

COMPARATIVE ANALYSIS OF *GAUSSIAN MIXTURE MODEL* AND *ROBUST TRIMMED GAUSSIAN MIXTURE MODEL* FOR CUSTOMER SEGMENTATION BASED ON CREDIT PAYMENT BEHAVIOR

ANALISIS PERBANDINGAN *GAUSSIAN MIXTURE MODEL* DAN *ROBUST TRIMMED GAUSSIAN MIXTURE MODEL* UNTUK SEGMENTASI NASABAH BERDASARKAN PERILAKU PEMBAYARAN KREDIT

Qindy Naura¹, Elardian Putera Ramadhan^{1‡}, Nisrina Suci Fadila¹, Shafi Faris Arif Rabbani¹, Muhammad Sabda Agama¹, and Septian Rahardiantoro¹

¹Department of Statistics, IPB University, Indonesia

[‡]corresponding author: puteraramadhan@apps.ipb.ac.id

Copyright © 2026 Qindy Naura, Elardian Putera Ramadhan, Nisrina Suci Fadila, Shafi Faris Arif Rabbani, Muhammad Sabda Agama, and Septian Rahardiantoro. This is an open-access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract

Customer credit segmentation based on credit payment behavior is important for identifying payment behavior patterns and supporting customer management strategies in the banking sector. However, financial transaction data often contain outliers that may reduce clustering performance and distort parameter estimation. This study aims to compare the performance of Gaussian Mixture Model (GMM) and Robust Trimmed Gaussian Mixture Model (RTGMM) for customer segmentation based on credit payment behavior. The study uses the Default of Credit Card Clients dataset from the UCI Machine Learning Repository consisting of 30,000 customers. Data preprocessing included ordinal cumulative probability encoding, Yeo–Johnson transformation, and standardization. Model evaluation was conducted using Bayesian Information Criterion (BIC), Integrated Completed Likelihood (ICL), Silhouette Score, Calinski–Harabasz Index, and Davies–Bouldin Index. The optimal GMM configuration was obtained at $K = 4$, while RTGMM achieved optimal performance at $K = 5$ with trimming proportion $\alpha = 0.10$. Results indicate that RTGMM produced more stable clustering and lower BIC and ICL

values, while GMM showed slightly better cluster separation quality based on internal validation metrics. These findings demonstrate the effectiveness of robust clustering approaches for handling complex and heterogeneous financial datasets containing extreme observations.

Keywords: Clustering, Customer Segmentation, Gaussian Mixture Model, Robust Trimmed Gaussian Mixture Model.

Abstrak

Segmentasi nasabah berdasarkan perilaku pembayaran kredit berperan penting dalam mengidentifikasi pola pembayaran serta mendukung strategi pengelolaan nasabah pada sektor perbankan. Namun, data transaksi keuangan umumnya mengandung pencilan yang dapat menurunkan kualitas klasterisasi dan mendistorsi estimasi parameter model. Penelitian ini bertujuan membandingkan performa *Gaussian Mixture Model* (GMM) dan *Robust Trimmed Gaussian Mixture Model* (RTGMM) untuk segmentasi nasabah berdasarkan perilaku pembayaran kredit. Data yang digunakan adalah *Default of Credit Card Clients Dataset* dari UCI Machine Learning Repository yang terdiri atas 30.000 nasabah kartu kredit. Tahap pra-pemrosesan meliputi *ordinal cumulative probability encoding*, transformasi Yeo–Johnson, dan standarisasi data. Evaluasi model dilakukan menggunakan *Bayesian Information Criterion* (BIC), *Integrated Completed Likelihood* (ICL), *Silhouette Score*, *Calinski–Harabasz Index*, dan *Davies–Bouldin Index*. Hasil penelitian menunjukkan bahwa GMM optimal diperoleh pada $K = 4$, sedangkan RTGMM optimal pada $K = 5$ dengan proporsi *trimming* $\alpha = 0,10$. RTGMM menghasilkan klaster yang lebih stabil serta nilai BIC dan ICL lebih rendah, sementara GMM memiliki kualitas separasi klaster yang lebih baik berdasarkan metrik validasi internal. Temuan ini menunjukkan bahwa pendekatan *robust* efektif digunakan pada data finansial yang kompleks, heterogen, dan mengandung amatan ekstrem.

Kata Kunci: *Gaussian Mixture Model*, Klasterisasi, *Robust Trimmed Gaussian Mixture Model*, Segmentasi nasabah.

