

Analisis Komparatif *Multinomial Naïve Bayes* dan *Logistic Regression* untuk Klasifikasi Sentimen Ulasan Pengguna Aplikasi TIX ID

Mochamad Miftah Rachmatullah¹, Bambang Irawan², Ahmad Faqih³, Denni Pratama⁴, Dian Ade Kurnia⁵

^{1,2,3}Program Studi Teknik Informatika, STMIK IKMI Cirebon, Indonesia

^{4,5}Program Studi Manajemen Informatika, STMIK IKMI Cirebon, Indonesia

mochamadmiftah34@gmail.com, bambangirawan2000@yahoo.com, ahmadfaqih367@gmail.com,

pratamadenni@gmail.com, dianade2012@gmail.com

Article Info

Article history:

Received November 22, 2025

Accepted November 28, 2025

Published January 1, 2026

Kata Kunci:

Analisis sentimen

Multinomial Naïve Bayes

Logistic Regression

TIX ID

Klasifikasi Multi-kelas

ABSTRAK

Penelitian ini membandingkan performa algoritma *Multinomial Naïve Bayes* (MNB) dan *Logistic Regression* (LR) dalam klasifikasi sentimen multi-kelas pada ulasan pengguna aplikasi TIX ID. Sebanyak 2.500 ulasan dikumpulkan dari Google Play Store dan diproses melalui tahapan *preprocessing*, termasuk pembersihan teks, *case folding*, *tokenisasi*, *stopword removal*, dan *stemming*. Dua metode ekstraksi fitur diterapkan, yaitu *CountVectorizer* dan TF-IDF, sebelum model dilatih menggunakan kedua algoritma. Proses *hyperparameter tuning* dilakukan dengan *GridSearchCV* untuk memperoleh konfigurasi terbaik. Hasil menunjukkan bahwa MNB dengan *CountVectorizer* sebelum *tuning* memberikan performa tertinggi dengan akurasi 84,80% dan *F1-score macro* terbaik. *Hyperparameter tuning* meningkatkan stabilitas model, namun tidak melampaui performa tersebut. Temuan ini menunjukkan bahwa kombinasi MNB dan *CountVectorizer* paling sesuai untuk teks ulasan berbahasa Indonesia yang bersifat *sparse* dan *repetitif*. Model terbaik kemudian diimplementasikan dalam sistem analisis sentimen berbasis web untuk klasifikasi ulasan secara real time.



Corresponding Author:

Mochamad Miftah Rachmatullah,

Program Studi Teknik Informatika,

STMIK IKMI Cirebon,

Email: mochamadmiftah34@gmail.com

1. PENDAHULUAN

Perkembangan teknologi digital yang pesat telah mendorong peningkatan penggunaan aplikasi *mobile* untuk berbagai kebutuhan sehari-hari, termasuk layanan hiburan seperti pemesanan tiket bioskop. TIX ID, sebagai salah satu *platform* pemesanan tiket terpopuler di Indonesia, menghasilkan ribuan ulasan pengguna di Google Play Store yang berisi informasi krusial mengenai tingkat kepuasan, keluhan, dan pengalaman pengguna (Awaludin, Nuryadi, et al., 2024). Untuk mengolah data teks dalam jumlah besar ini secara efisien, analisis sentimen berbasis *Machine Learning* dan *Natural Language Processing* (NLP) menjadi pendekatan yang sangat dibutuhkan. Sejalan dengan peningkatan volume data tersebut, pengembangan model otomatisasi yang akurat menjadi semakin mendesak untuk membantu pengembang memahami persepsi publik secara *real-time* (Hanafi & R, 2024; Siregar et al., 2024).

Meskipun analisis sentimen telah banyak diterapkan pada berbagai domain seperti transportasi dan layanan publik (Syafrizal et al., 2024; Triana, 2023), mayoritas penelitian masih terpaku pada

klasifikasi biner (positif dan negatif). Padahal, dalam konteks evaluasi aplikasi seperti TIX ID, keberadaan sentimen netral sangat signifikan untuk menangkap opini yang objektif atau datar (Awaludin & Wahono, 2015), sehingga pendekatan klasifikasi multi-kelas (*multi-class classification*) menjadi lebih relevan (Kundana & Chairani, 2023). Tantangan ini diperberat oleh karakteristik data ulasan yang seringkali tidak seimbang (*imbalanced dataset*) dan penggunaan bahasa Indonesia yang tidak baku atau informal, yang menuntut strategi pra-pemrosesan dan pemilihan algoritma yang tepat agar akurasi model tetap terjaga (Insan et al., 2023; Mandasari et al., 2022).

