif rendah. Namun nilai maksimum 0,83 mengindikasikan terdapat debitur yang mengalokasikan hingga 83% pendapatannya untuk cicilan, yang merupakan indikator risiko gagal bayar yang sangat tinggi.

Variabilitas tinggi pada atribut *income* (standar deviasi 52.535 jauh melebihi rata-rata 65.882) mengindikasikan distribusi pendapatan yang sangat tidak merata di antara debitur. Kondisi ini relevan dengan proses *K-Means Clustering* karena variabilitas tinggi berpotensi menciptakan *Cluster* yang tidak seimbang jika normalisasi tidak dilakukan dengan tepat [15]. Penerapan *Z-Transformation* dalam tahap preprocessing menjadi langkah krusial untuk memastikan seluruh atribut berkontribusi secara proporsional dalam penghitungan jarak Euclidean pada proses *Clustering*. Ketidakseimbangan kelas *loan_status* (78,18% vs 21,82%) perlu dicermati dalam interpretasi metrik evaluasi karena distribusi kelas yang tidak seimbang dapat menyebabkan model cenderung bias terhadap kelas mayoritas. Oleh karena itu, diperlukan teknik penyeimbangan data, seperti SMOTE, untuk meningkatkan kemampuan model dalam mengenali kelas minoritas [16].

3.2. Hasil *K-Means Clustering*

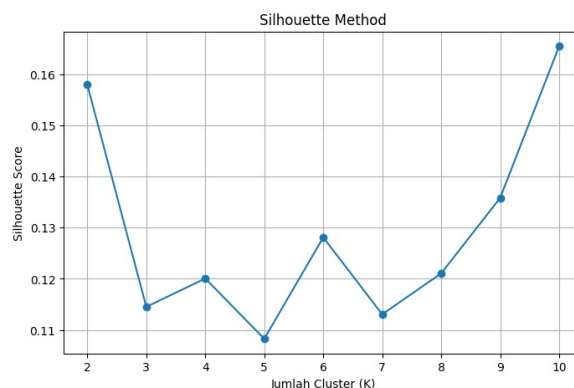
Sebelum melaksanakan eksperimen utama, dilakukan penilaian kuantitatif untuk menentukan $k=3$ dengan memakai *Elbow Method* serta *Silhouette Score* pada rentang $k=2-10$.



Gambar 2. *Elbow Method*

Pada *Elbow Method* (Gambar 2), kurva WCSS menurun secara bertahap tanpa terlihat sudut tajam. Penurunan paling signifikan terjadi antara $k=2$ dan $k=3$ dengan selisih kira-kira 60.000 (dari sekitar

760.000 menjadi 700.000), sementara penurunan pada $k=4$ hingga $k=10$ menjadi lebih datar dan hampir konstan. Hal ini menyiratkan bahwa menambah kluster di atas $k=3$ tidak memberi peningkatan kompaksi yang berarti secara proporsional.



Gambar 3. *Silhouette Score*

Pada *Silhouette Score* (Gambar 3), $k=2$ memperoleh nilai tertinggi (0,158), diikuti $k=4$ (0,120) dan $k=3$ (0,115). Walaupun $k=2$ secara statistik lebih baik, pilihan ini ditolak karena hanya menghasilkan dua segmen yang tidak dapat memisahkan debitur berisiko menengah dari yang berisiko rendah, perbedaan yang penting dalam kebijakan kredit bank. Sementara itu, $k=10$ menghasilkan *Silhouette Score* yang juga tinggi (0,165) namun tidak praktis bagi bisnis karena menghasilkan sepuluh segmen yang terlalu detail untuk dianalisis dan diterapkan dalam kebijakan kredit. Dengan menimbang keseimbangan antara validitas statistik, kemudahan interpretasi, serta relevansi bagi sektor perbankan yang biasanya memakai tiga sampai lima segmen risiko, $k=3$ diputuskan sebagai nilai optimal dalam studi ini. Algoritma *K-Means* ($k=3$, *seed* acak = 2001) membentuk tiga kluster seperti yang ditampilkan pada Gambar 2. Representasi pohon graf menunjukkan bahwa set akar terbagi menjadi tiga sub-kluster dengan ukuran proporsional yang berbeda.



Gambar 4. *Graph Tree Cluster Model*

Tabel 4. Distribusi Hasil *K-Means Clustering*

Cluster	Jumlah Item	Persentase	Profil Risiko
Cluster 0	10.782	33,10%	Risiko Rendah
Cluster 1	10.448	32,07%	Risiko Menengah
Cluster 2	11.346	34,83%	Risiko Tinggi
Total	32.576	100%	-

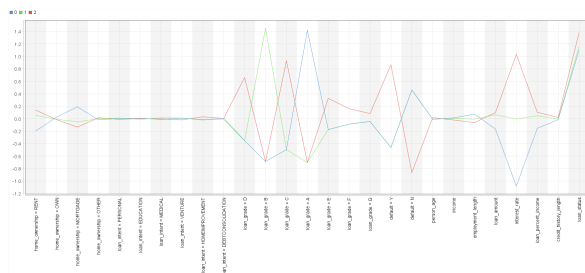
Penetapan profil risiko tiap *Cluster* didasarkan pada analisis *centroid* table pada Gambar 3 dan Tabel 5 berikut. Analisis ini merupakan kontribusi utama interpretasi *K-Means* dalam penelitian ini.