1. Pendahuluan

Sektor perbankan memiliki peran penting dalam pengelolaan dan penyaluran dana masyarakat. Perilaku pembayaran nasabah menjadi aspek penting dalam pengelolaan kredit karena keterlambatan pembayaran dapat memengaruhi kualitas pembiayaan dan stabilitas institusi keuangan (Melese *et al.*, 2023). Oleh karena itu, segmentasi nasabah berdasarkan pola pembayaran diperlukan untuk memahami karakteristik perilaku kredit nasabah. Seiring perkembangan teknologi, industri keuangan mulai memanfaatkan pendekatan pembelajaran mesin untuk meningkatkan akurasi dan transparansi analisis data pelanggan (Robisco & Martínez, 2022).

Analisis kluster berbasis model seperti *Gaussian Mixture Model* (GMM) banyak digunakan karena kemampuannya memodelkan populasi sebagai campuran distribusi normal multivariat dengan parameter bobot campuran, vektor rata-rata, dan matriks kovarians (Fraley & Raftery, 2002). Estimasi parameter dilakukan melalui algoritma *Expectation-Maximization* (EM) yang bekerja secara iteratif melalui tahap penghitungan probabilitas posterior keanggotaan dan pemutakhiran parameter hingga konvergen (Bishop, 2006). Penentuan jumlah kluster optimal menggunakan kriteria *Bayesian Information Criterion* (BIC) yang menyeimbangkan kecocokan model dengan kompleksitas parameter, serta *Integrated Completed Likelihood* (ICL) yang mempertimbangkan kualitas pemisahan antarkluster (Shatwan *et al.*, 2024).

Meskipun GMM fleksibel dalam memodelkan struktur data yang kompleks, algoritma EM sangat sensitif terhadap pencilon sehingga dapat mendistorsi estimasi parameter dan menghasilkan segmentasi yang tidak objektif (Cuesta-Albertos *et al.*, 2008; García-Escudero & Mayo-Iscar, 2024). Pada kajian perilaku pembayaran kredit, penelitian sebelumnya lebih berfokus pada pendekatan prediktif seperti *Logistic Regression* dan *Random Forest* untuk mendeteksi gagal bayar (Yeh & Lien, 2009), sedangkan penerapan pendekatan klusterisasi *robust* untuk segmentasi nasabah masih relatif terbatas. Selain itu, kajian yang membandingkan secara langsung performa GMM dan RTGMM pada data perilaku pembayaran kredit juga masih relatif terbatas. RTGMM dikembangkan sebagai solusi dengan mengeluarkan sebagian kecil observasi *ber-likelihood* terendah dari proses estimasi sehingga menghasilkan parameter yang lebih stabil dan *robust* terhadap observasi ekstrem (Coretto & Hennig, 2016; Fadhlia & Hendri, 2025). Berdasarkan hal tersebut, penelitian ini bertujuan membandingkan performa RTGMM dan GMM dalam segmentasi risiko kredit menggunakan dataset *Default of Credit Card Clients* dari UCI Machine Learning Repository. Perbandingan dilakukan untuk melihat pengaruh pendekatan *robust* terhadap kualitas struktur kluster, stabilitas segmentasi, serta kemampuan model dalam menangani pencilon. Hasil penelitian diharapkan dapat memberikan gambaran mengenai perbedaan karakteristik segmentasi antara model *robust* dan *non-robust* sebagai pertimbangan dalam pemilihan metode klusterisasi pada data perilaku pembayaran nasabah yang heterogen dan berpotensi mengandung pencilon.

2. Metodologi

2.1 Data

Data yang digunakan dalam penelitian merupakan *real-world dataset* yang berasal dari institusi perbankan di Taiwan dan tersedia secara publik pada UCI Machine Learning Repository dengan nama *Default of Credit Card Clients Dataset*. Dataset terdiri dari 30.000 nasabah kartu kredit, dengan total 23 peubah prediktor dan 1 peubah respon. Berikut merupakan daftar peubah yang digunakan dalam penelitian.

Tabel 1: Peubah-peubah yang digunakan

Kelompok Peubah	Nama Peubah	Tipe Data	Deskripsi
Status pembayaran	X1–X6	Ordinal	Riwayat status pembayaran April–September 2005. Skala: -1 = tepat waktu; 1–9 = tingkat keterlambatan pembayaran dalam bulan.
Jumlah tagihan	X7–X12	Integer	Jumlah tagihan kartu kredit April–September 2005 (NT Dollar).
Jumlah pembayaran	X13–X18	Integer	Jumlah pembayaran aktual April–September 2005 (NT Dollar).

2.2 Metode Penelitian

Penelitian dilakukan melalui empat tahap sistematis yaitu pra-pemrosesan data, eksplorasi data, pemodelan GMM dan RTGMM, serta evaluasi dan perbandingan model menggunakan bahasa pemrograman Python pada *Google Colaboratory*.