Dari sisi metodologi, *Multinomial Naïve Bayes* (MNB) dan *Logistic Regression* (LR) merupakan dua algoritma yang paling umum digunakan dalam klasifikasi teks (Awaludin, Yasin, et al., 2024). MNB dikenal dengan efisiensi komputasinya pada data berdimensi tinggi, sementara LR unggul dalam interpretabilitas probabilitas (Bahtiar et al., 2023; Rahman et al., 2025). Namun, penelitian yang secara komprehensif membandingkan performa kedua algoritma ini secara *head-to-head* pada dataset ulasan aplikasi hiburan lokal dengan melibatkan variasi teknik ekstraksi fitur (*CountVectorizer* vs TF-IDF) masih terbatas (Averina et al., 2022; Arief & Samsudin, 2023). Studi komparasi ini penting untuk mengisi kesenjangan literatur mengenai kombinasi model dan representasi fitur mana yang paling optimal untuk menangani karakteristik linguistik ulasan aplikasi di Indonesia.

Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk membangun dan membandingkan performa model analisis sentimen multi-kelas pada ulasan aplikasi TIX ID menggunakan algoritma *Multinomial Naïve Bayes* dan *Logistic Regression*. Penelitian ini secara spesifik akan menguji efektivitas kedua algoritma tersebut ketika dipadukan dengan dua teknik ekstraksi fitur yang berbeda, yaitu *CountVectorizer* dan TF-IDF. Selain itu, penelitian ini juga bertujuan untuk mengevaluasi dampak penerapan *Hyperparameter tuning* dalam meningkatkan stabilitas dan akurasi model, khususnya dalam menangani kelas sentimen yang minoritas.

Penelitian ini diharapkan memberikan kontribusi signifikan dengan menyediakan analisis komparatif yang mendalam antara MNB dan LR pada domain yang spesifik. Temuan empiris mengenai pengaruh teknik ekstraksi fitur terhadap performa model akan melengkapi literatur yang ada, sekaligus memberikan wawasan praktis bagi pengembang aplikasi dalam menangani data ulasan yang tidak seimbang. Hasil akhir dari penelitian ini berupa model klasifikasi terbaik yang dapat dijadikan rujukan atau *baseline* untuk pengembangan sistem analisis sentimen otomatis pada platform layanan digital lainnya di Indonesia.

2. METODE

Hasil Ringkasan Penelitian ditampilkan pada Tabel 1.

Tabel 1 Ringkasan Penelitian

No	Peneliti (Tahun)	Pembahasan (Masalah & Solusi)	Hasil Utama
1	Arief & Samsudin (2023)	Masalah: kesulitan mengklasifikasi kelas netral. Solusi: <i>hybrid</i> VADER + <i>Multinomial Logistic Regression</i> .	Akurasi meningkat signifikan terutama pada kelas netral.
2	Averina et al. (2022)	Masalah: variasi kosakata ulasan film. Solusi: <i>Logistic Regression</i> + TF-IDF.	LR memberikan generalisasi kuat untuk <i>multi-class</i> .
3	Bahtiar et al. (2023)	Masalah: perbandingan MNB vs LR dan <i>labeling</i> berbasis <i>rating</i> . Solusi: evaluasi komprehensif kedua model.	LR lebih presisi, MNB lebih cepat dan stabil.
4	Cahyani et al. (2022)	Masalah: kualitas teks buruk. Solusi: <i>preprocessing</i> bertahap.	<i>Preprocessing</i> meningkatkan akurasi model secara signifikan.
5	Hanafi & R.K. (2024)	Masalah: <i>noise</i> , <i>slang</i> , karakter tidak baku. Solusi: <i>Naïve Bayes</i> + <i>preprocessing</i> lengkap.	NB bekerja baik untuk ulasan Sirekap.
6	Herjanto & Carudin (2024)	Masalah: ketidakseimbangan dataset ulasan Sirekap. Solusi: <i>Random Forest</i> + <i>balancing</i> .	RF memberikan performa stabil dan akurasi tinggi.

