

ANALISIS SENTIMEN KUALITAS LAYANAN J&T MENGGUNAKAN METODE NAÏVE BAYES CLASSIFIER

Andhika Dias Fernanda¹, Karisna Wahyuningtyas², Bunga Amalia Putri³, Raihan Fadhilah⁴,
Marta Ardiyanto⁵

¹²³⁴⁵ Sistem Informasi, Ilmu Komputer, Universitas Duta Bangsa Surakarta

Jln. Bhayangkara No.55, Tipes, Kec. Serengan, Kota Surakarta, Jawa Tengah 57154
¹ 220101006@mhs.ud.ac.id, ² 220101021@mhs.udb.ac.id, ³ 220101009@mhs.udb.ac.id,
⁴ 220101070@mhs.udb.ac.id, ⁵ marta.ardiyanto@udb.ac.id

Abstract

This study investigates user sentiment regarding the services provided by J&T Express, a prominent logistics provider in Indonesia. The research data was gathered from user reviews on the J&T Express application available via the Google Play Store. To process and categorize the sentiments within this textual data, the Naïve Bayes algorithm was employed, recognized as an effective probabilistic method for sentiment analysis. The analysis revealed that the Naïve Bayes Classifier demonstrated strong performance in identifying and classifying user sentiments toward J&T's services, achieving an accuracy rate of 91.58%, precision of 84%, recall of 92%, and an F1-score of 88%. These findings suggest that the developed sentiment analysis model accurately categorizes user reviews of the J&T application, thereby proving its utility in assessing and enhancing service quality.

Keywords : *Sentiment Analysis, J&T Express, Google Play Store, Naïve Bayes Classifier.*

Abstrak

Studi ini berfokus pada analisis sentimen pengguna terhadap layanan ekspedisi J&T Express, salah satu penyedia jasa logistik terkemuka di Indonesia. Data penelitian diperoleh dari ulasan pengguna pada aplikasi J&T Express yang tersedia di Google Play Store. Untuk mengolah dan mengklasifikasikan sentimen pada data berbentuk teks, digunakan algoritma Naïve Bayes sebuah metode probabilistik yang efektif dalam analisis sentimen. Berdasarkan hasil analisis, algoritma Naïve Bayes Classifier menunjukkan performa yang cukup baik dalam mengidentifikasi dan mengklasifikasikan sentimen pengguna terhadap layanan J&T dengan tingkat akurasi 91,58%, precision 84%, recall 92% dan f1-score 88%. Hasil akhir dari penelitian ini menunjukkan bahwa model analisis sentimen yang dikembangkan dapat mengkategorikan ulasan pengguna aplikasi J&T dengan tepat, sehingga dapat digunakan secara efektif untuk menilai dan meningkatkan mutu layanan.

Kata kunci : *Analisis Sentimen, J&T Express, Google Play Store, Naïve Bayes Classifier*

1. PENDAHULUAN

Pesatnya perkembangan teknologi digital yang signifikan telah menginisiasi perubahan berbagai layanan tradisional menjadi layanan yang dijalankan melalui aplikasi digital. Kemajuan sektor logistik terlihat dari luasnya pemanfaatan platform digital, yang bertujuan untuk menyederhanakan alur kerja, mempercepat proses pengiriman, dan meningkatkan efisiensi operasional secara keseluruhan. Perusahaan jasa pengiriman barang

kini berlomba-lomba mengembangkan layanan berbasis aplikasi untuk meningkatkan efisiensi dan kenyamanan bagi pengguna. J&T Express sebagai salah satu penyedia layanan logistik di Indonesia telah memanfaatkan teknologi ini dengan menyediakan aplikasi yang memungkinkan pengguna memesan, melacak, dan memberikan ulasan secara langsung. Aplikasi J&T dimanfaatkan oleh berbagai pihak, termasuk pengirim, penerima, serta mitra logistik. Tanggapan atau ulasan yang diberikan oleh pengguna menjadi sumber data yang penting karena menggambarkan pandangan serta

tingkat kepuasan mereka terhadap kualitas layanan yang disediakan. Namun demikian banyaknya jumlah ulasan yang masuk dapat diolah untuk proses analisis sebagai pengambilan keputusan dan peningkatan kualitas layanan pada aplikasi. Sebagai respons terhadap tingginya ulasan pengguna, diperlukan pendekatan otomatis yang mampu menganalisis dan mengelompokkan opini dengan cepat dan tepat. Masukan yang diberikan oleh pengguna melalui ulasan dapat digunakan sebagai acuan dalam mengukur sejauh mana kepuasan dan pengalaman mereka terhadap suatu layanan. Namun dari banyaknya jumlah ulasan aplikasi J&T yang masuk dibutuhkan pendekatan otomatis untuk mengelompokkan ulasan berdasarkan analisis sentimen.

