

PENERAPAN ALGORITMA GAUSSIAN NAIVE BAYES UNTUK KLASIFIKASI TINGKAT RISIKO PELANGGAN POTENSIAL PENUNGGAK TAGIHAN LISTRIK PADA DATA PLN

Ari Ahmad Dahril^{1*}, Rizki Yusliana Bakti², Muhammad Faisal³

^{1,2,3} Universitas Muhammadiyah Makassar

¹105841111822@student.unismuh.ac.id, ²rizkiyusliana@unismuh.ac.id,

³muhfaisal@unismuh.ac.id

* Corresponding Author

Info Article

Submission : 08-08-2026
Review : 11-08-2026
Revised : 29-08-2026
Accepted : 02-09-2026
Published : 05-09-2026

Keywords:

Gaussian Naive Bayes,
Electricity Payment
Delinquency Risk, ANOVA F-
test, Bayesian Optimization,
Class Imbalance, PLN.

Abstrak

Penunggakan pembayaran tagihan listrik pascabayar merupakan salah satu permasalahan yang berdampak pada arus kas PT Perusahaan Listrik Negara (Persero). Mekanisme penanganan tunggakan yang masih bersifat reaktif menyebabkan biaya operasional yang tinggi sehingga diperlukan pendekatan prediktif berbasis data untuk mengidentifikasi pelanggan berisiko sejak dini. Penelitian ini menerapkan algoritma Gaussian Naive Bayes untuk mengklasifikasikan tingkat risiko pelanggan listrik pascabayar PLN UP3 Makassar Selatan ke dalam tiga kategori, yaitu Rendah, Sedang, dan Tinggi, berdasarkan pola historis pembayaran. Data riwayat pembayaran diagregasi menjadi 670 data pelanggan dengan delapan atribut prediktor. Tahapan penelitian meliputi preprocessing data, standardisasi menggunakan StandardScaler, seleksi fitur dengan SelectKBest (ANOVA F-test), pembagian data 80:20 menggunakan Stratified Split, optimasi hyperparameter *var_smoothing* melalui Bayesian Optimization, serta evaluasi akhir menggunakan 5-Fold Stratified Cross Validation. Hasil pengujian menunjukkan model memperoleh akurasi sebesar 85,07% dan Macro F1-score sebesar 0,79, dengan precision, recall, dan F1-score tertinggi pada kelas Rendah (0,90; 0,84; 0,87) dan performa lebih rendah pada kelas Tinggi (0,74; 0,57; 0,64). Rendahnya recall kelas Tinggi menunjukkan bahwa ketidakseimbangan kelas merupakan batasan utama penelitian karena 43% pelanggan yang sebenarnya termasuk risiko tinggi belum berhasil teridentifikasi. Oleh karena itu, penerapan teknik penanganan class imbalance, seperti resampling atau cost-sensitive learning, diperlukan sebelum model dipertimbangkan untuk implementasi pada skala produksi. Atribut jumlah pembayaran tepat waktu, jumlah keterlambatan, dan status pembayaran bulan terakhir merupakan fitur yang paling berpengaruh dalam proses klasifikasi.

How to Cite : Dahril, A. A., Bakti, R. Y., & Faisal, M. (2026). Penerapan Algoritma Gaussian Naive Bayes Untuk Klasifikasi Tingkat Risiko Pelanggan Potensial Penunggak Tagihan Listrik Pada Data PLN. *Journal of Computer Science and Information Technology (JCSIT)*, 3(4), 323-333

DOI : 10.70248/jcsit.v3i4.4434

This is an open access article under <https://creativecommons.org/licenses/by/4.0/>
Copyright© 2026 by Author. Published by Yayasan Nuraini Ibrahim Mandiri



PENDAHULUAN

Energi listrik merupakan elemen fundamental yang menopang hampir seluruh aktivitas kehidupan modern (Batubara et al., 2022). Di Indonesia, PT Perusahaan Listrik Negara (Persero) atau PLN memegang mandat sebagai penyedia layanan kelistrikan nasional yang menyuplai energi ke seluruh wilayah (Mustriyanto et al., 2022). Dalam mendistribusikan energi, PLN menerapkan skema layanan pascabayar yang memberikan kepercayaan kepada pelanggan untuk mengonsumsi energi terlebih

dan melunasi tagihannya pada bulan berikutnya. Skema ini memberikan fleksibilitas akses, namun menuntut tingkat kedisiplinan pembayaran yang tinggi dari pelanggan.

Pada praktiknya, tingginya angka penunggakan pembayaran rekening listrik menimbulkan kendala finansial yang signifikan bagi perusahaan. Percepatan pendapatan dari piutang rekening listrik merupakan salah satu upaya krusial bagi perusahaan energi untuk menjaga arus kas (*cash flow*) agar kegiatan operasional dan investasi pengembangan aset dapat terus berjalan (Nugraha et al., 2022). Ketika terjadi tunggakan yang masif, PLN mengalami defisit kas yang seharusnya digunakan untuk biaya pembangkitan dan pemeliharaan jaringan.