Attribute	cluster_0	cluster_1	cluster_2
home_ownership = RENT	-0.199	0.054	0.139
home_ownership = OWN	0.024	-0.009	-0.015
home_ownership = MORTGAGE	0.190	-0.050	-0.135
home_ownership = OTHER	-0.015	-0.001	0.015
loan_grade = PERSONAL	0.001	0.011	-0.011
loan_grade = EDUCATION	0.009	-0.011	0.001
loan_grade = MEDICAL	-0.012	0.014	-0.002
loan_grade = VENTURE	0.010	0.009	-0.018
loan_grade = HOMEIMPROVEMENT	-0.020	-0.013	0.030
loan_grade = DEBTCONSOLIDATION	0.008	-0.013	0.005
loan_grade = D	-0.353	-0.354	0.662
loan_grade = B	-0.687	1.455	-0.687
loan_grade = C	-0.497	-0.497	0.930
loan_grade = A	1.421	-0.703	-0.703
loan_grade = E	-0.174	-0.175	0.326
loan_grade = F	-0.086	-0.086	0.162
loan_grade = G	-0.044	-0.044	0.083
default = Y	-0.463	-0.463	0.866
default = N	0.463	0.463	-0.866
person_age	-0.009	-0.010	0.017
income	0.013	0.009	-0.020
employment_length	0.074	-0.008	-0.063
loan_amount	-0.105	0.094	0.099
interest_rate	-1.084	-0.005	1.035
loan_percent_income	-0.155	0.048	0.103
credit_history_length	-0.015	-0.005	0.019
loan_status	1.100	1.163	1.382

Gambar 5. Centroid Table K-Means Clustering (Nilai Z-score per Atribut per Cluster)

Tabel 5. Nilai Centroid Kunci per Cluster (Atribut Diskriminatif)

Atribut	Cluster 0	Cluster 1	Cluster 2
home_ownership=RENT	-0,199	0,054	0,139
home_ownership=MORTGAGE	0,190	-0,050	-0,135
loan_grade=A	1,421	-0,703	-0,703
loan_grade=B	-0,687	1,455	-0,687
loan_grade=C	-0,497	-0,497	0,930
loan_grade=D	-0,353	-0,354	0,662
loan_grade=E	-0,174	-0,175	0,326
loan_grade = F	-0,086	-0,086	0,162
loan_grade = G	-0,044	-0,044	0,083
default=Y	-0,463	-0,463	0,866
default=N	0,463	0,463	-0,866
income	0,013	0,009	-0,020
employment_length	0,074	-0,008	-0,063
loan_percent_income	-0,155	0,048	0,103
loan_status	1,100	1,163	1,382
interest_rate	-1,084	-0,005	1,035



Gambar 6. Plot Centroid K-Means Perbedaan Profil Tiap Cluster per Atribut

Berdasarkan tabel *centroid* dan visualisasi (Tabel 5 serta Gambar 6), sifat masing-masing *Cluster* dapat dijelaskan secara terperinci:

Cluster 0 (10.782 observasi, 33,10%) merupakan segmen risiko rendah dengan profil keuangan paling kuat: didominasi kepemilikan rumah hipotek (*home_ownership*=MORTGAGE: +0,190), peringkat kredit A tertinggi (*loan_grade*=A: +1,421), suku bunga terendah (*interest_rate*: -1,084), masa kerja di atas rata-rata (*employment_length*: +0,074), rasio cicilan terhadap pendapatan terendah (*loan_percent_income*:

-0,155), serta riwayat gagal bayar terendah (*default*=Y: -0,463). Segmen ini menggambarkan debitur mapan dengan aset properti, kredit berkualitas tinggi, dan beban keuangan ringan, profil risiko rendah yang paling ideal bagi pemberi pinjaman. *Cluster 1* (10.448 observasi, 32,07%) merupakan kelompok risiko menengah: didominasi peringkat kredit B (*loan_grade*=B: +1,455), status penyewa rumah (*home_ownership*=RENT: +0,054), suku bunga mendekati rata-rata populasi (*interest_rate*: -0,005), jumlah pinjaman sedikit di atas rata-rata (*loan_amount*: +0,064), dan riwayat gagal bayar setara dengan *Cluster 0* (*default*=Y: -0,463). Segmen ini mencerminkan debitur dengan stabilitas keuangan menengah yang tinggal di rumah sewaan dengan kredit berkualitas cukup baik. *Cluster 2* (11.346 observasi, 34,83%) memperlihatkan profil risiko tertinggi yang sangat jelas: riwayat gagal bayar tertinggi (*default*=Y: +0,866) menandakan konsentrasi debitur bermasalah; peringkat kredit substandar C (*loan_grade*=C: +0,930), D (+0,662), E (+0,326), F (+0,162), dan G (+0,083) yang seluruhnya dominan di *Cluster* ini; suku bunga tertinggi (*interest_rate*: +1,035) sebagai premium risiko dari pemberi pinjaman; serta rasio cicilan terhadap pendapatan tertinggi (*loan_percent_income*: +0,103) yang mengindikasikan beban keuangan paling berat. Segmen ini merepresentasikan debitur dengan profil kredit paling lemah dan probabilitas gagal bayar tertinggi.