Salah satu metode yang dapat digunakan adalah analisis sentimen, yaitu teknik dalam NLP yang berfungsi untuk mengenali, mengidentifikasi, dan mengelompokkan opini dalam bentuk teks ke dalam dua kategori sentimen, yaitu positif dan negatif. Penelitian ini menganalisis dan mengklasifikasikan data ulasan pengguna aplikasi J&T Express dengan metode Naïve Bayes Classifier (NBC), yakni algoritma berbasis probabilitas yang sederhana namun efektif untuk klasifikasi opini dalam bentuk teks. Metode ini telah diterapkan dalam sejumlah studi sebelumnya dan terbukti menunjukkan kinerja yang unggul dalam analisis data berbasis teks[1][2]. Penelitian yang dilakukan oleh Zulcharnain, Abdurrahman, dan Daryanto (2025) menerapkan algoritma Multinomial Naive Bayes untuk analisis sentimen ulasan aplikasi Duolingo[3], dengan tingkat akurasi yang tinggi. Sementara itu, Aini et al. (2022) membandingkan metode NBC, SVM, dan KNN dalam mengklasifikasikan ulasan aplikasi transportasi online[4], dan Hasanah & Sari (2024) menggunakan algoritma Naive Bayes Classifier (NBC) untuk menganalisis sentimen dari ulasan pengguna aplikasi ojek online Maxim yang diperoleh melalui platform Google Play Store. [5]. Penelitian mengenai analisis sentimen pengguna media sosial terhadap topik yang berbeda. Nurraharjo (2023) menerapkan Naive Bayes Classifier untuk mengevaluasi sentimen tweet mengenai peningkatan kasus Omicron. Meskipun metode ini terbukti efektif, penelitian tersebut menghadapi tantangan berupa ketidakseimbangan data dan adanya istilah medis yang menyulitkan analisis. Sementara itu, Zakasih dan Handoko (2022) juga menggunakan metode serupa untuk menganalisis sentimen pengguna Twitter terhadap Non-Fungible Tokens (NFT) [6]. Hasil penelitian mereka menunjukkan bahwa sebagian besar sentimen adalah netral, dengan kata kunci yang

sering muncul seperti "crypto", "jual", dan "koleksi"[7]. Meskipun metode Naive Bayes telah banyak digunakan, penelitian yang secara spesifik menganalisis ulasan pengguna terhadap layanan logistik digital, khususnya aplikasi J&T Express, masih relatif terbatas.

Penelitian ini memiliki kebaruan dari penerapan Naïve Bayes Classifier (NBC) secara khusus untuk menganalisis ulasan pengguna aplikasi J&T Express, yaitu platform logistik digital yang jarang dikaji secara spesifik atau mendalam. Dengan melalui pemanfaatan data ulasan Google Play Store, penelitian ini menyajikan analisis langsung terhadap pelanggan sektor logistik digital, berbeda dengan penelitian sebelumnya yang berfokus di e-commerce dan transportasi daring[8].

Penelitian yang dilakukan oleh Erfina dan Al-shufi (2022) analisis terhadap ulasan pengguna pada aplikasi kurir di Google Play Store memberikan gambaran awal mengenai persepsi pelanggan, termasuk J&T, JNE, SiCepat, Idexpress, dan Ninja Xpress[9]. Dalam penelitian ini, digunakan algoritma Naïve Bayes untuk mengklasifikasikan sentimen. Temuan yang mencapai tingkat keakuratan yang sangat tinggi, yaitu 100% untuk J&T, 98% untuk JNE, sekitar 97% untuk SiCepat dan Ninja Xpress, serta 94% untuk Idexpress. Model menghasilkan akurasi yang tinggi, sebagian besar menunjukkan ulasan negatif, khususnya mengenai keterlambatan pengiriman. Temuan ini menegaskan bahwa efektivitas algoritma Naive Bayes dalam mengolah sentimen teks, dengan dukungan proses text mining dan pra-pemrosesan menggunakan Python.