Mekanisme penyelesaian tunggakan yang diterapkan PLN saat ini masih bersifat reaktif. Pelanggan yang belum melunasi tagihan pada rentang tanggal tertentu akan diberikan surat peringatan, yang kemudian diikuti dengan tindakan Penertiban Pemakaian Tenaga Listrik (P2TL) berupa pemutusan aliran sementara hingga pembongkaran meteran secara permanen (Siregar, 2023). Tindakan reaktif ini terbukti tidak efisien karena menuntut alokasi biaya operasional yang besar untuk menerjunkan petugas ke lapangan. Di sisi lain, PLN memiliki volume data transaksi historis yang sangat besar, namun pemanfaatannya untuk langkah preventif masih terbatas (Maipasha, 2025).

Evolusi teknologi informasi menawarkan solusi melalui konsep Knowledge Discovery in Databases (KDD) atau data mining untuk mengekstraksi pola tersembunyi dari data historis pelanggan (Maipasha, 2025). Dengan menganalisis variabel seperti jumlah pembayaran tepat waktu, jumlah keterlambatan, dan riwayat pembayaran lainnya, perusahaan dapat memprediksi tingkat risiko pelanggan dan mengelompokkannya ke dalam kategori risiko rendah, sedang, atau tinggi untuk mendukung pengambilan keputusan yang lebih efektif (Batubara et al., 2022). Gaussian Naive Bayes dipilih pada penelitian ini karena memiliki proses komputasi yang sederhana dan efisien serta sesuai untuk mengolah atribut numerik hasil ekstraksi fitur riwayat pembayaran pelanggan.

Beberapa penelitian terdahulu telah menerapkan algoritma Naive Bayes pada permasalahan yang berkaitan dengan prediksi perilaku pembayaran. Algoritma Naive Bayes terbukti efektif memprediksi keterlambatan pembayaran sewa berdasarkan data historis penyewa (Rachman & Handayani, 2021). Varian Gaussian Naive Bayes juga berhasil diimplementasikan untuk segmentasi piutang pajak Perusahaan (Evardo, 2024), serta digunakan untuk memprediksi kategori nasabah potensial guna mendukung strategi peningkatan laba perusahaan (Fitrianah et al., 2021). Selain itu, Naive Bayes menunjukkan performa kompetitif untuk prediksi risiko kredit dibandingkan Random Forest (Prayesy et al., 2025), sedangkan penelitian pada prioritas pembayaran menunjukkan bahwa Naive Bayes dapat digunakan untuk menentukan prioritas tagihan (Rahayu et al., 2021). Studi-studi tersebut menunjukkan bahwa Naive Bayes relevan untuk permasalahan klasifikasi yang berkaitan dengan perilaku pembayaran dan risiko gagal bayar, sehingga mendukung pemilihan Gaussian Naive Bayes pada penelitian ini.

Urgensi penelitian ini didasarkan pada kebutuhan akan instrumen pendukung keputusan yang dapat berperan sebagai *Early Warning System*. Belum tersedianya sistem yang secara otomatis mengklasifikasikan pelanggan berdasarkan tingkat risiko menghambat PLN dalam mengambil langkah mitigasi dini (Mustriyanto et al., 2022). Berdasarkan permasalahan tersebut, penelitian ini bertujuan menerapkan algoritma *Gaussian Naive Bayes* untuk mengklasifikasikan tingkat risiko pelanggan listrik

pascabayar ke dalam tiga kategori, yaitu Rendah, Sedang, dan Tinggi, berdasarkan pola historis pembayaran, serta mengevaluasi kinerja model menggunakan *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian diharapkan dapat menjadi referensi bagi PLN dalam mengembangkan sistem identifikasi dini terhadap pelanggan yang berpotensi menunggak.

METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif berbasis machine learning dengan tujuan membangun model klasifikasi tingkat risiko tunggakan pelanggan listrik pascabayar. Data penelitian bersumber dari riwayat pembayaran pelanggan PLN UP3 Makassar Selatan yang telah dianonimkan sehingga tidak memuat identitas pribadi pelanggan. Data mentah yang diperoleh masih berupa catatan transaksi bulanan sehingga satu pelanggan dapat memiliki beberapa baris data. Data tersebut kemudian diagregasi berdasarkan identitas pelanggan (IDPEL) sehingga setiap pelanggan direpresentasikan oleh satu baris data (one customer one record), dengan nilai atribut numerik dihitung menggunakan rata-rata (mean) dan label risiko ditentukan berdasarkan kategori yang paling sering muncul (majority voting). Pendekatan majority voting digunakan untuk memperoleh satu label yang konsisten pada tingkat pelanggan. Namun, pendekatan ini memiliki keterbatasan karena memberikan bobot yang sama terhadap seluruh riwayat pembayaran. Akibatnya, keterlambatan yang terjadi pada bulan-bulan terakhir berpotensi tertutupi oleh status pembayaran tepat waktu pada periode sebelumnya. Kondisi tersebut dapat menyebabkan perubahan perilaku pembayaran terbaru belum sepenuhnya tercermin pada label risiko pelanggan. Keterbatasan ini menjadi pertimbangan untuk pengembangan penelitian berikutnya melalui pendekatan pelabelan yang lebih mempertimbangkan recency.

Proses agregasi menghasilkan dataset akhir berjumlah 670 data pelanggan dengan delapan atribut prediktor dan satu atribut target berupa kategori risiko Rendah, Sedang, dan Tinggi.