No	Peneliti (Tahun)	Pembahasan (Masalah & Solusi)	Hasil Utama
7	Insan et al. (2023)	Masalah: teks tidak konsisten pada ulasan BRImo. Solusi: NB + <i>stemming</i> + <i>cleansing</i> .	NB menghasilkan akurasi hingga 84%.
8	Kundana & Chairani (2023)	Masalah: interpretasi ulasan pariwisata. Solusi: NB + <i>preprocessing</i> .	NB efektif pada domain wisata.
9	Mandasari et al. (2022)	Masalah: banyak slang/emotikon pada ulasan Grab. Solusi: normalisasi & MNB.	MNB bekerja sangat baik untuk teks informal.
10	Mary & Rose (2023)	Masalah: <i>sparsity</i> pada fitur pendidikan. Solusi: TF-IDF + <i>n-gram</i> + <i>classifier multi-class</i> .	LR & SVM meningkat akurasi.
11	Nadira et al. (2023)	Masalah: pelabelan manual sulit. Solusi: <i>lexicon-based InSet</i> + NB.	Efektif untuk pelabelan otomatis.
12	Novantika (2022)	Masalah: performa SVM vs LR pada Google Meet. Solusi: membandingkan dua algoritma.	SVM unggul, LR kompetitif.
13	Putri & Kharisudin (2022)	Masalah: bandingkan SVM, NB, LR pada Tokopedia. Solusi: eksperimen multi-algoritma.	SVM tertinggi, NB efisien.
14	Rahman et al. (2023)	Masalah: data tidak seimbang pada media sosial. Solusi: <i>multi-tier classification</i> .	Recall meningkat pada kelas minoritas.
15	Rahman et al. (2025)	Masalah: slang dan konteks <i>gaming</i> . Solusi: NB + <i>preprocessing</i> .	NB bekerja baik untuk ulasan <i>eFootball</i> .
16	Siregar et al. (2024)	Masalah: performa rendah pada teks Vidio. Solusi: NB + pengolahan teks intensif.	Akurasi meningkat secara konsisten.
17	Syafrizal et al. (2024)	Masalah: perbandingan NB vs KNN pada PLN <i>Mobile</i> . Solusi: dua model & <i>preprocessing</i> .	NB lebih unggul dari KNN.
18	Triana (2023)	Masalah: <i>noise</i> pada ulasan KAI Access. Solusi: NB + <i>stemming</i> .	Akurasi meningkat setelah normalisasi.

Berdasarkan Tabel 1, penelitian terkait analisis sentimen dalam tiga tahun terakhir menunjukkan konsistensi penggunaan algoritma seperti *Multinomial Naïve Bayes*, *Logistic Regression*, dan *Random Forest* untuk memproses ulasan aplikasi digital berbahasa Indonesia. Sebagian besar studi menekankan pentingnya preprocessing komprehensif untuk mengatasi *noise*, *slang*, dan variasi penulisan pada teks ulasan (Cahyani et al., 2022; Mandasari et al., 2022; Insan et al., 2023). Di sisi lain, penelitian yang menggunakan *Logistic Regression* menemukan bahwa model ini bekerja optimal ketika dipadukan dengan TF-IDF dan *n-gram* pada dataset bersifat kompleks (Averina et al., 2022; Mary & Rose, 2023).

Sementara itu, beberapa penelitian melakukan komparasi antar-algoritma seperti SVM, *Naïve Bayes*, *Logistic Regression*, dan KNN, dan hasilnya menunjukkan bahwa masing-masing model memiliki keunggulan tersendiri bergantung pada struktur data dan fitur yang digunakan (Bahtiar et al., 2023; Putri & Kharisudin, 2022; Syafrizal et al., 2024). Namun, belum ada studi yang secara khusus membandingkan MNB dan LR dengan dua teknik ekstraksi fitur (*CountVectorizer* dan TF-IDF) pada domain aplikasi pemesanan tiket film seperti TIX ID. Selain itu, kajian mengenai *hyperparameter tuning* pada kedua model ini dalam konteks multi-class juga masih terbatas.

3. HASIL DAN PEMBAHASAN

3.1. Hasil Pra-Pemrosesan Data

Tahap pra-pemrosesan dilakukan untuk menstandarkan teks ulasan sehingga cocok digunakan sebagai input algoritma machine learning. Perubahan yang terjadi pada data dapat dilihat pada Tabel 2, yang menampilkan contoh teks asli dibandingkan dengan hasil *preprocessing* yang mencakup *cleaning*, *case folding*, *tokenizing*, *stopword removal*, dan *stemming*.