2.2.1 Pra-pemrosesan Data

Pra-pemrosesan diawali dengan pemeriksaan data hilang dan kesesuaian tipe data untuk memastikan kelayakan analisis (Maharana *et al.*, 2022). Liu *et al.* (2021) membahas pentingnya kuantifikasi dalam analisis asosiasi antarpeubah ordinal. Sejalan dengan prinsip tersebut, peubah ordinal X1–X6 ditransformasi menggunakan *ordinal cumulative probability encoding* guna mempertahankan urutan kategori status pembayaran dalam bentuk numerik yang proporsional. Sementara itu, peubah kontinu X7–X18 ditransformasi menggunakan transformasi Yeo-Johnson pada Persamaan 1 untuk mereduksi kemencengan distribusi data finansial sekaligus menangani nilai nol dan negatif (Yeo & Johnson, 2000; Medina *et al.* 2019). Seluruh peubah kemudian distandarisasi agar tidak ada peubah yang mendominasi proses pembentukan klaster.

$$\psi(\lambda, y) = \begin{cases} \frac{(y+1)^\lambda - 1}{\lambda}, & \lambda \neq 0, y \geq 0 \\ \log(y+1), & \lambda = 0, y \geq 0 \\ \frac{(1-y)^{2-\lambda} - 1}{\lambda - 2}, & \lambda \neq 2, y < 0 \\ -\log(y-1), & \lambda = 2, y < 0 \end{cases} \quad (1)$$

dengan y adalah nilai peubah kontinu serta λ parameter transformasi yang diestimasi.

2.2.2 Eksplorasi Data

Eksplorasi data dilakukan untuk memahami karakteristik distribusi, mendeteksi pencilan, dan mengamati hubungan linear antarpeubah melalui visualisasi histogram, boxplot, dan korelasi Pearson (Fadhli & Hendri, 2025).

2.2.3 Pemodelan

Pemodelan dilakukan dalam dua pendekatan, GMM sebagai model dasar dan RTGMM sebagai pemodelan yang lebih *robust* untuk mengurangi pengaruh pencilan pada proses klasterisasi. Berikut merupakan tahap pemodelan yang dilakukan:

a. Estimasi *Gaussian Mixture Model* (GMM)

GMM memodelkan populasi sebagai campuran distribusi normal multivariat dengan parameter bobot campuran (π_k), vektor rata-rata (μ_k), dan matriks kovarians (Σ_k). Jumlah klaster optimal ditentukan melalui *grid search* pada rentang $K = 2, 3, \dots, 10$ (Astina et al., 2026). Kemudian struktur kovarians yang digunakan adalah *tied covariance* guna menangkap keterkaitan antar peubah secara konsisten. Estimasi parameter dilakukan melalui algoritma EM yang terdiri dari E-step dan M-step secara iteratif hingga konvergen (Bishop, 2006). Pada E-step dihitung probabilitas posterior:

$$\gamma(z_{nk}) = \frac{\pi_k N(x_k | \mu_k, \Sigma_k)}{\sum_{j=1}^K \pi_j N(x | \mu_j, \Sigma_j)} \quad (2)$$

Parameter kemudian diperbarui pada M-Step:

$$\mu_k^{new} = \frac{1}{N_k} \sum_{n=1}^N \gamma(z_{nk}) x_n \quad (3)$$

$$\Sigma_k^{new} = \frac{1}{N_k} \sum_{n=1}^N \gamma(z_{nk}) (x_n - \mu_k^{new})(x_n - \mu_k^{new})^T \quad (4)$$

$$\pi_k^{new} = \frac{N_k}{N} \quad (5)$$

dengan:

$$N_k = \sum_{n=1}^N \gamma(z_{nk}) \quad (6)$$

Iterasi berlanjut hingga *log-likelihood* konvergen:

$$\ln p(X | \mu, \Sigma, \pi) = \sum_{n=1}^N \ln \left\{ \sum_{k=1}^K \pi_k N(X_n | \mu_k, \Sigma_k) \right\} \quad (7)$$

dengan $\gamma(z_{nk})$ adalah estimasi nilai E-step, π adalah bobot sampel, N adalah jumlah observasi, μ adalah rata-rata sampel, x_n adalah data ke- n , k indeks klaster tertentu, K jumlah klaster keseluruhan, dan Σ adalah matriks kovarians.

b. Estimasi *Robust Trimmed Gaussian Mixture Model*

RTGMM merupakan pengembangan GMM yang lebih *robust* terhadap penciran melalui penambahan *trimming step* diantara algoritma E-step dan M-step (Fadhli & Hendri, 2025). Pada setiap iterasi, sebanyak α_n observasi dengan *likelihood* terendah dikeluarkan dari estimasi parameter sehingga hanya subset $H_\alpha = \{i: D_i \text{ termasuk top } n(1 - \alpha) \text{ likelihood tertinggi}\}$ digunakan dalam M-step:

$$D_i = \sum_{k=1}^K \pi_k N(x_i | \mu_k, \Sigma_k) \quad (8)$$

Parameter diperbaharui eksklusif menggunakan H_α :

$$\pi_k^{new} = \frac{1}{|H_\alpha|} \sum_{i \in H_\alpha} \gamma_{ik} \quad (9)$$

$$\mu_k^{new} = \frac{\sum_{i \in H_\alpha} \gamma_{ik} x_i}{\sum_{i \in H_\alpha} \gamma_{ik}} \quad (10)$$

$$\Sigma_k^{new} = \frac{\sum_{i \in H_\alpha} \gamma_{ik} (x_i - \mu_k^{new})(x_i - \mu_k^{new})^T}{\sum_{i \in H_\alpha} \gamma_{ik}} \quad (11)$$