Perbedaan paling signifikan antar *Cluster* terletak pada tiga atribut diskriminatif utama. Pertama, atribut *default*=Y menunjukkan kontras yang tajam: *centroid Cluster 2* sebesar 0,866 (di atas rata-rata populasi) berbanding *centroid Cluster 0* dan *Cluster 1* yang sama-sama berada di -0,463. Hal ini mengindikasikan bahwa *Cluster 2* didominasi oleh debitur dengan riwayat gagal bayar sebelumnya sebuah prediktor kuat terhadap risiko kredit macet di masa mendatang [17]. Kedua, atribut *interest_rate* menunjukkan *centroid Cluster 2* (+1,035) jauh di atas *Cluster 0* (-1,084) dan *Cluster 1* (-0,005) dalam skala Z-score, mengindikasikan bahwa debitur di *Cluster 2* cenderung mendapatkan suku bunga yang lebih tinggi yang umumnya ditetapkan oleh pemberi pinjaman sebagai kompensasi atas profil risiko yang lebih buruk.

Ketiga, pola *loan_grade* menunjukkan polarisasi yang jelas: *Cluster 0* memiliki *centroid loan_grade*=A tertinggi (+1,421) dan kepemilikan rumah berstatus hipotek di atas rata-rata (*home_ownership*=MORTGAGE = +0,190), mencerminkan profil debitur dengan kredit berkualitas tinggi dan stabilitas finansial yang baik. *Cluster 1* didominasi oleh *loan_grade*=B (*centroid* = +1,455) yang menggambarkan debitur dengan kualitas kredit baik namun belum setara *Cluster 0*. Sebaliknya, *Cluster 2* memiliki *centroid loan_grade*=C (+0,930) dan *loan_grade*=D (+0,662) tertinggi, yang merepresentasikan debitur dengan peringkat kredit substandar. Profil *centroid* ini secara konsisten

mendukung penetapan *Cluster 2* sebagai segmen risiko tinggi, *Cluster 1* sebagai risiko menengah, dan *Cluster 0* sebagai risiko rendah.

3.3. Hasil *Decision Tree*

Model *Decision Tree* dengan kriteria *Gain Ratio* dan kedalaman maksimum 10 dilatih menggunakan partisi data latih (70%) yang telah melalui proses penyeimbangan kelas menggunakan *SMOTE Upsampling* (*number_of_neighbours=5*, *equalize_classes=true*, *auto_detect_minority_class=true*), kemudian diuji pada partisi data uji (30%) yang tetap mencerminkan distribusi data asli tanpa penyeimbangan. Atribut yang memberikan kontribusi tertinggi dalam pemisahan *node* menurut *Gain Ratio* meliputi *interest_rate*, *loan_percent_income*, *loan_grade*, serta *home_ownership* — sesuai dengan nilai *centroid* yang mengindikasikan perbedaan yang jelas antar *Cluster* pada atribut-atribut tersebut.

Penggunaan *Gain Ratio* sebagai kriteria percabangan sebagai pengganti *Information Gain* konvensional memberikan keunggulan dalam menangani atribut dengan banyak nilai kategorikal seperti *loan_grade* (7 kategori A–G) dan *loan_intent* (6 kategori). *Gain Ratio* menormalisasi *Information Gain* dengan *Split Information* sehingga atribut dengan banyak nilai tidak secara otomatis mendominasi pemilihan *node* percabangan [18]. Hal ini menghasilkan pohon keputusan yang lebih generalisabel dan tidak *overfit* terhadap atribut kategorikal berkardinitas tinggi.

Mekanisme *prepruning* dengan parameter *minimal_gain=0,01* dan *minimal_size_for_split=4* memastikan percabangan hanya terjadi jika memberikan peningkatan informasi yang berarti dan jumlah data pada *node* mencukupi, sedangkan *post-pruning* berbasis *confidence=0,1* dengan *minimal_leaf_size=2* memangkas cabang yang terlalu spesifik terhadap data latih. Kombinasi kedua teknik *pruning* ini berkontribusi pada akurasi 89,26% pada data uji, tidak menunjukkan indikasi *overfitting* yang ekstrem berdasarkan performa pada data uji yang signifikan meskipun model dilatih pada data latih yang telah diseimbangkan melalui *SMOTE*. Penerapan *SMOTE* pada data latih berhasil meningkatkan kemampuan model dalam mendeteksi kelas minoritas (kredit macet) dengan *Recall* sebesar 73,05%, yang lebih relevan dalam konteks sistem peringatan dini NPL dibandingkan sekadar mengejar akurasi keseluruhan pada distribusi kelas yang tidak seimbang [19].