Tabel 2 Sebelum dan Sesudah Text *Preprocessing*

No	Kalimat Asli (Sebelum)	Setelah Cleaning	Setelah Case folding	Setelah Tokenizing	Setelah Stopword removal	Setelah Stemming
1	Aplikasi TIX ID-nya bagus banget! Tapi kadang loading lama 😊	Aplikasi TIX IDnya bagus banget Tapi kadang loading lama	aplikasi tix idnya bagus banget tapi kadang loading lama	[aplikasi, tix, idnya, bagus, banget, tapi, kadang, loading, lama]	[tix, idnya, bagus, loading, lama]	[tix, id, bagus, load, lama]
2	Sangat membantu untuk beli tiket bioskop tanpa antri di tempat!	Sangat membantu untuk beli tiket bioskop tanpa antri di tempat	sangat membantu untuk beli tiket bioskop tanpa antri di tempat	[sangat, membantu, untuk, beli, tiket, bioskop, tanpa, antri, di, tempat]	[membantu, beli, tiket, bioskop, antri, tempat]	[bantu, beli, tiket, bioskop, antri, tempat]
3	Gagal login terus padahal koneksi internet lancar :(	Gagal login terus padahal koneksi internet lancar	gagal login terus padahal koneksi internet lancar	[gagal, login, terus, padahal, koneksi, koneksi, internet, internet, lancar]	[gagal, login, koneksi, internet, lancar]	[gagal, login, koneksi, internet, lancar]
4	Kenapa setiap mau checkout malah error terus ya?	Kenapa setiap mau checkout malah error terus ya	kenapa setiap mau checkout malah error terus ya	[kenapa, setiap, mau, checkout, checkout, malah, error, terus, ya]	[checkout, error]	[checkout, error]
5	Tampilan aplikasinya menarik, tapi sering keluar sendiri.	Tampilan aplikasinya menarik tapi sering keluar sendiri	tampilan aplikasinya menarik tapi sering keluar sendiri	[tampilan, aplikasinya, menarik, tapi, sering, keluar, keluar, sendiri]	[tampilan, aplikasinya, menarik, keluar, sendiri]	[tampil, aplikasi, tarik, keluar, diri]

Berdasarkan Tabel 2, terlihat bahwa proses *preprocessing* mampu mengurangi sejumlah besar kata tidak penting seperti kata sambung (dan, atau, yang), emotikon, serta variasi kata tidak baku. Perubahan ini membuat teks menjadi lebih ringkas dan fokus pada kata-kata bermakna. Temuan tersebut selaras dengan penelitian Mandasari et al. (2022), yang menegaskan bahwa *preprocessing* yang kuat sangat penting untuk teks berbahasa Indonesia karena banyak mengandung bahasa informal dan slang.

3.2. Hasil Ekstraksi Fitur dan Distribusi Data

Setelah *preprocessing*, data diubah menjadi representasi numerik menggunakan *CountVectorizer* dan TF-IDF. *CountVectorizer* menghasilkan 3.482 fitur, sementara TF-IDF menghasilkan jumlah fitur yang sama namun dengan bobot yang ter-normalisasi. Teknik ini penting karena dataset ulasan aplikasi seperti TIX ID memiliki karakteristik *sparse*, yaitu terdiri dari banyak kata unik yang jarang muncul.

Dataset kemudian dibagi menjadi 80% data latih dan 20% data uji. Distribusi kelas menunjukkan bahwa kelas positif mendominasi, sementara kelas netral lebih sedikit, sehingga

menyebabkan ketidakseimbangan data. Kondisi ini sering ditemukan pada ulasan aplikasi lokal, sebagaimana dijelaskan oleh Syafrizal et al. (2024) dan Kundana & Chairani (2023).

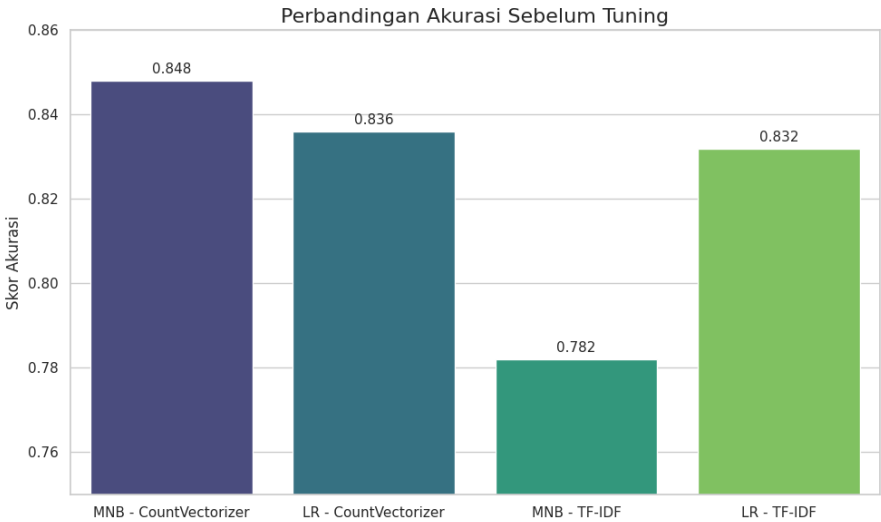
3.3. Evaluasi Model Awal (Sebelum *Tuning*)

Kinerja awal model diuji menggunakan kombinasi algoritma *Multinomial Naïve Bayes* (MNB) dan *Logistic Regression* (LR) dengan kedua teknik ekstraksi fitur. Hasilnya ditampilkan pada Tabel 3.