Parameter diuji pada rentang K dan kovarians yang sama dengan GMM serta $\alpha \in \{0,05; 0,10\}$. Nilai trimming tersebut dipilih untuk menjaga keseimbangan antara ketahanan algoritma terhadap penciran dan retensi informasi data sehingga sekitar 90%–95% observasi tetap dipertahankan dalam proses pembentukan kluster

2.2.4 Evaluasi Model

Evaluasi model dilakukan menggunakan nilai BIC dan ICL untuk setiap kombinasi K dan struktur kovarians (Shatwan *et al.*, 2024):

$$BIC = 2 \log(L) - k \log(n) \quad (12)$$

$$ICL = \log(L) - k \log(n) + 2 \times \text{entropy} \quad (13)$$

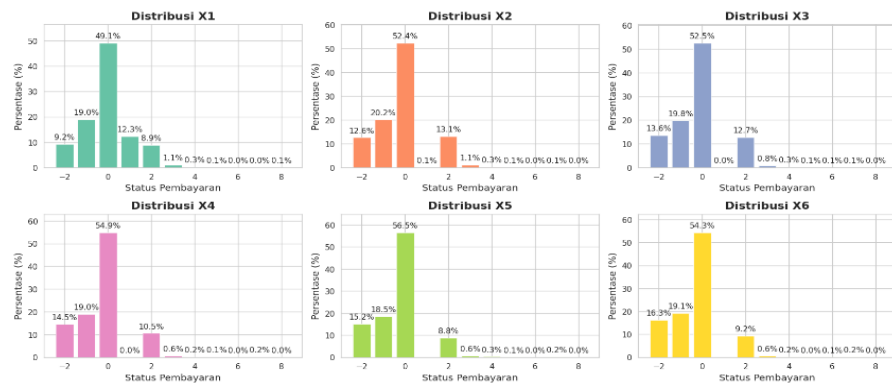
dengan k adalah jumlah parameter model, n adalah jumlah observasi, dan L nilai maksimum *likelihood* model. Selain menggunakan BIC dan ICL sebagai metrik utama, kualitas kluster juga dievaluasi menggunakan tiga metrik validasi internal yaitu *Silhouette Score* (lebih tinggi lebih baik), *Calinski-Harabasz Index* (lebih tinggi lebih baik), dan *Davies-Bouldin Index* (lebih kecil lebih baik). Kombinasi kelima metrik memastikan pemilihan model optimal baik secara statistik maupun kualitas segmentasi. Selain itu, pemilihan K optimal juga berdasarkan kemudahan interpretasi.

3 Hasil dan Pembahasan

3.1 Eksplorasi Data

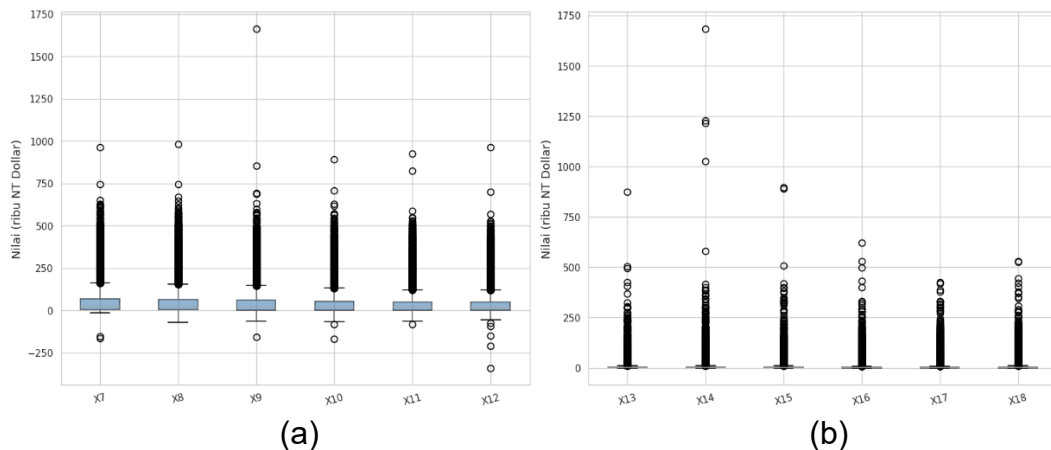
Analisis distribusi peubah status pembayaran bulanan (X1–X6) pada Gambar 1 menunjukkan dominansi nasabah pada kategori pembayaran minimum atau *revolving credit* sebesar 49,10%–56,50%. Nasabah yang melunasi tagihan penuh berada pada kisaran 18,50%–20,20%, sedangkan

saldo transaksi nol berkisar 9,20%–16,30%. Hasil ini mengindikasikan bahwa mayoritas nasabah cenderung mempertahankan saldo utang berjalan dibandingkan mengalami penunggakan kronis.