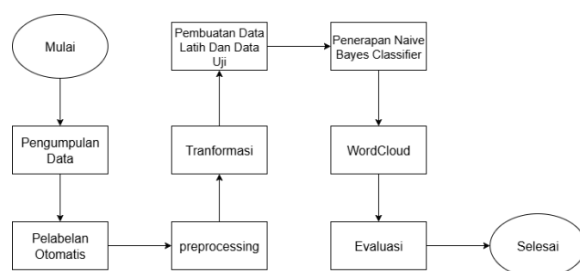
Penelitian selanjutnya yang dilakukan oleh Rahman, H. A., Santoso, R., & Widiarihi, T. (2023) melakukan perbandingan algoritma Multinomial Naive Bayes dengan Support Vector Machine (SVM) untuk analisis sentimen pengguna terhadap layanan J&T Express berdasarkan 2.500 data tweet[10]. Hasil penelitian yang telah dilakukan menunjukkan bahwa SVM memiliki tingkat akurasi yang lebih unggul, yakni 82,40%, sedangkan Naive Bayes hanya mencapai 72,80%. Kata "paket" menjadi kata yang paling sering muncul dalam analisis. Penelitian ini juga mengidentifikasi ketidakseimbangan distribusi sentimen sebagai faktor yang menurunkan akurasi, sehingga penggunaan metode SMOTE direkomendasikan untuk meningkatkan kualitas klasifikasi.

Irawansyah, R. S., & Wiriasto, G. W. (2023) menerapkan metode Naive Bayes Classifier untuk analisis sentimen terhadap program Merdeka

Belajar-Kampus Merdeka (MBKM) berdasarkan data dari Twitter [11]. Dari keseluruhan 1.175 tweet yang dianalisis, mayoritas mengandung sentimen positif sebesar 53,44%, sementara sentimen negatif hanya terdeteksi pada 12,08% tweet. Hasil klasifikasi mencapai akurasi 79,66%, dan ditampilkan melalui aplikasi visualisasi berbasis web. Salah satu tantangan utama dalam penelitian ini adalah keterbatasan jumlah data yang merepresentasikan sentimen negatif, yang berpotensi mempengaruhi kestabilan dan kemampuan model dalam mengklasifikasikan sentimen secara seimbang.

2. METODOLOGI PENELITIAN

Knowledge Discovery in Database (KDD) adalah data besar diolah secara bertahap mengungkap pengetahuan tersembunyi dan relevan dengan kebutuhan analisis. Dalam konteks analisis sentimen terhadap kualitas layanan sistem logistik J&T, dengan penerapan metode Naive Bayes Classifier, proses KDD diawali dengan pemilihan data relevan berupa ulasan dan komentar pelanggan terkait pengalaman mereka menggunakan layanan J&T. Pemodelan Knowledge Discovery in Database memiliki beberapa tahapan yaitu pemilihan data, pra pemrosesan, transformasi data, proses pengembangan data, serta tahap evaluasi. Pada tahapan pra pemrosesan dilakukan pembersihan teks, case folding (mengubah ke huruf kecil), tokenisasi, penghapusan stopwords, dan stemming. Sedangkan pada tahapan transformasi dilakukan pengubah teks menjadi representasi numerik menggunakan metode TF-IDF. Berikut gambaran lebih lanjut mengenai setiap tahap dalam metodologi secara rinci.



Gambar 1. Tahap Penelitian

2.1 Pengumpulan Data

Pada tahapan ini, data ulasan pengguna terhadap layanan aplikasi J&T dikumpulkan dari platform Google Play Store dan digunakan sebagai dataset untuk analisis sentimen. Teknik web scraping digunakan untuk mengumpulkan data, dengan

JupyterLab sebagai lingkungan pengembangan berbasis Python yang mendukung pemrosesan data secara online. Dengan menggunakan teknik ini, berhasil dikumpulkan 1.000 ulasan dengan kategori paling relevan.

2.2 Pelabelan Otomatis

Dataset ulasan pengguna yang telah diasumsikan bahwa kemudian ditinjau dan diproses menggunakan fase pelabelan otomatis dengan nilai cek dengan nilai bintang 4 dan 5 diasumsikan mencerminkan sentimen positif. Sementara ulasan dengan rating bintang 3 dikategorikan sebagai netral dan tidak disertakan dalam proses pelabelan sentimen. Sementara itu, ulasan yang memperoleh bintang 1 dan 2 mewakili sentimen negatif.

2.3 Preprocessing

Proses case folding dilakukan untuk mengonversi semua huruf dalam data teks menjadi bentuk huruf kecil. Tahapan berikutnya, pembersihan teks dari elemen yang tidak diperlukan, seperti angka, simbol, dan tanda baca. Selanjutnya dilakukan stopwords removal, yaitu menghilangkan kata umum yang tidak memiliki makna terhadap teks. Proses terakhir yaitu tokenisasi memisahkan atau proses memecah kata dalam kalimat menjadi unit yang lebih kecil.