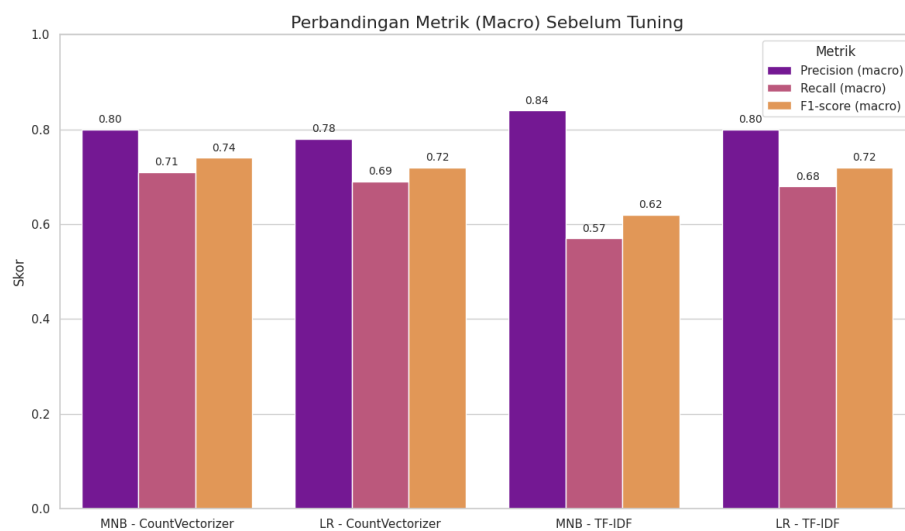
Tabel 3 Perbandingan Kinerja Model Sebelum Tuning

Model	Fitur Ekstraksi	Akurasi (%)	Precision (macro)	Recall (macro)	F1-score (macro)	Precision (weighted)	Recall (weighted)	F1-score (weighted)
MNB (MNB)	CountVectorizer	84.80%	0.80	0.71	0.74	0.84	0.85	0.84
LR (LR)	CountVectorizer	83.60%	0.78	0.69	0.72	0.82	0.84	0.82
MNB (MNB)	TF-IDF	78.20%	0.84	0.57	0.62	0.80	0.78	0.75
LR (LR)	TF-IDF	83.20%	0.80	0.68	0.72	0.82	0.83	0.82

Berdasarkan Tabel 3, model MNB dengan *CountVectorizer* merupakan kombinasi terbaik sebelum *tuning* dengan akurasi 84.80% dan F1-score *Macro* 0.74. Penguatan performa ini dipahami karena *CountVectorizer* menghitung frekuensi kata secara langsung sehingga sesuai dengan asumsi MNB yang bekerja optimal pada data berbasis hitungan (Rahman et al., 2025).



Gambar 1 Perbandingan Akurasi Model Sebelum Tuning



Gambar 2 Perbandingan Kinerja *Macro* Sebelum Tuning

Hal ini juga terlihat pada Gambar 1, yang menunjukkan bahwa model MNB *CountVectorizer* unggul dibandingkan kombinasi lainnya. Selanjutnya, Gambar 2 memperlihatkan bahwa kinerja *macro-average*, terutama *recall*, menurun drastis pada kelas netral mengindikasikan tantangan pada dataset *imbalanced*.

Model MNB dengan TF-IDF menunjukkan performa terendah (78.20%), sejalan dengan temuan Bahtiar et al. (2023) yang menyatakan bahwa TF-IDF sering menurunkan bobot kata yang sebenarnya penting bagi MNB, sehingga distribusi probabilitas menjadi tidak stabil.

Logistic Regression menunjukkan performa lebih stabil pada TF-IDF, mendukung hasil penelitian Averina et al. (2022) yang menyebutkan bahwa LR lebih responsif terhadap bobot fitur yang terdistribusi secara logaritmik.

3.4. Evaluasi Model Setelah *Tuning*

Hasil *tuning* menggunakan *GridSearchCV* ditampilkan pada Tabel 4.

Tabel 4 Perbandingan Kinerja Model Setelah Tuning

Model	Fitur Ekstraksi	Akurasi (%)	Precision (macro)	Recall (macro)	F1-score (macro)	Precision (weighted)	Recall (weighted)	F1-score (weighted)
MNB (MNB)	<i>CountVectorizer</i>	83.80%	0.77	0.72	0.74	0.74	0.83	0.84
LR (LR)	<i>CountVectorizer</i>	83.60%	0.78	0.69	0.69	0.72	0.82	0.84
MNB (MNB)	TF-IDF	83.80%	0.78	0.57	0.70	0.73	0.83	0.84
LR (LR)	TF-IDF	83.00%	0.76	0.68	0.70	0.72	0.82	0.83

Berdasarkan Tabel 3, model MNB *CountVectorizer* dan MNB TF-IDF menunjukkan akurasi tertinggi sebesar 83.80%. Meskipun akurasinya sedikit turun dari model awal, terjadi peningkatan pada stabilitas nilai precision dan *recall*, terutama untuk kelas positif.