Gambar 1: Distribusi persentase peubah status pembayaran bulanan (X1–X6).

Karakteristik peubah kontinu yang mencakup jumlah tagihan dan pembayaran aktual dievaluasi menggunakan visualisasi boxplot pada Gambar 2. Hasil pemeriksaan menunjukkan adanya distribusi menceng ke kanan (*right-skewed*) serta keberadaan pencilan ekstrem sehingga diperlukan pendekatan klasterisasi yang *robust* agar estimasi parameter tidak terdistorsi oleh observasi ekstrem.



Gambar 2: Visualisasi boxplot peubah representatif: X7–X12(a) dan X13–X18(b).

Hasil analisis korelasi Pearson pada Tabel 2 menunjukkan korelasi internal yang sangat kuat pada peubah jumlah tagihan bulanan (X7–X12) dengan rentang 0,80–0,95, serta sedang hingga kuat pada peubah status pembayaran (X1–X6) dengan rentang 0,47–0,82.

Tabel 2: Ringkasan hasil korelasi Pearson.

Kelompok Peubah	Korelasi Internal
Status pembayaran (X1-X6)	0,47-0,82
Jumlah tagihan (X7-X12)	0,80-0,95
Jumlah pembayaran (X13-X18)	0,15-0,31

3.2 Pemodelan *Gaussian Mixture Model* (GMM)

Pemodelan GMM dilakukan untuk membentuk segmentasi nasabah berdasarkan pola perilaku pembayaran dan transaksi kredit. Hasil evaluasi seluruh konfigurasi jumlah kluster disajikan pada Tabel 3.

Tabel 3: Metrik evaluasi model GMM.

K	BIC	ICL	<i>Silhouette</i>	<i>Calinski</i>	<i>Davies</i>
2	1.004.411,26	1.004.411,26	0,32	11.524,9 1	1,35
3	969.554,63	969.559,52	0,23	6.661,78	2,19
4	956.499,92	958.332,82	0,25	8.064,48	1,86
5	949.479,78	951.321,19	0,22	6.243,65	2,17
6	937.031,18	940.096,61	0,22	6.230,10	1,92
7	937.031,180	893.375,26	0,21	5.373,04	2,09
8	734.957,05	738.012,36	0,12	4.142,17	2,44
9	871.821,54	874.445,68	0,14	4.127,27	2,09
10	740.094,69	746.121,40	0,12	3.593,24	2,27

Hasil evaluasi pada Tabel 3 menunjukkan nilai BIC dan ICL minimum diperoleh pada $K = 8$, namun konfigurasi $K = 4$ dipilih sebagai K optimal karena memberikan keseimbangan terbaik antara *goodness-of-fit*, kualitas separasi, kestabilan, dan interpretabilitas hasil, dengan *Silhouette Score* tertinggi (0,25), *Calinski-Harabasz* sebesar 8.064,48, dan *Davies-Bouldin* sebesar 1,86. Model GMM $K = 4$ menghasilkan distribusi anggota kluster sebesar 18,98%, 7,69%, 58,85%, dan 14,48%. Profil kluster dihitung berdasarkan data asli menggunakan rata-rata peubah status pembayaran, jumlah tagihan, dan jumlah pembayaran. Peluang keterlambatan dihitung berdasarkan proporsi status pembayaran dengan nilai ≥ 1 , sedangkan rata-rata tagihan dan pembayaran diperoleh dari rata-rata peubah tagihan dan pembayaran pada setiap kluster. Hasil profil kluster disajikan pada Tabel 4.

Tabel 4: Profil kluster model GMM.

Kluster	Peluang terlambat (%)	Rataan Jumlah tagihan	Rataan Jumlah Pembayaran
1	5,32	11.042,26	7.150,30
2	12,90	216.056,16	12.648,35
3	18,98	43.238,07	4.100,00
4	5,03	5.659,89	3.678,06

Berdasarkan profil klaster GMM, terbentuk empat kelompok nasabah dengan karakteristik berbeda. Klaster 1 menunjukkan penggunaan kredit rendah dengan peluang keterlambatan kecil. Klaster 2 memiliki rata-rata tagihan dan pembayaran tertinggi yang mencerminkan aktivitas kredit sangat aktif, meskipun peluang keterlambatannya lebih tinggi. Klaster 3 memiliki rata-rata pembayaran relatif rendah dibandingkan tagihan serta peluang keterlambatan tertinggi, sehingga menunjukkan pola pembayaran kurang stabil. Sementara itu, klaster 4 merupakan kelompok paling disiplin dengan rata-rata tagihan, pembayaran, dan peluang keterlambatan terendah.