2.4 Transformasi

TF - IDF adalah metode pembobotan kata dalam sebuah dokumen yang digunakan dalam teks untuk menggambarkan sebuah kata dalam dokumen. TF (Term Frequency) menghitung seberapa sering suatu kata muncul dalam dokumen atau ulasan, sedangkan IDF (Inverse Document Frequency) memberikan nilai yang lebih tinggi pada kata-kata yang sering muncul. TF-IDF banyak digunakan dalam berbagai aplikasi seperti pencarian informasi,

klasifikasi teks, dan sistem rekomendasi berbasis konten.

2.5 Pembuatan data latih dan data uji

Merupakan salah satu tahap krusial dalam proses data mining. Kedua jenis data ini memiliki penting dalam membangun serta mengevaluasi kinerja model. Pada tahap pelatihan (training), data berlabel digunakan untuk membangun model prediktif, sedangkan data uji dimanfaatkan pada tahap pengujian (testing) untuk mengevaluasi kinerja dan tingkat akurasi model terhadap data.

2.6 Penerapan naive bayes classifier

Algoritma Naive Bayes Classifier dirancang secara sederhana, namun mampu memberikan hasil yang optimal dan kinerja yang efektif. Meskipun strukturnya relatif sederhana, metode ini efektif dalam memproses data secara cepat, serta menghasilkan akurasi yang tinggi dalam berbagai kasus klasifikasi. Metode ini menggunakan prinsip teorema Bayesian dengan menggabungkan peluang awal dan bersyarat untuk menentukan kelas data yang paling mungkin. Pada proses klasifikasi, algoritma Naive Bayes menghitung seberapa besar kemungkinan suatu data termasuk ke dalam masing-masing kelas, berdasarkan nilai-nilai fitur yang dimilikinya. Kelas dengan probabilitas tertinggi yang akan dipilih sebagai hasil prediksi akhir. Meskipun pendekatannya bersifat sederhana dan mengasumsikan bahwa antar fitur saling bebas secara statistik (naïf), algoritma ini menunjukkan hasil yang efektif dalam berbagai permasalahan klasifikasi, seperti analisis sentimen, pengenalan teks, serta deteksi spam. Berikut dapat dilihat di persamaan (1).

Rumus Naïve Bayes Classifier:

$$P(c|X) = (P(X|c) \cdot P(c)) / P(X) \dots(1)$$

$P(c|X)$: Probabilitas kelas c setelah melihat data $\rightarrow$ hasil prediksi.

$P(X|c)$: Seberapa besar kemungkinan data X muncul di kelas c .

$P(c)$: Seberapa umum kelas c muncul (prior).

$P(X)$: Probabilitas umum data $X \rightarrow$ sering diabaikan saat klasifikasi.

2.7 WordCloud

Pada proses pembuatan word cloud melibatkan analisis frekuensi kata untuk mengidentifikasi istilah yang paling sering muncul dalam ulasan pengguna. Tujuan langkah ini untuk menggambarkan secara visual fokus utama dan pola yang ada dalam pandangan pengguna. Word

cloud berguna untuk menemukan kata-kata kunci yang paling terlihat, baik yang bersifat positif maupun negatif, tanpa harus membaca semua teks dengan cermat. Hasil dari word cloud akan menjadi langkah awal dalam memahami tema yang sering diangkat, seperti masalah pengiriman, pelayanan kurir, atau kualitas aplikasi, serta memberikan pemahaman awal sebelum melanjutkan ke analisis yang lebih rinci, seperti klasifikasi sentimen.

2.8 Evaluasi

Tahap evaluasi merupakan bagian penting mengevaluasi kinerja model klasifikasi untuk menilai kinerja serta validitas dari model yang telah dikembangkan. Evaluasi dilakukan melalui perhitungan metrik-metrik penting, seperti accuracy, precision, dan recall, guna menilai sejauh mana model mampu mengklasifikasikan data secara efektif. Perhitungan accuracy, precision, dan recall didasarkan pada confusion matrix, yang mencakup empat elemen penting, yakni True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Akurasi menunjukkan proporsi hasil klasifikasi yang benar sebanyak keseluruhan sampel pada data uji. Precision menentukan tingkat akurasi model dalam memprediksi kelas positif, yaitu rasio antara prediksi positif yang akurat (True Positive) dengan total seluruh prediksi positif. Sementara itu, recall mencerminkan kemampuan model dalam mendeteksi seluruh data positif yang sebenarnya, dengan menghitung rasio jumlah data positif yang berhasil diidentifikasi dengan benar terhadap total data positif aktual. Evaluasi dilakukan dengan membandingkan performa klasifikasi antara penerapan Frequency-Inverse Document Frequency (TF-IDF) sebagai seleksi fitur dan klasifikasi tanpa seleksi fitur. Perbandingan ini untuk mengukur efektivitas TF-IDF dalam meningkatkan representasi fitur teks, yang diharapkan dapat berkontribusi pada peningkatan akurasi dalam pengklasifikasian sentimen. Penelitian yang dilakukan Agustina et al. (2022) [12] dan Salsabila et al. (2022) [13] juga menunjukkan bahwa penggunaan TF-IDF sebagai metode transformasi teks mampu meningkatkan ketepatan klasifikasi dibandingkan metode frekuensi biasa [14][15].