Hyperparameter terbaik yang ditemukan adalah:

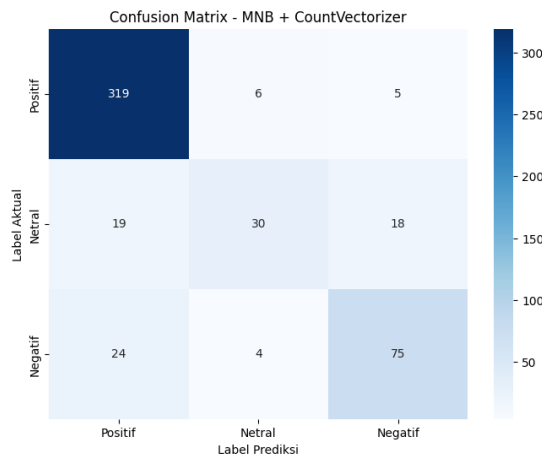
1. MNB: $\alpha = 1.0$
2. LR: $C = 1.0$, *penalty* L2, *solver* liblinear

Temuan ini sejalan dengan penelitian Rahman et al. (2025), yang menunjukkan bahwa *tuning* dapat meningkatkan generalisasi meskipun tidak selalu meningkatkan akurasi secara signifikan.

Namun, seluruh model tetap menunjukkan performa rendah pada kelas netral. Hal ini menguatkan temuan Insan et al. (2023) dan Mandasari et al. (2022) bahwa kelas netral pada ulasan aplikasi Indonesia seringkali memiliki ekspresi linguistik yang tumpang tindih dengan sentimen positif, sehingga sulit dibedakan oleh model klasik seperti MNB dan LR.

3.5. Analisis Mendalam dan Confusion Matrix

Model terbaik (MNB + *CountVectorizer*) dianalisis lebih mendalam menggunakan *confusion matrix* yang ditampilkan pada **Gambar 3**.



Gambar 3 *Confusion matrix* MNB Teknik *CountVectorizer*

Berdasarkan Gambar 3, terlihat bahwa model sangat baik dalam memprediksi kelas positif, yang merupakan kelas dominan. Namun, terjadi kesalahan yang cukup besar pada kelas netral, karena banyak diprediksi sebagai kelas positif. Beberapa faktor penyebabnya adalah:

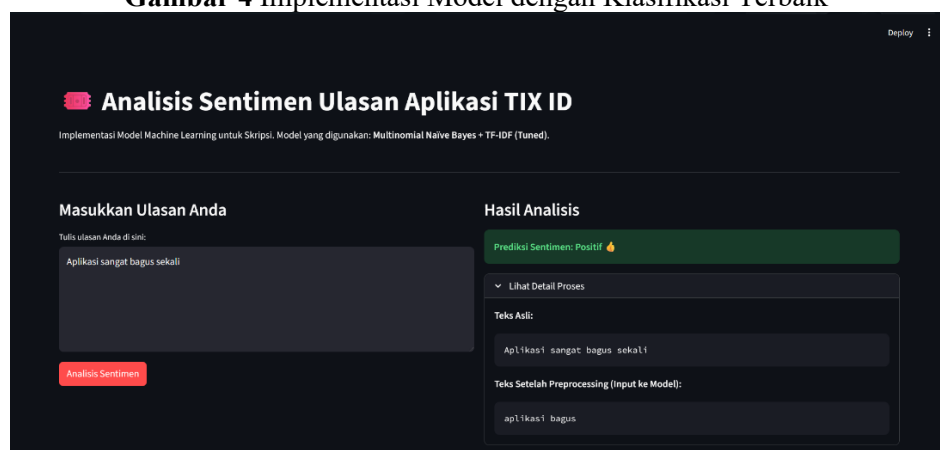
1. Ketidakseimbangan data, membuat model belajar lebih banyak pola pada kelas positif.
2. Kemiripan linguistik kelas netral dengan kelas positif, misalnya kata “lumayan”, “cukup bagus”, “oke”, yang cenderung condong positif.
3. *Sparse features*, yang membuat pemisahan antar kelas lebih sulit pada model linear dan probabilistik.