3.3 Pemodelan *Robust Trimmed Gaussian Mixture Model* (RTGMM)

Pemodelan RTGMM dilakukan setelah pemodelan GMM sebagai pengembangan pemodelan yang lebih *robust* terhadap keberadaan pencilon pada data. Hasil evaluasi seluruh kombinasi parameter RTGMM disajikan pada Tabel 5.

Tabel 5: Metrik evaluasi model RTGMM.

α	K	BIC	ICL	<i>Silhouette</i>	<i>Calinski</i>	<i>Davies</i>
0,05	2	615.576,73	615.577,37	0,32	11.524,46	1,36
	3	581.937,96	581.950,93	0,23	6.665,16	2,19
	4	567.610,53	567.823,02	0,22	4.877,32	2,63
	5	317.321,30	318.721,10	0,12	5.558,61	2,85
	6	309.069,12	310.468,56	0,13	4.991,35	2,70
	7	503.010,29	504.488,11	0,23	4.636,26	2,35
	8	301.572,51	303.929,41	0,11	3.632,32	2,79
	9	472.108,98	474.583,78	0,18	4.082,21	2,46
	10	246.945,53	249.331,10	0,14	3.401,82	2,54
	2	500.012,28	500.012,87	0,32	11.529,86	1,36
0,1	3	468.157,19	468.173,77	0,22	6.659,18	2,19
	4	454.166,23	454.339,35	0,22	4.873,98	2,63
	5	433.999,18	436.479,21	0,18	6.592,62	1,91
	6	186.209,18	187.617,13	0,13	4.975,84	2,70
	7	201.558,19	205.443,01	0,11	4.500,50	2,44
	8	193.944,36	197.190,99	0,09	4.136,99	2,64
	9	361.215,11	363.723,28	0,17	4.106,97	2,48
	10	129.513,15	132.294,76	0,14	3.807,53	2,51

Pemilihan konfigurasi optimal mempertimbangkan BIC dan ICL sebagai metrik utama, didukung *Silhouette Score*, *Calinski-Harabasz*, dan *Davies-Bouldin* sebagai evaluasi kualitas struktur klaster. Berdasarkan Tabel 5, nilai BIC dan ICL minimum diperoleh pada $K = 10$ dan $\alpha = 0,10$, namun konfigurasi tersebut tidak diikuti peningkatan kualitas klaster, terlihat dari *Silhouette* yang cenderung menurun dan *Davies-Bouldin* yang tidak stabil. Oleh karena itu, model RTGMM dengan $K = 5$ dan $\alpha = 0,10$ dipilih sebagai konfigurasi optimal dengan BIC = 433.999,18, ICL = 436.479,21, dan *Davies-Bouldin* = 1,91. Konfigurasi ini menunjukkan keseimbangan antara *goodness-of-fit*, kualitas pemisahan klaster, dan

interpretabilitas hasil, sekaligus menghindari fragmentasi segmen yang berlebihan. Model ini menghasilkan distribusi anggota klaster yang relatif proporsional, yaitu masing-masing sebesar 19,14%, 13,11%, 4,63%, 46,80% dan 16,31%. Hasil profil klaster disajikan pada tabel 6.

Tabel 6: Profil klaster model RTGMM.

Klaster	Peluang terlambat (%)	Rataan Jumlah tagihan	Rataan Jumlah Pembayaran
1	5,76	11.977,40	6.852,61
2	15,55	117.142,52	8.297,02
3	13,31	246.179,21	15.836,50
4	19,69	31.454,91	3.469,44
5	5,71	7.326,38	3.176,31

RTGMM membentuk lima kelompok nasabah dengan karakteristik yang lebih rinci. Klaster 1 dan klaster 5 merepresentasikan nasabah dengan penggunaan kredit relatif rendah serta peluang keterlambatan yang rendah sehingga menunjukkan pola pembayaran yang cukup disiplin. Klaster 2 dan klaster 3 merupakan kelompok dengan rata-rata tagihan dan pembayaran sangat tinggi, yang mencerminkan intensitas penggunaan kredit yang aktif, namun peluang keterlambatannya cukup tinggi. Klaster 3 secara khusus memiliki rata-rata tagihan dan pembayaran tertinggi di antara seluruh klaster. Sementara itu, klaster 4 memiliki rata-rata pembayaran relatif rendah dibandingkan tagihan serta peluang keterlambatan tertinggi, sehingga mencerminkan pola pembayaran yang kurang stabil dibandingkan klaster lainnya.