Tabel 1. Atribut yang digunakan dalam penelitian

No.	Atribut	Keterangan
1	JML_TEPAT_WAKTU	Jumlah pembayaran tepat waktu
2	JML_TERLAMBAT	Jumlah keterlambatan pembayaran
3	JML_TELAT_BULAN	Jumlah bulan menunggak
4	RATA_HARI_TERLAMBAT	Rata-rata hari keterlambatan
5	MAKS_HARI_TERLAMBAT	Maksimal hari keterlambatan
6	TELAT_BERUNTUN	Jumlah keterlambatan beruntun
7	MAKS_TELAT_BERUNTUN	Maksimal keterlambatan beruntun
8	STATUS_BULAN_TERAKHIR	Status pembayaran pada bulan terakhir
-	LABEL_RISIKO (target)	Kategori risiko: Rendah, Sedang, atau Tinggi

Label risiko pelanggan dikategorikan menjadi tiga kelas berdasarkan kondisi aktual riwayat pembayaran. Kategori Rendah diberikan kepada pelanggan dengan perilaku pembayaran yang relatif tepat waktu. Kategori Sedang diberikan kepada pelanggan yang mengalami keterlambatan, namun pembayaran masih dilakukan dalam periode bulan berjalan. Kategori Tinggi diberikan kepada pelanggan yang mengalami keterlambatan hingga melewati periode bulan pembayaran atau menunjukkan adanya tunggakan.

Tahapan *preprocessing* dilakukan sebelum pelatihan model, meliputi pemeriksaan *missing value* dan duplikasi data, penyeragaman format kolom, pembentukan atribut berdasarkan riwayat pembayaran, standardisasi data, serta seleksi fitur. Ringkasan hasil *preprocessing* ditampilkan pada Tabel 2.

Tabel 2. Ringkasan tahapan preprocessing data

Tahap Preprocessing	Hasil
Pemeriksaan missing value	Tidak ditemukan nilai kosong
Pemeriksaan data duplikat	Tidak ditemukan data duplikat
Agregasi data (one customer one record)	Data digabung berdasarkan IDPEL; atribut numerik dirata-ratakan, label ditentukan melalui majority voting
Standardisasi data	StandardScaler, mean = 0 dan standar deviasi = 1
Seleksi fitur	SelectKBest dengan ANOVA F-test ($k = 3$); nilai k dipilih berdasarkan tiga atribut dengan skor F tertinggi. Validasi ablasi $k=1-8$ perlu ditambahkan pada eksperimen lanjutan.

Dataset yang telah diproses selanjutnya dibagi menggunakan fungsi `train_test_split` pada pustaka Scikit-learn dengan metode Stratified Split, sehingga proporsi setiap kelas risiko tetap terjaga pada data latih maupun data uji. Pembagian data dilakukan dengan perbandingan 80% data latih dan 20% data uji. Pada penelitian ini, teknik resampling seperti SMOTE atau Random Over-sampling tidak diterapkan pada eksperimen utama karena penelitian diarahkan untuk terlebih dahulu mengukur performa dasar Gaussian Naive Bayes pada distribusi pelanggan yang sebenarnya. Dengan demikian, hasil evaluasi merepresentasikan kondisi ketidakseimbangan kelas sebagaimana terdapat pada dataset asli. Keputusan tersebut sekaligus menjadi batasan penelitian karena kelas Tinggi hanya terdiri atas 30 pelanggan (4,48%), sehingga kemampuan model dalam mengenali kelas minoritas menjadi terbatas. Teknik resampling dan pendekatan cost-sensitive learning direkomendasikan sebagai eksperimen lanjutan untuk menekan False Negative pada kelas Tinggi.

Tabel 3. Pembagian data latih dan data uji

Jenis Data	Persentase	Jumlah Data
Data Latih (Training Data)	80%	536
Data Uji (Testing Data)	20%	134
Total	100%	670

Model klasifikasi dibangun menggunakan algoritma Gaussian Naive Bayes, salah satu varian Naive Bayes yang digunakan untuk mengolah data numerik dengan asumsi setiap atribut dalam masing-masing kelas mengikuti pendekatan distribusi Gaussian dan bersifat independen antaratribut (Khoirunnisa, 2025). Algoritma ini bekerja berdasarkan Teorema Bayes, yang menghitung probabilitas posterior suatu kelas berdasarkan probabilitas prior dan likelihood dari atribut yang dimiliki data (Rachman & Handayani, 2021). Gaussian Naive Bayes dipilih pada penelitian ini karena seluruh atribut yang digunakan berupa data numerik hasil pengolahan riwayat pembayaran pelanggan, berbeda dengan varian Multinomial atau Bernoulli Naive Bayes yang lebih sesuai untuk data diskrit atau biner (Mahendra et al., 2025). Penggunaan Gaussian Naive Bayes tidak didasarkan pada asumsi bahwa seluruh data mentah harus berdistribusi normal secara sempurna; namun, pemeriksaan skewness dan kurtosis pada setiap atribut numerik perlu dilaporkan sebagai bukti empiris kesesuaian distribusi.