3. HASIL DAN PEMBAHASAN

TABEL I. RINGKASAN TAHAPAN

No	Tahap	Uraian Singkat	Tujuan
1	Pengumpulan Data	Mengambil 1.000 ulasan dari Google Play Store menggunakan teknik web scraping.	Menyediakan data ulasan untuk dianalisis.
2	Pelabelan Sentimen	Memberi label sentimen berdasarkan rating bintang (1-2 negatif, 4-5 positif, 3 diabaikan).	Mengelompokkan data ke dalam kelas sentimen.

3	Pra-pemrosesan Teks	Melakukan normalisasi dan pembersihan teks (case folding, tokenisasi, stopword, stemming).	Mempersiapkan teks agar layak untuk dianalisis.
4	Transformasi Fitur	Mengubah data teks menjadi format numerik menggunakan metode TF-IDF.	Memberi bobot relevansi pada fitur kata.
5	Pembagian Dataset	Memisahkan data menjadi 80% latih dan 20% uji.	Melatih dan menguji model klasifikasi.
6	Klasifikasi	Menerapkan algoritma Naïve Bayes untuk mengklasifikasikan data berdasarkan probabilitas.	Memprediksi polaritas sentimen secara otomatis.
7	Evaluasi Model	Menggunakan metrik akurasi, presisi, recall, dan F1-score.	Menilai efektivitas performa model klasifikasi.
8	Visualisasi	Membuat Word Cloud berdasarkan frekuensi kata dalam ulasan.	Mengidentifikasi kata kunci dominan secara visual.

3.1 Pengumpulan Data

Tahap pertama dilakukan dengan mengumpulkan data melalui proses scraping dari tautan aplikasi di Google Play Store dengan mengambil data 1000 ulasan pelanggan J&T, berikut hasil proses pengumpulan data ulasan pelanggan.

	userName	score	at	content
0	Cahaya Jayatama	1	2025-05-26 16:07:43	udah order pick up jam 10.30, jam 15.25 dpt wa...
1	Basuki Endro	1	2025-05-26 14:24:07	kurir tidak mau mengantar kan paket bahkan tid...
2	Aisyah Utami	3	2025-05-25 01:11:13	jnt skrng telat ya pengiriman nya gak sprti du...
3	Meri Anggreyani	1	2025-05-24 21:29:05	sangat* buruk order pickup sampai 5 hari tapi ...
4	Soni Aditama	5	2025-05-24 13:36:46	pengirimannya sangat cepat 🍀🍀🍀
...	...	...	...	...
995	Pengguna Google	1	2018-11-08 11:45:24	Selama saya pengirim pakai yang lain belum p...
996	Pengguna Google	1	2018-10-26 14:00:28	Sekarang aplikasi nya payah masa sering bgt ma...
997	Pengguna Google	2	2018-10-25 00:04:52	aplikasi update yg sekarang malah bikin gak ny...
998	Pengguna Google	2	2018-09-26 16:11:36	Tolong BUGS pada maps dan petanya diperbaiki k...
999	Pengguna Google	5	2018-09-20 07:18:15	Tampilan barunya bagus, cek resi mudah, tapi s...

Gambar 2. Pengumpulan Data

3.2 Pelabelan Otomatis

Proses kedua adalah proses pelabelan otomatis dari data pengumpulan data untuk membedakan hasil maka 1 dan 2 akan dinyatakan sebagai negative, 3 akan diberi label netral, sedangkan 4 dan 5 sebagai positif. Hasil pelabelan dapat dilihat sebagai berikut.

No	userName	score	at	content	label
0	Cahaya Jayatama	1	2025-05-26 16:07:43	udah order pick up jam 10.30, jam 15.25 dpt wa...	negatif
1	Basuki Endro	1	2025-05-26 14:24:07	kurir tidak mau mengantar kan paket bahkan tid...	negatif
2	Meri Anggreyani	1	2025-05-25 21:29:05	Sangat* buruk order pickup sampai 5 hari tapi...	netral
3	Soni Aditama	1	2025-05-25 14:36:45	pengirimannya sangat cepat 🍀🍀🍀	positif
4	Arid Ridwan	5	2025-05-24 11:09:41	Udah dua kali paket cuman sampe gudang, terus...	negatif
...	...	...	...	...	...
999	Pengguna Google	5	2018-09-20 07:18:15	Tampilan barunya bagus, cek resi mudah, tapi s...	positif