3.4. Evaluasi Performa *DecisionTree*

accuracy: 89.26%

	true false	true true	class precision
pred. false	7178	570	92.64%
pred. true	480	1545	76.30%
class recall	93.73%	73.05%	

precision: 92.64% (positive class: false)

	true true	true false	class precision
pred. true	1545	480	76.30%
pred. false	570	7178	92.64%
class recall	73.05%	93.73%	

recall: 93.73% (positive class: false)

	true true	true false	class precision
pred. true	1545	480	76.30%
pred. false	570	7178	92.64%
class recall	73.05%	93.73%	

f_measure: 93.18% (positive class: false)

	true true	true false	class precision
pred. true	1545	480	76.30%
pred. false	570	7178	92.64%
class recall	73.05%	93.73%	

Gambar 7. *PerformanceVector* (Accuracy, Precision, Recall, dan F-Measure)

Tabel 6. Hasil Evaluasi Performa *Decision Tree*

Metrik	Nilai	Positive Class
Accuracy	89,26%	-
Precision	76,30%	true (Macet)
Recall	73,05%	true (Macet)
F-Measure	74,58%	true (Macet)

Akurasi sebesar 89,26% menandakan bahwa model berhasil mengklasifikasikan hampir sembilan dari sepuluh debitur dengan tepat. *Precision* sebesar 76,30% (*positive class: true*) berarti dari total prediksi kredit macet yang dilakukan model, 76,30% adalah benar, dengan 480 prediksi merupakan *false positive*, sehingga proporsi alarm palsu berada pada tingkat yang dapat dikendalikan. Nilai ini penting dalam sektor perbankan karena dapat membantu meminimalkan penolakan kredit yang tidak seharusnya terjadi. *Recall* yang tercatat 73,05% mengindikasikan kemampuan model untuk menemukan 1.545 dari total 2.115 kasus macet yang sebenarnya (menyisakan 570 *false negative*). Keseimbangan antara *Precision* dan *Recall* yang relatif lebih baik dibandingkan model tanpa penyeimbangan kelas merupakan hasil dari penerapan *SMOTE* pada data latih, yang memungkinkan model mendapatkan eksposur yang lebih proporsional terhadap kelas minoritas (kredit macet). *F-Measure* sebesar 74,58% sebagai rata-rata harmonik mencerminkan keseimbangan yang wajar antara kemampuan deteksi dan ketepatan prediksi. *Recall* untuk kelas negatif (kredit lancar) mencapai 93,73%, memperlihatkan keandalan model dalam mengidentifikasi debitur yang tidak macet.

Jika dibandingkan dengan studi-studi serupa, hasil ini tergolong kompetitif. Nazifahet al. melaporkan bahwa *Decision Tree* menghasilkan akurasi antara 82 % hingga 88 % pada dataset kredit lain [20]. Kristiani & Jatmiko mencatat akurasi C4.5 sebesar 88,4 % pada data perbankan Indonesia [21]. Pencapaian akurasi 89,26% pada penelitian ini berada di atas rentang yang dilaporkan kedua studi tersebut, dengan keunggulan tambahan berupa kemampuan deteksi kredit macet (*Recall* 73,05%) yang lebih eksplisit dilaporkan sebuah dimensi evaluasi yang tidak tersedia pada studi-studi pembandingan tersebut..

Tabel 7. *Confusion Matrix Decision Tree* (Data Uji 9.773 Record)

	Pred. false	Pred. true (Macet)
Aktual false (Lancar)	TN = 7.178	FP = 480
Aktual true (Macet)	FN = 570	TP = 1.545
Class Recall	93,73%	73,05%

Analisis lebih lanjut terhadap *confusion matrix* mengungkap karakteristik prediksi model yang relevan untuk aplikasi perbankan. Nilai FN=570 yang lebih besar dari FP=480 menunjukkan bahwa model masih cenderung lebih sering meloloskan debitur bermasalah (*false negative*) dibandingkan menolak debitur yang sebenarnya layak (*false positive*), meskipun selisih keduanya relatif lebih kecil dibandingkan model klasifikasi tanpa penyeimbangan kelas. Dalam perspektif manajemen risiko kredit, kondisi ini bermakna model memiliki kecenderungan *type II error* yang lebih tinggi yaitu kegagalan mendeteksi kredit yang berpotensi macet [21]. Implikasinya, institusi keuangan yang mengadopsi model ini perlu menerapkan lapisan pengamanan tambahan seperti verifikasi manual untuk debitur yang diprediksi lancar namun memiliki skor risiko batas (*borderline*).

Keseimbangan antara *Recall* kelas macet (73,05%) dan *Recall* kelas lancar (93,73%) mencerminkan hasil dari penerapan SMOTE pada data latih yang memungkinkan model mendapatkan eksposur lebih proporsional terhadap kelas minoritas. Meskipun demikian, *Precision* kelas macet sebesar 76,30% menunjukkan bahwa ketika model memutuskan suatu kredit termasuk macet, keputusan tersebut memiliki tingkat keandalan yang memadai dengan 1.545 prediksi benar dari total 2.025 prediksi macet yang dilakukan (TP=1.545, FP=480). Kombinasi *Precision* dan *Recall* yang lebih seimbang ini menghasilkan *F-Measure* 74,58%, yang menunjukkan model lebih cocok digunakan sebagai sistem peringatan dini tahap pertama yang diikuti dengan verifikasi manual, bukan sebagai penentu keputusan tunggal.