Pemeriksaan karakteristik distribusi dilakukan menggunakan nilai skewness dan kurtosis terhadap delapan atribut prediktor sebelum proses pemodelan. Hasil pemeriksaan menunjukkan bahwa sebagian besar atribut memiliki nilai skewness yang relatif rendah, seperti `JML_TEPAT_WAKTU` sebesar 0,2023, `JML_TERLAMBAT` sebesar 0,1928, `RATA_HARI_TERLAMBAT` sebesar 0,2655, `MAKS_HARI_TERLAMBAT` sebesar 0,1719, `TELAT_BERUNTUN` sebesar 0,0669, dan `MAKS_TELAT_BERUNTUN` sebesar

STATUS_BULAN_TERAKHIR memiliki skewness sebesar -0,7806, sedangkan JML_TELAT_BULAN memiliki skewness sebesar 1,0883 yang menunjukkan kemencengan positif lebih tinggi. Nilai kurtosis seluruh atribut berada pada rentang -1,0044 hingga 0,6441. Hasil tersebut menunjukkan bahwa sebagian besar atribut memiliki karakteristik distribusi yang relatif moderat, meskipun tidak seluruh atribut sepenuhnya mengikuti distribusi normal. Oleh karena itu, asumsi Gaussian pada Gaussian Naive Bayes digunakan sebagai pendekatan pemodelan dan kinerja model selanjutnya diverifikasi secara empiris menggunakan 5-Fold Stratified Cross Validation.

Tabel 4. Tabel Skewness dan Kurtosis

Atribut	Skewness	Kurtosis
JML_TEPAT_WAKTU	0,2023	-1,0044
JML_TERLAMBAT	0,1928	-0,7015
JML_TELAT_BULAN	1,0883	0,6441
RATA_HARI_TERLAMBAT	0,2655	-0,4670
MAKS_HARI_TERLAMBAT	0,1719	-0,4075
TELAT_BERUNTUN	0,0669	-0,4966
MAKS_TELAT_BERUNTUN	0,2483	-0,7780
STATUS_BULAN_TERAKHIR	-0,7806	-0,1539

Sebelum pelatihan, dilakukan standardisasi data menggunakan StandardScaler agar seluruh atribut numerik memiliki skala yang setara. Selanjutnya dilakukan seleksi fitur menggunakan SelectKBest dengan uji ANOVA F-test (f_{classif}) untuk memilih tiga atribut ($k = 3$) yang memiliki hubungan paling kuat terhadap label risiko. Berdasarkan hasil seleksi, tiga atribut tersebut adalah JML_TEPAT_WAKTU, JML_TERLAMBAT, dan STATUS_BULAN_TERAKHIR. Pemilihan $k = 3$ didasarkan pada peringkat skor F tertinggi dari delapan atribut. Namun, untuk memenuhi pengujian ablasional yang disarankan reviewer, performa model idealnya dibandingkan pada $k = 1$ sampai $k = 8$ dan disajikan dalam grafik accuracy atau Macro F1-score. Hasil numerik variasi k tersebut belum tersedia pada eksperimen yang tersimpan sehingga tidak direka atau ditambahkan sebagai hasil penelitian. Optimasi hyperparameter `var_smoothing` dilakukan menggunakan Bayesian Optimization melalui algoritma BayesSearchCV pada ruang pencarian 10^{-12} hingga 10^0 dengan distribusi log-uniform, dilakukan sebanyak 30 iterasi menggunakan 5-fold cross validation. Proses ini menghasilkan nilai `var_smoothing` terbaik sebesar 0,011185625 dengan akurasi cross validation sebesar 85,64%. Seluruh tahapan standardisasi, seleksi fitur, dan klasifikasi disatukan dalam sebuah pipeline untuk menghindari data leakage.

Evaluasi akhir kinerja model dilakukan menggunakan metode 5-Fold Stratified Cross Validation, yaitu membagi dataset menjadi lima bagian dengan proporsi kelas yang tetap terjaga pada setiap fold. Pada setiap iterasi, satu fold digunakan sebagai data validasi dan empat fold lainnya digunakan sebagai data pelatihan, sehingga seluruh data memperoleh kesempatan menjadi data validasi. Kinerja model diukur menggunakan confusion matrix, accuracy, precision, recall, F1-score, dan Macro F1-score. Macro F1-score dihitung sebagai rata-rata aritmetika F1-score dari tiga kelas tanpa mempertimbangkan jumlah sampel pada masing-masing kelas, sehingga memberikan gambaran performa model yang lebih seimbang antar-kelas. Accuracy menunjukkan proporsi prediksi benar terhadap seluruh data uji, precision menggambarkan ketepatan model saat memprediksi suatu kelas (Yusuf Ramadhan Nasution et al., 2024), recall menunjukkan kemampuan model menemukan seluruh data aktual suatu kelas (Rachman & Handayani, 2021), dan F1-score digunakan sebagai ukuran gabungan antara precision dan recall (Prayesy et al., 2025).