Gambar 3. Pelabelan Otomatis

3.3 Pra Processing

Dalam tahap preprocessing akan dilakukan proses data filtering untuk menghapus ulasan 3 yang dikategorikan sebagai sentimen netral bertujuan untuk menyederhanakan proses klasifikasi menjadi dua kelas utama, yaitu sentimen positif dan negative. Pra-pemrosesan mencakup pembersihan teks dari unsur-unsur non-informatif, termasuk angka, tanda baca, simbol tertentu, serta spasi yang tidak diperlukan. Setelah proses case folding yang menyelaraskan huruf menjadi kecil, teks diproses lebih lanjut melalui tokenisasi, yaitu pemecahan kalimat menjadi token. Kata-kata umum tanpa makna signifikan, seperti “yang”, “dari”, dan “dan”, kemudian disaring melalui tahap stopword removal. Proses stemming dilakukan pada tahap akhir untuk memulihkan kata-kata ke bentuk asalnya, sehingga kata-kata turunan seperti “berjalan”, “berjalanlah”, atau “perjalanan” disederhanakan menjadi “jalan”.

3.3.1 Pembersihan Teks

Pada proses pembersihan teks menghapus elemen yang tidak di perlukan seperti angka, tanda baca, dan karakter khusus untuk mempersiapkan data agar layak digunakan dalam proses analisis, pemodelan, atau pemrosesan bahasa alami (Natural Language Processing/NLP). Contoh sebelum dan sesudah pembersihan teks ditunjukkan pada gambar berikut.

No	userName	score	at	content	clean_content	label
0	Cahaya Jayatama	1	2025-05-26 16:07:43	udah order pick up jam 10.30, jam 15.25 dpt wa...	udah order pick up jam jam dpt wa dari cs loka...	negatif
1	Basuki Endro	1	2025-05-26 14:24:07	kurir tidak mau mengantar kan paket bahkan tid...	kurir tidak mau mengantar kan paket bahkan tid...	negatif
2	Meri Anggreyani	1	2025-05-25 21:29:06	Sangat' buruk order pickup sampai 5 hari tapi...	sangat buruk order pickup sampai 5 hari tapi tid...	negatif
3	Soni Aditama	1	2025-05-25 14:36:45	pengirimannya sangat cepat 444	pengirimannya sangat cepat	positif
4	Arid Ridwan	5	2025-05-24 11:09:41	Udah dua kali paket cuman sampe udang terus...	udah dua kali paket cuman sampe udang terus...	negatif
...	...	...	...	...	...	...
999	Pengguna Google	5	2018-09-20 07:18:15	Tampilan barunya bagus cek resi mudah tapi say...	tampilan barunya bagus cek resi mudah tapi say...	positif

Gambar 4. Pembersihan Teks

3.3.2 Case Folding

Case folding adalah teknik text preprocessing untuk menyeragamkan teks agar perbedaan huruf kapital dan huruf kecil tidak mempengaruhi proses analisis. Berikut tampilan hasil dari proses case folding.

No	clean_content	casefold_content
0	udah order pick up jam jam dpt wa dari cs loka...	udah order pick up jam jam dpt wa dari cs loka...
1	kurir tidak mau mengantar kan paket bahkan tid...	kurir tidak mau mengantar kan paket bahkan tid...
2	sangat buruk order pickup sampai hari tapi tid...	sangat buruk order pickup sampai hari tapi tid...
3	pengirimannya sangat cepat	pengirimannya sangat cepat
4	udah dua kali paket cuman sampe gudang terus n...	udah dua kali paket cuman sampe gudang terus n...
...	...	...
999	tampilan barunya bagus cek resi mudah tapi say...	tampilan barunya bagus cek resi mudah tapi say...

Gambar 5. Case Folding

3.3.3 Tokenisasi

Pada tahap tokenisasi dilakukan untuk memecah teks ulasan menjadi bagian -bagian kecil, dalam bentuk kata -kata individual. Tujuan proses ini bahwa setiap kata dapat dianalisis secara individual melalui proses analisis sentimen. Misalnya "pengirimannya sangat cepat" dapat dikonversi menjadi token [pengirimannya, sangat, cepat] yang memungkinkan sistem lebih mudah memproses makna setiap kata. Berikut hasil sebelum dan sesudah tokenisasi.

index	casefold_content	tokens
0	udah order pick up jam jam dpt wa dari cs loka...	[udah, order, pick, up, jam, jam, dpt, wa, dari, cs, loka, ...]
1	kurir tidak mau mengantar kan paket bahkan tid...	[kurir, tidak, mau, mengantar, kan, paket, bahkan, tid, ...]
2	sangat buruk order pickup sampai hari tapi tid...	[sangat, buruk, order, pickup, sampai, hari, tapi, tid, ...]
3	pengirimannya sangat cepat	[pengirimannya, sangat, cepat]
4	udah dua kali paket cuman sampe udang terus n...	[udah, dua, kali, paket, cuman, sampe, udang, terus, n, ...]
...	...	...
999	tampilan barunya bagus cek resi mudah tapi say...	[tampilan, barunya, bagus, cek, resi, mudah, tapi, say, ...]