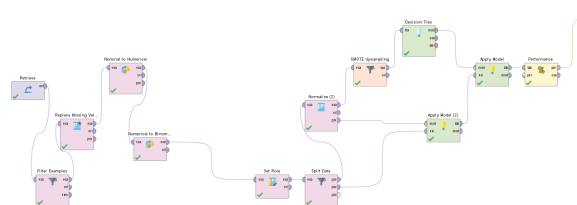
Nilai *Recall* kelas macet sebesar 73,05% perlu dipahami dalam konteks distribusi kelas data dan mekanisme SMOTE yang diterapkan. Secara teoretis, *Recall* yang belum mencapai angka optimal disebabkan oleh tiga faktor yang saling berinteraksi. Pertama, meskipun SMOTE berhasil menyeimbangkan distribusi data latih, pola sintesis yang dihasilkan tidak sepenuhnya merepresentasikan karakteristik debitur macet di dunia nyata, sehingga model *Decision Tree* tetap mengalami kesulitan dalam menggeneralisasi pola kegagalan kredit pada data uji yang tidak melalui augmentasi. Kedua, algoritma *Decision Tree* dengan parameter *Gain Ratio* cenderung membangun pohon keputusan yang lebih mengutamakan ketepatan prediksi secara keseluruhan (*Accuracy*) dibandingkan sensitivitas terhadap kelas minoritas, sehingga batas keputusan yang terbentuk masih berpotensi bias ke arah kelas mayoritas (kredit lancar). Ketiga, ketidakseimbangan kelas residual pada data uji (78,18% lancar berbanding 21,82% macet) yang tidak

melalui SMOTE tetap memengaruhi distribusi prediksi model pada tahap inferensi. Dalam konteks *early warning system* NPL, *Recall* 73,05% bermakna bahwa dari setiap 100 debitur yang sesungguhnya akan mengalami kredit macet, model berhasil mengidentifikasi 73 di antaranya sejak dini. Meskipun 27 kasus lolos dari deteksi (*false negative*), angka ini masih memberikan nilai preventif yang signifikan dibandingkan pendekatan konvensional tanpa model prediktif, dengan catatan bahwa institusi keuangan perlu menerapkan mekanisme verifikasi berlapis khususnya pada debitur yang diprediksi lancar namun memiliki satu atau lebih atribut risiko tinggi seperti *loan_grade C* ke atas atau riwayat *default=Y* untuk meminimalkan risiko kredit macet yang tidak terdeteksi oleh model.

Untuk meningkatkan performa deteksi kredit macet pada penelitian mendatang, eksplorasi terhadap penyesuaian nilai ambang batas klasifikasi (*threshold tuning*) dari nilai *default* 0,5 ke nilai yang lebih rendah berpotensi meningkatkan *Recall* kelas macet tanpa harus mengubah arsitektur model. Alternatif lain yang dapat dipertimbangkan adalah penerapan *cost-sensitive learning* yang secara eksplisit memberikan bobot penalti lebih tinggi pada kesalahan *false negative*, serta perbandingan dengan algoritma *ensemble* seperti *Random Forest* yang umumnya lebih robust terhadap ketidakseimbangan kelas.

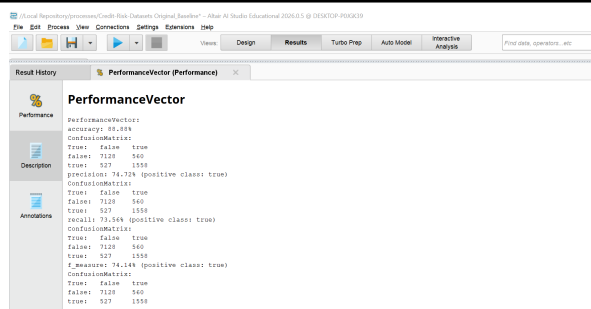
3.5. Analisis Kontribusi Komponen K-Means dalam Framework Paralel

Untuk mengevaluasi kontribusi komponen K-Means dalam *framework* yang diusulkan, dilakukan eksperimen ablasi dengan menjalankan *Decision Tree* + SMOTE secara mandiri tanpa integrasi K-Means pada dataset dan konfigurasi parameter yang identik. Proses dilakukan dengan menghapus operator *Multiply*, *Normalize* pada cabang K-Means, dan operator *Clustering*, sedangkan alur klasifikasi tetap dipertahankan. Pada skenario ini, proses normalisasi dilakukan hanya pada data latih setelah Split Data, kemudian parameter normalisasi diterapkan pada data uji melalui *Apply Model*, sehingga prosedur evaluasi tetap konsisten dengan model utama dan terhindar dari data *leakage*. Visualisasi alur proses ditampilkan pada Gambar 8.



Gambar 8. Alur Proses Ablation Study

Output yang dihasilkan berupa *PerformanceVektor* yang dihasilkan oleh operator *Performance* berisi informasi nilai *Accuracy*, *Precision*, *Recall*, dan *F-Measure* yang dihasilkan oleh model.



Gambar 9. PerformanceVector Ablation Study

Gambar 9 menunjukkan bahwa pada data uji sebanyak 9.773 record, model dengan SMOTE menghasilkan confusion matrix dengan TP=1.558, TN=7.128, FP=527, dan FN=560, sehingga diperoleh Accuracy 88,88%, Precision 74,72%, Recall 73,56%, dan F-Measure 74,14% (positive class: true). Hasil perbandingan lengkap antara model tanpa SMOTE dan dengan SMOTE disajikan pada Tabel 8.