HASIL DAN PEMBAHASAN

Dataset penelitian yang telah melalui proses agregasi berdasarkan IDPEL menghasilkan 670 data pelanggan dengan distribusi label risiko yang tidak seimbang (imbalanced). Kategori Rendah memiliki jumlah pelanggan terbanyak, yaitu 328 pelanggan (48,96%), diikuti kategori Sedang sebanyak 312 pelanggan (46,57%), sedangkan kategori Tinggi hanya berjumlah 30 pelanggan (4,48%). Ketimpangan ini menunjukkan bahwa kelas Tinggi hanya memiliki sebagian kecil observasi dibandingkan kelas lainnya dan berpotensi menyebabkan model lebih banyak belajar pola kelas mayoritas. Kondisi tersebut menjadi batasan utama penelitian karena kesalahan False Negative pada kelas Tinggi berarti pelanggan yang sebenarnya berisiko tinggi diklasifikasikan sebagai risiko yang lebih rendah.

Tabel 5. Distribusi label risiko pelanggan

Kategori Risiko	Jumlah Pelanggan	Persentase
Rendah	328	48,96%
Sedang	312	46,57%
Tinggi	30	4,48%
Total	670	100%

Hasil seleksi fitur menggunakan ANOVA F-test menunjukkan bahwa atribut JML_TEPAT_WAKTU, JML_TERLAMBAT, dan STATUS_BULAN_TERAKHIR memperoleh skor F tertinggi dibandingkan lima atribut lainnya, sehingga dipilih sebagai tiga fitur ($k = 3$) yang digunakan dalam pembentukan model. Secara interpretabilitas bisnis, STATUS_BULAN_TERAKHIR dapat memiliki hubungan yang kuat dengan risiko karena merepresentasikan kondisi pembayaran paling aktual. JML_TEPAT_WAKTU dan JML_TERLAMBAT juga secara langsung menggambarkan frekuensi perilaku pembayaran pelanggan. Sebaliknya, TELAT_BERUNTUN dan JML_TELAT_BULAN dapat memperoleh skor lebih rendah karena informasi yang dikandungnya sebagian beririsan dengan frekuensi keterlambatan dan status pembayaran, serta tidak selalu membedakan pelanggan antar-kelas secara konsisten. Dengan demikian, skor ANOVA yang lebih rendah tidak berarti fitur tersebut tidak relevan secara bisnis, tetapi menunjukkan bahwa hubungan univariatnya dengan label pada dataset ini lebih lemah dibandingkan tiga fitur teratas.

Tabel 6. Peringkat atribut berdasarkan skor relevansi ANOVA F

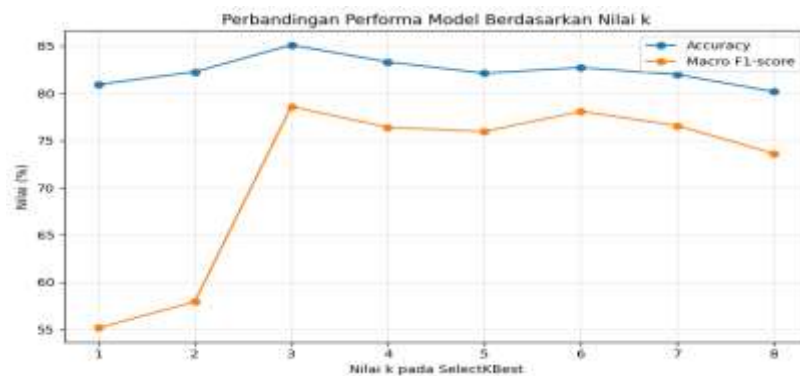
Peringkat	Atribut
1	JML_TEPAT_WAKTU
2	JML_TERLAMBAT
3	STATUS_BULAN_TERAKHIR
4	RATA_HARI_TERLAMBAT
5	JML_TELAT_BULAN
6	TELAT_BERUNTUN
7	MAKS_HARI_TERLAMBAT
8	MAKS_TELAT_BERUNTUN

Untuk memvalidasi pemilihan jumlah fitur yang digunakan dalam model, dilakukan pengujian ablasi dengan memvariasikan nilai k pada SelectKBest dari $k=1$ hingga $k=8$. Setiap nilai k diuji menggunakan pipeline yang sama, yaitu StandardScaler, SelectKBest dengan ANOVA F-test, dan Gaussian Naive Bayes menggunakan 5-Fold Stratified Cross Validation. Pengujian dilakukan untuk mengetahui pengaruh jumlah fitur terhadap performa klasifikasi serta memastikan bahwa pemilihan tiga fitur terbaik didukung oleh hasil eksperimen.

Hasil Pengujian Variasi Nilai k pada SelectKBest

k	Accuracy	Macro F1-score
1	80,90%	55,17%
2	82,24%	57,93%
3	85,07%	78,57%
4	83,28%	76,38%
5	82,09%	75,94%
6	82,69%	78,06%
7	81,94%	76,54%
8	80,15%	73,60%

Berdasarkan hasil pengujian, nilai $k=3$ menghasilkan performa terbaik dengan accuracy sebesar 85,07% dan Macro F1-score sebesar 78,57%. Nilai tersebut lebih tinggi dibandingkan penggunaan seluruh delapan fitur pada $k=8$ yang menghasilkan accuracy sebesar 80,15% dan Macro F1-score sebesar 73,60%. Hasil ini menunjukkan bahwa penggunaan seluruh fitur tidak selalu meningkatkan performa model. Oleh karena itu, tiga fitur dengan skor ANOVA F tertinggi, yaitu JML_TEPAT_WAKTU, JML_TERLAMBAT, dan STATUS_BULAN_TERAKHIR, dipilih sebagai fitur akhir dalam pemodelan karena memberikan performa terbaik berdasarkan hasil pengujian ablasi.