Analisis ini memperkuat temuan penelitian Syafrizal et al. (2024), bahwa kesalahan klasifikasi terbesar pada analisis sentimen multi-kelas terjadi pada kelas netral.

3.6. Hasil Implementasi Sistem

Pengujian implementasi dilakukan menggunakan aplikasi berbasis web yang dibangun dengan Streamlit. Antarmuka sistem ditampilkan pada Gambar 4.

Gambar 4 Implementasi Model dengan Klasifikasi Terbaik



Berdasarkan Gambar 4, sistem mampu memberikan hasil prediksi sentimen secara *real-time*, lengkap dengan tahap *preprocessing* otomatis. Hal ini menunjukkan bahwa model tidak hanya kuat secara teoretis, tetapi juga layak diterapkan dalam skenario dunia nyata, mendukung tujuan penelitian dan memberi nilai praktis bagi pengembang aplikasi TIX ID.

4. KESIMPULAN

Penelitian ini membandingkan performa algoritma *Multinomial Naïve Bayes* (MNB) dan *Logistic Regression* (LR) dalam klasifikasi sentimen multi-kelas pada ulasan pengguna aplikasi TIX ID dengan menggunakan dua teknik ekstraksi fitur, yaitu *CountVectorizer* dan TF-IDF. Hasil eksperimen menunjukkan bahwa model MNB dengan *CountVectorizer* sebelum dilakukan *Hyperparameter tuning* merupakan kombinasi terbaik dengan akurasi 84,80% serta F1-score *Macro* tertinggi. Kinerja ini dipengaruhi oleh kesesuaian karakteristik algoritma MNB terhadap data teks yang bersifat sparse dan berbasis frekuensi kata. Sementara itu, proses *tuning* memberikan peningkatan stabilitas performa bagi sebagian model, namun tidak mampu melampaui performa MNB *CountVectorizer* pada baseline. Evaluasi lebih lanjut melalui *confusion matrix* menunjukkan bahwa tantangan utama terletak pada prediksi kelas netral akibat ketidakseimbangan distribusi data dan kemiripan leksikal dengan kelas lain.

Penelitian ini memberikan kontribusi penting berupa analisis komparatif yang komprehensif terhadap dua algoritma populer dalam NLP untuk bahasa Indonesia serta menunjukkan peran signifikan pilihan teknik ekstraksi fitur dalam meningkatkan akurasi model. Selain itu, model terbaik berhasil diimplementasikan ke dalam aplikasi analisis sentimen berbasis web, yang mampu memproses ulasan pengguna secara *real-time*. Temuan penelitian ini dapat menjadi acuan bagi pengembangan sistem analisis sentimen pada aplikasi layanan digital lain di Indonesia serta membuka peluang penelitian lanjutan terkait penyeimbangan data dan pemanfaatan fitur berbasis *embedding*.