Tabel 8. Perbandingan Performa dengan dan Tanpa Clustering K-

Metrik	Means		Selisih
	K-Means + DT (Framework Usulan)	DT Tanpa K-Means	
Accuracy	89,26%	88,88%	+0,38%
Precision (Macet)	76,30%	74,72%	+1,58%
Recall (Macet)	73,05%	73,56%	-0,51%
F-Measure	74,58%	74,14%	+0,44%

Hasil pada Tabel 8 menunjukkan bahwa penerapan K-Means Clustering memberikan kontribusi positif terhadap performa Decision Tree, meskipun selisihnya relatif kecil. Accuracy meningkat sebesar 0,38%, Precision kelas macet meningkat sebesar 1,58%, dan F-Measure meningkat sebesar 0,44%. Satu-satunya metrik yang mengalami penurunan adalah Recall kelas macet sebesar 0,51%, yang mengindikasikan bahwa penambahan tahap Clustering sedikit mengurangi kemampuan model dalam mendeteksi kasus macet secara menyeluruh. Dan secara keseluruhan, pada temuan ini mengonfirmasi bahwa integrasi K-Means dalam framework paralel tidak merugikan performa Decision Tree, melainkan memberikan peningkatan positif pada Precision, Accuracy, dan F-Measure, meskipun disertai penurunan marginal pada Recall (-0,51%). Mengingat Recall merupakan metrik yang diprioritaskan dalam konteks early warning system NPL sebagaimana ditetapkan pada bagian Metode, penurunan ini perlu dicermati meskipun besarnya tergolong kecil dan berada dalam rentang yang wajar akibat variasi proses pelatihan model. Secara keseluruhan, integrasi K-Means tetap dapat dipertimbangkan sebagai komponen yang bernilai tambah pada framework yang diusulkan, terutama karena kontribusinya pada aspek segmentasi debitur yang tidak dapat diperoleh dari Decision Tree secara mandiri, sementara trade-off kecil pada Recall menjadi salah satu pertimbangan yang perlu diperhitungkan institusi keuangan sebelum mengadopsi pendekatan paralel ini secara penuh.

Hasil ablasi ini mengonfirmasi bahwa tahap K-Means Clustering dalam framework yang diusulkan bukan sekadar pelengkap, melainkan memberikan kecenderungan peningkatan performa terhadap kualitas prediksi model. Operator Multiply yang memungkinkan kedua algoritma memproses representasi data yang identik terbukti menghasilkan konsistensi preprocessing yang berdampak positif pada performa klasifikasi, sekaligus menghasilkan nilai tambah berupa segmentasi debitur yang tidak dapat diperoleh dari Decision Tree secara mandiri.

3.6. Implikasi Manajerial

Hasil paralel K-Means dan Decision Tree memberikan nilai strategis ganda: (1) Segmentasi tiga Cluster memungkinkan kebijakan kredit terdiferensiasi, Cluster 2 memerlukan persyaratan agunan ketat, suku bunga premium, dan monitoring intensif sebagai segmen risiko tinggi; Cluster 1 dapat dikenai persyaratan standar dengan pemantauan berkala sebagai segmen risiko menengah; sementara Cluster 0 dapat ditawarkan produk kredit kompetitif dengan suku bunga preferensial sebagai reward atas profil risiko rendah yang ditunjukkan oleh loan_grade A dominan dan riwayat pembayaran bersih. (2) Model Decision Tree dengan akurasi 89,26% dan Precision kelas macet 76,30% siap diintegrasikan ke sistem scoring kredit existing sebagai lapisan analisis tambahan, dengan threshold yang dapat disesuaikan berdasarkan risk appetite institusi. Mengingat model menunjukkan kecenderungan false negative (FN=570) yang lebih besar dari false positive (FP=480), institusi disarankan untuk menerapkan verifikasi manual pada debitur yang diprediksi lancar namun berada di perbatasan skor risiko (borderline) guna meminimalkan risiko kredit macet yang tidak terdeteksi.

4. Kesimpulan

Penelitian ini berhasil mengimplementasikan alur proses paralel K-Means Clustering dan Decision Tree pada Credit Risk Dataset (32.576 record) menggunakan RapidMiner 2026.0.5 dengan operator Multiply sebagai titik percabangan. K-Means (k=3, Mixed Euclidean Distance) menghasilkan tiga segmen debitur yang dapat diinterpretasikan secara bisnis: Cluster 0 (10.782 debitur, 33,10%) sebagai segmen risiko rendah dengan loan_grade A dominan dan suku bunga terendah (interest_rate = -1,084); Cluster 1 (10.448 debitur, 32,07%) sebagai segmen risiko menengah dengan loan_grade B dominan dan suku bunga rata-rata (interest_rate = -0,005); serta Cluster 2 (11.346 debitur, 34,83%) sebagai segmen risiko tinggi yang ditandai suku bunga tertinggi (interest_rate = +1,035), dominasi loan_grade C-G, dan konsentrasi riwayat gagal bayar terbesar (default=Y = +0,866). Decision Tree (Gain Ratio, max_depth=10) dengan penerapan SMOTE pada data latih mencapai akurasi 89,26%, Precision 76,30%, Recall 73,05%, dan F-Measure 74,58% pada data uji sebesar 30%, dengan confusion matrix menghasilkan TP=1.545, TN=7.178,

FP=480, dan FN=570, serta *class Recall* kredit lancar sebesar 93,73%.