Gambar 1. Perbandingan Performa Model Berdasarkan Nilai K

Sebelum evaluasi akhir dilakukan, proses optimasi *hyperparameter* menghasilkan nilai *var_smoothing* sebesar 0,011185625 dengan akurasi *cross validation* sebesar 85,64%. Nilai tersebut selanjutnya digunakan sebagai parameter pada model akhir yang dievaluasi menggunakan *5-Fold Stratified Cross Validation*.

Hasil evaluasi akhir menunjukkan bahwa model Gaussian Naive Bayes memperoleh akurasi sebesar 85,07%, yang menunjukkan model mampu mengklasifikasikan sekitar 85 dari setiap 100 pelanggan dengan benar. Macro F1-score sebesar 0,79 menunjukkan bahwa performa rata-rata model pada ketiga kelas masih cukup baik, tetapi lebih rendah daripada accuracy karena Macro F1 memberikan bobot yang sama kepada setiap kelas, termasuk kelas Tinggi yang jumlah datanya hanya 30 pelanggan. Rincian nilai precision, recall, dan F1-score pada setiap kelas ditampilkan pada Tabel 8.

Tabel 8. Hasil evaluasi model Gaussian Naive Bayes

Kelas	Precision	Recall	F1-Score	Support
Rendah	0,90	0,84	0,87	328
Sedang	0,81	0,88	0,85	312
Tinggi	0,74	0,57	0,64	30
Accuracy	-	-	0,85	670
Macro F1-Score	-	-	0,79	670

Pada kelas Rendah, model memperoleh precision sebesar 0,90, recall sebesar 0,84, dan F1-score sebesar 0,87, yang menunjukkan bahwa model mampu mengenali pelanggan dengan riwayat pembayaran disiplin secara konsisten. Pada kelas Sedang, precision sebesar 0,81 dan recall sebesar 0,88 menunjukkan bahwa sebagian besar pelanggan risiko sedang berhasil dikenali, meskipun precision yang lebih rendah mengindikasikan adanya pelanggan kategori lain yang diprediksi sebagai kategori sedang akibat karakteristik pembayaran yang relatif mirip. Kelas Tinggi memperoleh precision sebesar 0,74, recall sebesar 0,57, dan F1-score sebesar 0,64, nilai yang lebih rendah dibandingkan dua kelas lainnya karena jumlah data pelanggan kategori Tinggi jauh lebih sedikit. Recall 0,57 berarti 17 dari 30 pelanggan Tinggi berhasil dikenali, sedangkan 13 pelanggan Tinggi atau sekitar 43% masih diklasifikasikan sebagai kelas Sedang. Dalam konteks operasional PLN, False Negative tersebut penting diperhatikan karena pelanggan berisiko tinggi yang tidak terdeteksi sejak awal dapat menyebabkan keterlambatan tindakan penagihan, meningkatnya piutang, serta tambahan kebutuhan operasional untuk penanganan tunggakan. Oleh karena itu, accuracy 85,07% tidak dapat menjadi satu-satunya dasar penilaian kelayakan model.

Pola kesalahan klasifikasi model dianalisis lebih lanjut menggunakan *confusion matrix* hasil agregasi *5-Fold Stratified Cross Validation*, sebagaimana ditampilkan pada Tabel 7. Dari 328 pelanggan aktual kategori Rendah, 277 pelanggan diprediksi benar. Pada kategori Sedang, 276 dari 312 pelanggan diprediksi benar, sedangkan pada kategori Tinggi, 17 dari 30 pelanggan diprediksi benar. Total prediksi benar mencapai 570 dari 670 data (85,07%), sedangkan 100 data mengalami kesalahan prediksi.

Tabel 9. Confusion matrix model Gaussian Naive Bayes

Aktual / Prediksi	Rendah	Sedang	Tinggi
Rendah	277	50	1
Sedang	31	276	5
Tinggi	0	13	17

Kesalahan klasifikasi paling banyak terjadi antara kelas Rendah dan Sedang, dengan 50 pelanggan Rendah diprediksi sebagai Sedang dan 31 pelanggan Sedang diprediksi sebagai Rendah. Kondisi ini dapat dijelaskan karena kedua kelompok pelanggan tersebut memiliki pola pembayaran yang relatif berdekatan. Sementara itu, kelas Tinggi tidak pernah diprediksi sebagai Rendah (0 data), namun 13 dari 30 pelanggan Tinggi diprediksi sebagai Sedang. Kesalahan tersebut merupakan False Negative pada kelas risiko tinggi dan menunjukkan bahwa model masih mengalami kesulitan membedakan pelanggan Tinggi dari Sedang akibat keterbatasan jumlah data kelas Tinggi. Karena konsekuensi operasional False Negative lebih besar daripada kesalahan pada kelas mayoritas, evaluasi lanjutan sebaiknya tidak hanya berfokus pada accuracy, tetapi juga mempertimbangkan recall kelas Tinggi, Macro F1-score, serta pendekatan cost-sensitive learning.