DAFTAR PUSTAKA

- Arief, M. R., & Samsudin, R. (2023). Hybrid approach with VADER and Multinomial Logistic Regression for multiclass sentiment analysis in online customer review. *International Journal of Advanced Computer Science and Applications*, 14(2). <https://doi.org/10.14569/IJACSA.2023.0140210>
- Awaludin, M., Nuryadi, H., & Pribadi, G. N. (2024). *Sistem Otomatisasi Laporan untuk Optimalisasi Pelaporan Data Penelitian dan Pengabdian kepada Masyarakat di Universitas Dirgantara Marsekal Suryadarma*. 9675, 1–7.
- Awaludin, M., & Wahono, R. S. (2015). Penerapan Metode Distance Transform Pada Linear Discriminant Analysis Untuk Kemunculan Kulit Pada Deteksi Kulit. *Journal of Intelligent Systems*, 1(1), 48–54.
- Awaludin, M., Yasin, V., & Risyda, F. (2024). The Influence of Artificial Intelligence Technology, Infrastructure and Human Resource Competence on Internet Access Networks. *Inform : Jurnal Ilmiah Bidang Teknologi Informasi Dan Komunikasi*, 9(2), 111–120. <https://doi.org/10.25139/inform.v9i2.8109>
- Averina, N., Nurmalsari, Y., & Nugroho, A. (2022). Analisis sentimen multi-kelas untuk film berbasis teks ulasan menggunakan model regresi logistik. *Teknika*, 8(2), 101–110. <https://doi.org/10.34148/teknika.v8i2.625>
- Bahtiar, S., Kurniawan, D., & Rahmawati, H. (2023). Comparison of Naïve Bayes and Logistic Regression in sentiment analysis on marketplace reviews using rating-based labeling. *Journal of Information Systems and Informatics*, 5(1), 35–45. <https://doi.org/10.33557/journalisi.v5i1.198>
- Cahyani, D., Fajri, I., & Saputra, R. (2022). Penerapan preprocessing teks pada analisis sentimen berbahasa Indonesia. *Jurnal Teknologi Informasi dan Komputer*, 8(3), 245–253. <https://doi.org/10.33330/jutik.v8i3.1568>
- Hanafi, M. R., & R. K., R. (2024). Analisis sentimen pada ulasan aplikasi Sirekap di Google Play menggunakan algoritma Naïve Bayes. *MALCOM*, 4(4), 112–121. <https://doi.org/10.57152/malcom.v4i4.1693>
- Herjanto, Y. A., & Carudin, C. (2024). Analisis sentimen ulasan pengguna aplikasi Sirekap pada Play Store menggunakan algoritma Random Forest Classifier. *Jurnal Informatika dan Teknik Elektro Terapan*, 12(2), 210–218. <https://doi.org/10.23960/jitet.v12i2.4192>
- Insan, M. T., Rakhman, A., & Wibowo, F. S. (2023). Analisis sentimen aplikasi Brimo pada ulasan pengguna di Google Play menggunakan algoritma Naïve Bayes. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 11(1), 50–58. <https://doi.org/10.36040/jati.v11i1.6377>
- Kundana, D., & Chairani, C. (2023). Data driven analysis of Borobudur ticket sentiment using Naïve Bayes. *Aptisi Transactions on Technopreneurship*, 5(2), 170–178. <https://doi.org/10.34306/att.v5i2sp.353>
- Mandasari, A. F., Pratama, D., & Yazid, M. (2022). Analisis sentimen pengguna transportasi online terhadap layanan Grab Indonesia menggunakan Multinomial Naïve Bayes Classifier. *J-SISKO TECH*, 5(1), 55–63. <https://doi.org/10.53513/jsk.v5i1.3832>
- Mary, S., & Rose, J. (2023). Multifaceted sentiment detection system to avoid dropout in virtual learning environment using multi-class classifiers. *International Journal of Advanced Computer Science and Applications*, 14(3), 41–49. <https://doi.org/10.14569/IJACSA.2023.0140305>

- Nadira, S., Ramadhan, J., & Fauzi, A. (2023). Analisis sentimen pada ulasan aplikasi mobile banking menggunakan metode Naïve Bayes dengan kamus Inset. *Informatic and Computational Intelligent Journal*, 5(2), 80–90. <https://doi.org/10.32764/icij.v5i2.1291>
- Novantika, L. (2022). Analisis sentimen ulasan pengguna aplikasi video conference Google Meet menggunakan metode SVM dan Logistic Regression. *Prosiding Seminar Nasional Matematika*, 9(1), 77–84. <https://doi.org/10.31227/osf.io/h6a2g>
- Putri, L. D., & Kharisudin, A. (2022). Analisis sentimen pengguna aplikasi marketplace Tokopedia pada situs Google Play menggunakan metode SVM, Naïve Bayes, dan Logistic Regression. *Prosiding Seminar Nasional Matematika*, 8(1), 120–127. <https://doi.org/10.31227/osf.io/58mdu>
- Rahman, M. A., Nurlaila, R., & Lubis, N. (2023). Multi-tier sentiment analysis of social media text using supervised machine learning. *Computers, Materials & Continua*, 74(1), 1301–1317. <https://doi.org/10.32604/cmc.2023.034821>
- Rahman, R. N., Rahim, A., & Pranoto, W. J. (2025). Analisis sentimen ulasan game eFootball 2024 pada Playstore menggunakan algoritma Naïve Bayes. *Jurnal Ilmiah Informatika*, 13(1), 45–54. <https://doi.org/10.33884/jif.v13i01.9913>
- Siregar, M. Y., Wiranata, A. D., & Saputra, R. A. (2024). Analisis sentimen pada ulasan pengguna aplikasi streaming Vidio menggunakan metode Naïve Bayes. *KLIK*, 4(5), 302–310. <https://doi.org/10.30865/klik.v4i5.1787>
- Syafrizal, S., Afdal, M., & Novita, R. (2024). Analisis sentimen ulasan aplikasi PLN Mobile menggunakan Naïve Bayes Classifier dan K-Nearest Neighbor. *MALCOM*, 4(1), 50–60. <https://doi.org/10.57152/malcom.v4i1.983>
- Triana, T. (2023). Implementasi algoritma Naïve Bayes untuk klasifikasi sentimen ulasan pengguna KAI Access. *JUTISI: Jurnal Ilmiah Teknik Informatika dan Sistem Informasi*, 14(1), 18–25. <https://doi.org/10.35889/jutisi.v14i1.2437>