3.4 Evaluasi dan Perbandingan Model

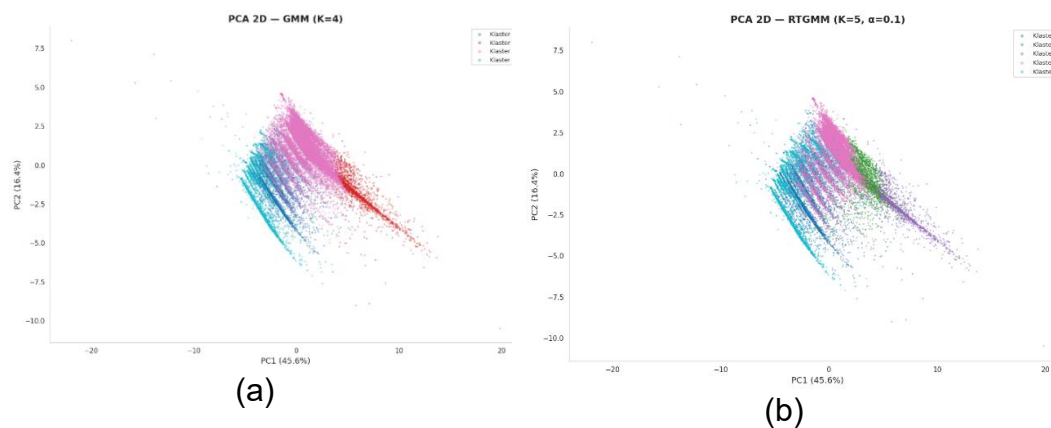
Evaluasi performa model GMM dan RTGMM dilakukan menggunakan lima metrik yang disajikan pada Tabel 7 berikut.

Tabel 7: Perbandingan metrik evaluasi model GMM dan RTGMM.

Metrik	GMM ($K = 4$)	RTGMM ($K = 5, \alpha = 0,1$)
BIC	956.499,89	433.999,18
ICL	958.332,64	436.479,21
<i>Silhouette</i>	0,25	0,18
<i>Calinski-Harabasz</i>	8.064,48	6.592,62
<i>Davies-Bouldin</i>	1,86	1,91

Hasil evaluasi menunjukkan bahwa RTGMM menghasilkan nilai BIC dan ICL lebih rendah dibanding GMM, mengindikasikan keseimbangan *goodness-of-fit* dan kompleksitas parameter lebih optimal serta estimasi lebih stabil terhadap pencilaan. Di sisi lain, GMM lebih unggul pada metrik validitas struktur klaster, yaitu *Silhouette*, *Calinski-Harabasz*, dan *Davies-Bouldin*, meskipun selisihnya tidak substansial. Dengan demikian, RTGMM tetap menghasilkan struktur klaster memadai meskipun data mengandung

pencilan. Visualisasi PCA 2D pada Gambar 3 menunjukkan bahwa kedua model masih memiliki overlap antarklaster yang mencerminkan sifat *soft clustering* berbasis probabilistik. Secara kuantitatif, GMM menghasilkan satu klaster dominan sebesar 58,85%, sedangkan RTGMM menghasilkan distribusi klaster lebih seimbang dengan dominasi klaster terbesar sebesar 46,80%. Kondisi ini mengindikasikan bahwa mekanisme *trimming* pada RTGMM mampu mengurangi pengaruh observasi ekstrem sehingga struktur klaster lebih stabil, meskipun GMM masih menunjukkan kualitas separasi klaster sedikit lebih baik.



Gambar 3: Visualisasi hasil pengelompokan menggunakan PCA 2D: (a) GMM dan (b) RTGMM

Secara umum, GMM unggul dalam kemudahan implementasi, efisiensi komputasi, dan separasi klaster, namun lebih rentan terhadap pencilan yang dapat menurunkan representativitas segmen. Sebaliknya, RTGMM menghasilkan estimasi parameter lebih *robust* dan struktur klaster lebih stabil melalui mekanisme *trimming*, meskipun memiliki kompleksitas komputasi lebih tinggi serta memerlukan penentuan parameter α . Oleh karena itu, RTGMM lebih sesuai digunakan pada data perilaku pembayaran kartu kredit yang mengandung pencilan dan distribusi non-ideal Gaussian.