Hasil penelitian ini sejalan dengan studi sebelumnya yang menunjukkan bahwa Naive Bayes dapat digunakan pada permasalahan yang berkaitan dengan perilaku pembayaran dan risiko kredit (Prayesy et al., 2025; Rachman & Handayani, 2021), serta pada segmentasi piutang (Evardo, 2024). Akurasi sebesar 85,07% menunjukkan bahwa Gaussian Naive Bayes mampu mengolah atribut numerik hasil ekstraksi riwayat pembayaran pada dataset penelitian ini. Namun, hasil tersebut harus dipahami bersama dengan Macro F1-score sebesar 0,79 dan recall kelas Tinggi sebesar 0,57. Ketidakseimbangan distribusi kelas Tinggi menjadi faktor penting yang membatasi

kemampuan model dalam mendeteksi pelanggan berisiko tinggi. Oleh karena itu, penerapan teknik resampling, seperti SMOTE atau Random Over-sampling, serta pendekatan cost-sensitive learning perlu diuji pada penelitian berikutnya untuk meningkatkan sensitivitas terhadap kelas Tinggi.

Secara keseluruhan, hasil penelitian menunjukkan bahwa algoritma Gaussian Naive Bayes dapat digunakan sebagai model awal untuk mengklasifikasikan tingkat risiko tunggakan pelanggan listrik pascabayar berdasarkan pola historis pembayaran. Model memiliki accuracy 85,07% dan Macro F1-score 0,79, tetapi recall kelas Tinggi yang hanya 0,57 menunjukkan bahwa model belum cukup sensitif terhadap pelanggan berisiko tinggi. Dengan demikian, model belum sebaiknya langsung diterapkan pada skala produksi sebagai satu-satunya mekanisme identifikasi risiko. Model lebih tepat diposisikan sebagai baseline atau pendukung Early Warning System, dengan syarat dilakukan penanganan class imbalance dan validasi tambahan sebelum implementasi operasional.

KESIMPULAN

Penelitian ini berhasil menerapkan algoritma Gaussian Naive Bayes untuk mengklasifikasikan tingkat risiko tunggakan pelanggan listrik pascabayar berdasarkan delapan atribut riwayat pembayaran ke dalam tiga kategori, yaitu Rendah, Sedang, dan Tinggi. Dataset yang digunakan berjumlah 670 data pelanggan hasil agregasi berdasarkan IDPEL dan telah melalui tahapan preprocessing, standardisasi data, seleksi fitur menggunakan SelectKBest (ANOVA F-test), optimasi hyperparameter var_smoothing melalui Bayesian Optimization, serta evaluasi menggunakan 5-Fold Stratified Cross Validation. Hasil pengujian menunjukkan model memperoleh akurasi sebesar 85,07% dan Macro F1-score sebesar 0,79, dengan performa terbaik pada kelas Rendah (precision 0,90; recall 0,84; F1-score 0,87) dan performa yang lebih rendah pada kelas Tinggi (precision 0,74; recall 0,57; F1-score 0,64). Confusion matrix menunjukkan 17 dari 30 pelanggan kelas Tinggi berhasil dikenali, sedangkan 13 pelanggan atau sekitar 43% masih diklasifikasikan sebagai Sedang. Kondisi tersebut menegaskan bahwa class imbalance merupakan keterbatasan utama penelitian. Dengan demikian, Gaussian Naive Bayes dapat digunakan sebagai baseline dan model pendukung identifikasi awal, tetapi belum dapat dinyatakan siap untuk implementasi produksi tanpa penanganan ketidakseimbangan data. Penelitian selanjutnya disarankan menguji SMOTE atau Random Over-sampling, cost-sensitive learning untuk menekan False Negative pada kelas Tinggi, variasi nilai k pada SelectKBest melalui pengujian ablasional, serta penambahan variabel seperti golongan tarif dan lama berlangganan.

PERNYATAAN PENGGUNAAN KECERDASAN BUATAN

Kecerdasan buatan (AI) digunakan dalam penyusunan naskah ini. Penulis menggunakan alat bantu kecerdasan buatan semata-mata untuk membantu perbaikan bahasa, pemeriksaan tata bahasa, penyuntingan, serta penyusunan naskah agar lebih jelas dan mudah dipahami. Penggunaan AI tidak menggantikan proses analisis dan pengambilan keputusan ilmiah dalam penelitian. Penulis tetap bertanggung jawab penuh atas akurasi, orisinalitas, integritas, dan isi akademik naskah ini. Seluruh konten yang dibantu AI telah ditinjau dan diverifikasi oleh penulis sebelum naskah dikirimkan.

KONFLIK KEPENTINGAN

Penulis menyatakan bahwa tidak terdapat konflik kepentingan dalam penelitian ini.

KONTRIBUSI PENIULIS

Konseptualisasi, Ari Ahmad Dahril; pengumpulan data, pengolahan dan preprocessing data, analisis dan pemodelan Gaussian Naive Bayes, evaluasi model, Rizki Yusliana Bakti dan Muhammad Faisal; penulisan—draf awal, penulisan—revisi dan penyuntingan, Rizki Yusliana Bakti dan Muhammad Faisal; visualisasi, AAD. Seluruh penulis telah membaca dan menyetujui versi akhir naskah ini.