Gambar 6. Tokenisasi

3.3.4 Stopword Removal

Stopword Removal atau penghapusan stopwords adalah proses menghilangkan kata umum dalam tulisan tetapi tidak memiliki arti yang

penting dalam analisis, seperti "yang", "dari", "dan", atau "adalah". Pada tahapan stopwords removal dilakukan agar analisis dapat fokus pada

kata yang jelas menunjukkan opini pengguna, sehingga hasil analisis menjadi lebih tepat. Berikut hasil tampilan stopwords removal.

index	tokens	filtered_stopwords
0	[udah, order, pick, up, jam, jam, dpt, wa, dar...]	[udah, order, pick, up, jam, jam, dpt, wa, dar...]
1	[kurir, tidak, mau, mengantar, kan, paket, bah...]	[kurir, mengantar, paket, kabar, paket, nya, ka...]
2	[sangat, buruk, order, pickup, sampai, hari, t...]	[buruk, order, pickup, kunjung, jemput, niat, ...]
3	[pengirimannya, sangat, cepat]	[pengirimannya, cepat]
4	[udah, dua, kali, paket, cuman, sampe, gudang, ...]	[udah, kali, paket, cuman, sampe, gudang, ngad...]
...	...	...
999	[tampilan, barunya, bagus, cek, resi, mudah, t...]	[tampilan, barunya, bagus, cek, resi, mudah, s...]

Gambar 7. Stopword Removal

3.3.5 Stemming

Pada proses Stemming langkah merubah kata yang memiliki imbuhan menjadi bentuk dasar atau akar. Proses ini dilakukan dengan menghilangkan afiks, baik itu awalan, akhiran, maupun sisipan dari suatu kata. Misalnya, kata "pengiriman" diubah menjadi "kirim". Proses stemming membantu sistem dalam mengelompokkan kata yang memiliki makna yang sama. Tahap ini penting dalam analisis sentimen sistem tidak membedakan antara bentuk kata yang berbeda namun memiliki arti yang serupa, sehingga hasil analisis menjadi lebih tepat dan efisien. Berikut proses hasil stemming.

no	filtered_stopwords	stemmed_tokens
0	[udah, order, pick, up, jam, jam, dpt, wa, cs, ...]	[udah, order, pick, up, jam, jam, dpt, wa, cs, ...]
1	[kurir, mengantar, paket, kabar, paket, nya, k...]	[kurir, antar, paket, kabar, paket, nya, kurir, ...]
2	[buruk, order, pickup, kunjung, jemput, niat, ...]	[buruk, order, pickup, kunjung, jemput, niat, ...]
3	[pengirimannya, cepat]	[kirim, cepat]
4	[udah, kali, paket, cuman, sampe, gudang, ngad...]	[udah, kali, paket, cuman, sampe, gudang, ngad...]
...	...	...
999	[tampilan, barunya, bagus, cek, resi, mudah, s...]	[tampil, baru, bagus, cek, resi, mudah, sayang...]

Gambar 8. Stemming

3.4 Pembobotan IF-IDF

Pembobotan TF-IDF (Term Frequency-Inverse Document Frequency) merupakan teknik pemberian bobot pada kata berfungsi untuk mengevaluasi tingkat kepentingan kata dalam sebuah ulasan dibandingkan dengan ulasan lainnya. Kata-kata yang muncul dengan frekuensi tinggi dalam satu ulasan, namun sering kali tidak terlihat dalam tinjauan lainnya, hal ini memiliki tingkat penting yang lebih tinggi. Metode ini membantu sistem lebih menekankan pada kata-kata yang paling berkaitan dalam mengidentifikasi

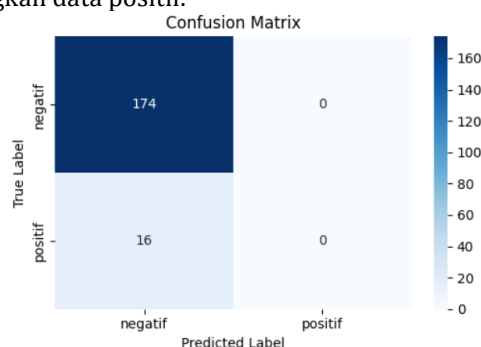
sentimen. Berikut adalah proses pembobotan TF-IDF.