Kontribusi utama penelitian ini adalah kerangka *profiling* risiko kredit berbasis interpretasi kuantitatif tabel *centroid K-Means* yang menghasilkan tiga segmen debitur dengan karakteristik terukur dan dapat langsung diterjemahkan menjadi kebijakan kredit yang terdiferensiasi, sesuatu yang tidak dapat dihasilkan oleh model klasifikasi tunggal. Eksperimen ablasi yang dilakukan dengan menjalankan *Decision Tree* + SMOTE tanpa komponen K-Means menghasilkan akurasi 88,88%, *Precision* 74,72%, *Recall* 73,56%, dan *F-Measure* 74,14%, mengkonfirmasi bahwa integrasi K-Means dalam *framework* paralel memberikan kontribusi positif pada ketepatan prediksi khususnya *Precision* kelas macet (+1,58%) meskipun dengan selisih yang marginal.

Penelitian ini memiliki beberapa keterbatasan yang perlu dipertimbangkan dalam menginterpretasikan hasil dan generalisasinya. Pertama, terkait cakupan perbandingan metode, evaluasi kontribusi *K-Means* dilakukan melalui eksperimen ablasi terkontrol antara *framework* paralel *K-Means* dan *Decision Tree* dengan *Decision Tree* mandiri tanpa komponen *Clustering*, sehingga posisi komparatif *framework* ini terhadap algoritma yang secara fundamental berbeda, seperti *Logistic Regression*, *Random Forest*, atau *XGBoost*, belum dapat ditetapkan. Hasil ablasi menunjukkan kontribusi positif *K-Means* terhadap *Accuracy*, *Precision*, dan *F-Measure* dengan margin kecil (0,38%–1,58%), sementara *Recall* kelas macet mengalami penurunan marginal (–0,51%) yang perlu dicermati mengingat *Recall* merupakan metrik yang diprioritaskan sebagaimana ditetapkan pada bagian Metode. Signifikansi statistik dari selisih performa tersebut belum diuji secara formal, sehingga perbandingan multimetode disertai uji statistik seperti uji t berpasangan pada hasil *cross-validation* menjadi agenda penelitian mendatang yang diprioritaskan.

Kedua, evaluasi efisiensi komputasi dari implementasi paralel melalui *operator Multiply*, seperti perbandingan waktu eksekusi terhadap pendekatan sekuensial atau pengujian pada skala data lebih besar, belum menjadi fokus penelitian ini. Penelitian ini juga menggunakan satu dataset tunggal (*Credit Risk Dataset*), sehingga generalisasi model pada data kredit institusi keuangan lain masih perlu diuji lebih lanjut.

Ketiga, terkait *Recall* kelas macet sebesar 73,05%, penyesuaian ambang batas klasifikasi (*threshold tuning*) disarankan dieksplorasi pada penelitian mendatang guna meningkatkan sensitivitas deteksi. Angka ini secara kualitatif dapat dipandang memberikan nilai preventif yang berarti, karena 73 dari setiap 100 kasus macet berhasil terdeteksi sejak dini dibandingkan tanpa mekanisme deteksi sama sekali, meskipun klaim ini bersifat interpretatif dan belum diuji secara empiris terhadap pendekatan konvensional. Institusi keuangan yang mengadopsi model ini

disarankan menerapkan verifikasi berlapis, khususnya pada debitur yang diprediksi lancar namun memiliki atribut risiko tinggi seperti *loan_grade* C ke atas atau riwayat *default=Y*. Pengembangan sistem berbasis *web* yang mengintegrasikan model ini dengan data kredit *real-time* institusi keuangan menjadi langkah implementasi praktis pada tahap selanjutnya.