UCAPAN TERIMA KASIH

Penelitian ini didukung oleh Program Studi Informatika, Fakultas Teknik, Universitas Muhammadiyah Makassar, serta PT PLN (Persero) UP3 Makassar Selatan yang telah memberikan akses terhadap data historis pembayaran pelanggan yang digunakan dalam penelitian ini. Penulis menyampaikan terima kasih atas dukungan, fasilitas, dan sumber daya yang diberikan sehingga penelitian ini dapat terlaksana dengan baik.

DAFTAR PUSTAKA

- Batubara, D. N., Windarto, A. P., & Irawan, E. (2022). Analisis Prediksi Keterlambatan Pembayaran Listrik Menggunakan Komparasi Metode Klasifikasi Decision Tree dan Support Vector Machine. *JURIKOM (Jurnal Riset Komputer)*, 9(1), 102. <https://doi.org/10.30865/jurikom.v9i1.3833>
- Evardo, wigi julio. (2024). *Implementasi Metode Gaussian Naive Bayes Untuk Segmentasi Tagihan Pajak Piutang Perusahaan Studi Kasus (KPP Pratama Jakarta Kalideres)*.
- Fitrianah, D., Dwiasnati, S., H, H. H., & Baihaqi, K. A. (2021). Penerapan Metode Machine Learning untuk Prediksi Nasabah Potensial menggunakan Algoritma Klasifikasi Naïve Bayes. *Faktor Exacta*, 14(2), 92. <https://doi.org/10.30998/faktorexacta.v14i2.9297>
- Khoirunnisa, T. (2025). *mplementasi Algoritma Naïve Bayes Untuk Klasifikasi Kesulitan Mahasiswa Dalam Pembelajaran Jarak Jauh (PJJ)*. 2, 306–312.
- Mahendra, I. B., Sunarya, I. M. G., & Wirawan, I. M. A. (2025). Comparison of Multinomial, Bernoulli, and Gaussian Naïve Bayes for Complaint Classification in Pro Denpasar Application. *JUITA: Jurnal Informatika*, 13(1), 77–86. <https://doi.org/10.30595/juita.v13i1.24828>
- Maipasha, B. (2025). *IMPLEMENTASI DATA MINING K-MEANS CLUSTERING TUNGGAKAN REKENING LISTRIK PASCABAYAR (Studi Kasus : PT PLN Persero ULP Tualang)*. 6(0), 167–186.
- Mustriyanto, abiyoga bagus, Habibi, M., Subekti, D., & Fajar, S. (2022). Perbandingan Metode Decision Tree Dan Naive Bayes Classifier Pada Analisis Sentimen Pengguna Layanan Pt Perusahaan Listrik Negara (Pln). *Teknomatika: Jurnal Informatika Dan Komputer*, 15(2), 53–61. <https://doi.org/10.30989/teknomatika.v15i2.1131>
- Nugraha, R. H., Purwitasari, D., & Raharjo, A. B. (2022). K-Means Dan XGBoost Untuk Analisis Perilaku Pembayaran Rekening Listrik Pelanggan (Studi Kasus : PLN ULP PANAKKUKANG). *JUTI: Jurnal Ilmiah Teknologi Informasi*, 20(2), 84–98.
- Prayesy, P., Pujakesuma, A., & Qisthiano, M. R. (2025). Evaluasi Kinerja Random Forest dan Naïve Bayes untuk Prediksi Risiko Kredit Berdasarkan Pekerjaan Debitur.

- 15(2), 114–120. <https://doi.org/10.35200/ex.v15i2.175>
- Rachman, R., & Handayani, R. N. (2021). Klasifikasi Algoritma Naive Bayes Dalam Memprediksi Tingkat Kelancaran Pembayaran Sewa Teras UMKM. *Jurnal Informatika*, 8(2), 111–122. <https://doi.org/10.31294/ji.v8i2.10494>
- Rahayu, woro isti, Prianto, C., & Novia, ema ainun. (2021). *PERBANDINGAN ALGORITMA K-MEANS DAN NAÏVE BAYES UNTUK MEMPREDIKSI PRIORITAS PEMBAYARAN TAGIHAN RUMAH SAKIT BERDASARKAN TINGKAT KEPENTINGAN PADA PT. 13*(2), 1–8.
- Siregar, B. J. (2023). PENYELESAIAN TUNGGAKAN REKENING LISTRIK PELANGGAN MELALUI PELAKSANAAN PENERTIBAN PEMAKAIAN TENAGA LISTRIK (P2TL) OLEH PT. PLN PERSERO. *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, 53(1), 1–19. <http://publications.lib.chalmers.se/records/fulltext/245180/245180.pdf%0Ahttp://hdl.handle.net/20.500.12380/245180%0Ahttp://dx.doi.org/10.1016/j.jsames.2011.03.003%0Ahttps://doi.org/10.1016/j.gr.2017.08.001%0Ahttp://dx.doi.org/10.1016/j.precamres.2014.12>
- Yusuf Ramadhan Nasution, Suhardi Suhardi, & Ilham Hafiz Satrio. (2024). Penerapan Algoritma Klasifikasi Naïve Bayes Untuk Analisis Sentimen Tentang Pemilu 2024. *Elkom: Jurnal Elektronika Dan Komputer*, 17(2), 495–502. <https://doi.org/10.51903/elkom.v17i2.2053>