4 Simpulan dan Saran

Penelitian ini menerapkan GMM dan RTGMM untuk segmentasi nasabah menggunakan *Default of Credit Card Clients Dataset*. Hasil eksplorasi menunjukkan adanya distribusi menceng, korelasi antarpeubah, dan pencilan pada peubah tagihan serta pembayaran sehingga dilakukan perbandingan antara GMM dan RTGMM. Hasil menunjukkan GMM optimal diperoleh pada $K = 4$, sedangkan RTGMM optimal pada $K = 5$ dengan proporsi *trimming* $\alpha = 0,10$. RTGMM menghasilkan nilai BIC dan ICL yang lebih rendah serta lebih stabil terhadap pencilan, sementara GMM menunjukkan kualitas separasi klaster yang sedikit lebih baik. Penelitian selanjutnya dapat mempertimbangkan variasi proporsi *trimming* dan perbandingan dengan metode *robust clustering* lainnya.

Daftar Pustaka

- Astina, P. N. P., Supriana, I. W., & Bimantara, I. M. S. (2026). Perbandingan metode clustering K-Means, GMM, dan DBSCAN berbasis fitur RFM. *Jurnal Nasional Teknologi Informasi dan Aplikasinya*, 4(2): 295–306.
- Bishop, C. M. (2006). *Pattern Recognition and Machine Learning*. Springer.
- Bussmann, N., Giudici, P., Marinelli, D., & Papenbrock, J. (2021). Explainable Machine Learning in Credit Risk Management. *Computational Economics*, 57(1), 203–216. <https://doi.org/10.1007/s10614-020-10042-0>
- Coretto, P., & Hennig, C. (2016). Robust Improper Maximum Likelihood: Tuning, Computation, and a Comparison With Other Methods for Robust Gaussian Clustering. *Journal of the American Statistical Association*, 111(516), 1648–1659. <https://doi.org/10.1080/01621459.2015.1100996>
- Cuesta-Albertos, J. A., Matrán, C., & Mayo-Iscar, A. (2008). Robust estimation in the normal mixture model based on robust clustering. *Journal of the Royal Statistical Society. Series B: Statistical Methodology*, 70(4), 779–802. <https://doi.org/10.1111/j.1467-9868.2008.00657.x>
- Fadhli, S., & Hendri, E. P. (2025). Time Series Clustering of Rice Productivity Using Trimming Gaussian Mixture Models. *Parameter: Jurnal Matematika, Statistika dan Terapannya*, 4(3), 381–394. <https://doi.org/10.30598/parameter.v4i3pp381-394>
- Fraley, C., & Raftery, A. E. (2002). Model-based clustering, discriminant analysis, and density estimation. *Journal of the American Statistical Association*, 97(458), 611–631. <https://doi.org/10.1198/016214502760047131>
- García-Escudero, L. A., & Mayo-Iscar, A. (2024). Robust clustering based on trimming. *Wiley Interdisciplinary Reviews: Computational Statistics*, 16(4), 1–17. <https://doi.org/10.1002/wics.1658>
- García-Escudero, L. A., Mayo-Iscar, A., & Riani, M. (2022). Constrained parsimonious model-based clustering. *Statistics and Computing*, 32(1), 1–15. <https://doi.org/10.1007/s11222-021-10061-3>
- Liu, D., Li, S., Yu, Y., & Moustaki, I. (2020). Assessing partial association between ordinal variables: Quantification, visualization, and hypothesis testing. *Journal of the American Statistical Association*, 0(0), 1–14. <https://doi.org/10.1080/01621459.2020.1796394>

- Maharana, K., Mondal, S., & Nemade, B. (2022). A review: Data pre-processing and data augmentation techniques. *Global Transitions Proceedings*, 3(1), 91–99. <https://doi.org/10.1016/j.gltp.2022.04.020>
- Medina, L., Kreutzmann, A. K., Rojas-Perilla, N., & Castro, P. (2019). The R package trafo for transforming linear regression models. *The R Journal*, 11(2), 512–532. <https://digitalcommons.unl.edu/r-journal/210/>
- Melese, T., Berhane, T., Mohammed, A., & Walelgn, A. (2023). Credit-Risk Prediction Model Using Hybrid Deep—Machine-Learning Based Algorithms. *Scientific Programming*, 2023, 1–13. <https://doi.org/10.1155/2023/6675425>
- Robisco, A., & Martínez, J. M. C. (2022). Measuring the model risk-adjusted performance of machine learning algorithms in credit default prediction. *Financial Innovation*, 8(1), 1–35. <https://doi.org/10.1186/s40854-022-00366-1>
- Shatwan, A. B., Abdalla, A., Mohammed, H., Ismaeil, B. E., & Mami, A. M. (2024). On Model Selection Criterion for Finite Gaussian Mixture Models. *International Journal of Sciences: Basic and Applied Research (IJSBAR)*, 73(1), 68–89.
- Yeh, I. C., & Lien, C. H. (2009). The comparisons of data mining techniques for the predictive accuracy of probability of default of credit card clients. *Expert Systems with Applications*, 36(2), 2473–2480. <https://doi.org/10.1016/j.eswa.2007.12.020>