Daftar Rujukan

- [1] P. Anisa, "Kinerja Keuangan dan Risiko Kredit pada Sektor Perbankan Indonesia," *J. Ilm. Ekon. dan Keuang.*, vol. 1, no. 1, pp. 24–31, 2025, [Online]. Available: <https://pustakajurnal.web.id/index.php/jiek/article/view/56>
- [2] H. Hendri, H. Andrianof, R. Robianto, H. Awal, and O. A. Putra, "A hybrid data mining for predicting scholarship recipient students by combining K-means and C4.5 methods," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 33, no. 3, pp. 1726–1735, 2024, doi: 10.11591/ijeecs.v33.i3.pp1726-1735.
- [3] M. Ahmed, R. Seraj, S. Mohammed, and S. Islam, "The k-means Algorithm: A Comprehensive Survey and Performance Evaluation," *Electronics*, vol. 9, no. 8, pp. 1–12, 2020, doi: 10.3390/electronics9081295.
- [4] B. T. Jijo and A. M. Abdulazeed, "Classification Based on Decision Tree Algorithm for Machine Learning," *J. Appl. Sci. Technol. TRENDS*, vol. 02, no. 01, pp. 20–28, 2021, doi: 10.38094/jastt20165.
- [5] A. Ali *et al.*, "applied sciences Financial Fraud Detection Based on Machine Learning: A Systematic Literature Review," *Appl. Sci.*, vol. 12, no. 19, 2022, doi: 10.3390/app12199637.
- [6] A. W. Kesumaningtias and B. Wicaksana, "Penerapan Metode K-Means Clustering untuk Menentukan Debitur Existing Penerima Pinjaman KUR," *J. SAINTEKOM (Sains dan Teknol. Komputasi)*, vol. 01, no. 03, pp. 16–25, 2025, doi: 10.36350/jskom.v1i3.58.
- [7] L. N. Rahmadhani and F. Ardiani, "RANCANG BANGUN SISTEM PREDIKSI KELAYAKAN PINJAMAN Abstraksi Keywords: Pendahuluan Tinjauan Pustaka," *J. Inf. Syst. Manag. (JOISM)*, vol. 7, no. 2, pp. 178–185, 2026, doi: 10.24076/joism.2026v7i2.2423.
- [8] T. P. Wahjuningsih, T. A. Setiawan, A. Ilyas, and A. Subagyo, "OPTIMASI SISTEM CREDIT SCORING PEMBIAYAAN PADA LEMBAGA KEUANGAN MIKRO DENGAN NEURAL NETWORK INTEGRASI SAMPLE BOOTSTRAPPING DAN WEIGHTED PCA," *DINAMIK*, vol. 31, no. 1, pp. 267–275, 2026, doi: 10.35315/dinamik.v31i1.10381.
- [9] S. A. Arnomo, A. A. Fajrin, Y. Siyamto, S. Fairuz, and N. Sadikin, "Evaluasi Model Decision Tree Pada Keputusan Kelayakan Kredit," *J. DESAIN DAN Anal. Teknol.*, vol. 2, no. 2, pp. 200–206, 2023, doi: <https://doi.org/10.58520/jddat.v2i2.39>.
- [10] M. M. Mutoffar, E. Retnoningsih, Y. L. Yasik, and Eliza, *Decoding Intelligence Algoritma Machine Learning dalam Aksi dan Bisnis*. Pt Kimhsafi Alung Cipta, 2025.
- [11] N. Suhadolnik, J. Ueyama, and S. Da Silva, "Machine Learning for Enhanced Credit Risk Assessment: An Empirical Approach," *J. Risk Financ. Manag.*, vol. 16, no. 12, 2023, doi: 10.3390/jrfm16120496.
- [12] A. A. Montevechi, R. de C. Miranda, A. L. Medeiros, and J. A. B. Montevechi, "Advancing credit risk modelling with Machine Learning: A comprehensive review of the state-of-the-art," *Eng. Appl. Artif. Intell.*, vol. 137, 2024, doi: 10.1016/j.engappai.2024.109082.
- [13] A. Puspitaloka and A. Herman Bedi, "PROFILING RISIKO KREDIT NASABAH BERBASIS K MEANS CLUSTERING DENGAN OPTIMASI SILHOUETTE COEFFICIENT PADA BANK MALUKUMALUT MASOHI," 2026. [Online]. Available: <https://repository.itpln.ac.id/id/eprint/5572>
- [14] M. Herfina and Y. P. Muchda, "Non-Performing Loan (NPL) Dan Kinerja Keuangan: Studi Pada PT Bank Negara Indonesia (Persero) Tbk Tahun 2022 – 2024," *J. Akad. Akunt. Indones.*

- Padang*, vol. 5, no. 1, pp. 186–195, 2025, doi: 10.31933/13e84r24.
- [15] M. Fajar, N. Rahaningsih, R. D. Dana, T. Informatika, and K. Akuntansi, “ANALISIS POLA PENJUALAN OBAT DI APOTEK AN-NAAFI MENGGUNAKAN METODE K-MEANS CLUSTERING,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 486–492, 2024, doi: 10.36040/jati.v8i1.8395.
- [16] M. I. A. R. G. Dwilestari and N. Suarna, “PREDIKSI PERSETUJUAN PINJAMAN MENGGUNAKAN DATASET LOAN APPROVAL MENGGUNAKAN ALGORITMA KLASIFIKASI,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 9, no. 1, pp. 1342–1347, 2025, doi: 10.36040/jati.v9i1.12486.
- [17] J. Alvi, I. Arif, and K. Nizam, “Heliyon Advancing financial resilience : A systematic review of *default* prediction models and future directions in credit risk management,” *Heliyon*, vol. 10, no. 21, p. e39770, 2024, doi: 10.1016/j.heliyon.2024.e39770.
- [18] Q. A. S and Z. Fatah, “KLASIFIKASI KELULUSAN MAHASISWA MENGGUNAKAN METODE DECISION TREE MENGGUNAKAN APLIKASI RAPIDMINER,” *J. Ris. Tek. Komput.*, vol. 1, no. 4, pp. 58–64, 2024, doi: 10.69714/em8qnw54.
- [19] C. N. Syahputri and M. S. Hasibuan, “OPTIMASI KLASIFIKASI DECISION TREE DENGAN TEKNIK PRUNING UNTUK MENGURANGI OVERFITTING,” *JSiI | J. Sist. Inf.*, vol. 11, no. 2, pp. 87–96, 2024, doi: 10.30656/jsii.v11i2.9161.
- [20] N. Nazifah and C. Prianto, “Analisis Perbandingan Decision Tree Algoritma C4 . 5 dengan algoritma lainnya : Sistematic Literature Review,” *J. Inform. Dan Teknol. Komput. (J-ICOM)*, vol. 04, no. 02, pp. 57–64, 2023, doi: 10.55377/j-icom.v4i2.7719.
- [21] P. A. Kristianti, S. Jatmiko, P. Lunak, S. Informasi, U. Gunadarma, and J. Pusat, “KLASIFIKASI PEMILIHAN PRODUK PERBANKAN DENGAN ALGORITMA DECISION TREE ID3 DAN C4 . 5,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 9, no. 4, pp. 5999–6005, 2025, doi: 10.36040/jati.v9i4.13907.