no	stemmed_tokens	clean_review
0	[udah, order, pick, up, jam, jam, dpt, wa, cs,...]	udah order pick up jam jam dpt wa cs lokal nan...
1	[kurir, antar, paket, kabar, paket, nya, kurir,...]	kurir antar paket kabar paket nya kurir pulang
2	[buruk, order, pickup, kunjung, jemput, niat,...]	buruk order pickup kunjung jemput niat adakan ...
4	[kirim, cepat]	kirim cepat
5	[udah, kali, paket, cuman, sampe, gudang, ngadu,...]	udah kali paket cuman sampe gudang ngadu call ...
...	...	...
999	[tampil, baru, bagus, cek, resi, mudah, sayang,...]	tampilan barunya bagus cek resi mudah tapi sayang...

Gambar 9. Pembobotan TF-IDF

3.5 Pembuatan data latih dan data uji

Setelah tahap pra-pemrosesan, langkah selanjutnya adalah membagi data menjadi dua bagian, yakni data latihan dan data tes. Data latihan dipakai untuk merancang model agar bisa mengenali pola dan sifat sentimen dalam teks, sedangkan data pengujian untuk menilai kinerja model dalam mengklasifikasikan data baru dengan tepat. Melakukan data latih 80% dan data uji 20%, sehingga hasil evaluasi model menjadi lebih akurat. Data yang dipakai berjumlah 948, terbagi menjadi 758 data pelatihan dan 190 data pengujian. Secara keseluruhan, distribusi label memperlihatkan bahwa data berlabel negatif mendominasi dengan jumlah 868, sedangkan yang positif 80. Data pelatihan terdapat 694 data negatif dan 64 data positif. Sementara itu, pada data pengujian, terdapat 174 data negatif dan 16 data positif. Kondisi ini menunjukkan ketidakseimbangan dataset pada jumlah data negatif yang jauh lebih banyak dibandingkan data positif.



Gambar 10. Confusion Matrik

3.6 Penerapan Naive Bayes Classifier

Penerapan Naive Bayes Classifier (NBC) untuk membangun model klasifikasi sentimen menggunakan perhitungan probabilitas kemunculan kata dalam teks berdasarkan dataset pelatihan. Untuk meningkatkan akurasi perhitungan, bobot kata dihitung dengan menggunakan metode TF-IDF. Setelah melalui

proses pelatihan, model dievaluasi melalui pendekatan data uji untuk mengukur sejauh mana kemampuan dalam mengklasifikasikan ulasan ke dalam sentimen positif maupun negatif. Hasil dari pengujian menunjukkan bahwa model mencapai tingkat akurasi sebesar 91,58%, yang menunjukkan bahwa model tersebut secara umum cukup konsisten dan efektif dalam menganalisis sentimen dari ulasan pengguna. Namun meskipun tingkat akurasi cukup tinggi, model ini menunjukkan ketidakseimbangan kinerja dalam membedakan kedua jenis sentiment. Pada laporan klasifikasinya, label negatif menunjukkan precision 0,92, recall 1,00, dan f1-score 0,96 dari 174 data yang diuji. Sebaliknya, untuk label positif, model ini tidak dapat melakukan klasifikasi dengan baik, karena semua metrik precision, recall, dan f1-score bernilai 0,00 dari 16 data yang diuji. Ketidakseimbangan data terlihat pada confusion matrix, di mana seluruh data berlabel positif salah diklasifikasikan sebagai negatif, sedangkan semua data negatif berhasil diidentifikasi dengan tepat. Dengan total 1000 data, melalui tahap penghilangan ulasan netral menjadi 948 data, terdapat 868 yang berlabel negatif dan 80 yang berlabel positif. Oleh karena itu, disarankan untuk memperbaiki distribusi data atau menggunakan metode penyeimbang data agar performa model dapat meningkat di masa mendatang.

TABEL II. HASIL PENGUKURAN KINERJA MODEL

	Nilai
Akurasi	91,58%
Recall	92%
Precision	84%
F1-score	88%

3.7 Word Cloud

Word Cloud menunjukkan bahwa komentar dari pengguna didominasi istilah seperti "paket", "barang", "kurir", "alamat", dan "pengiriman", yang menunjukkan perhatian utama pada proses pengantaran. Kata-kata dengan konotasi negatif seperti "gak", "buruk", dan "kecewa" juga muncul secara signifikan, menandakan adanya banyak keluhan mengenai keterlambatan, kesalahan dalam alamat, serta layanan kurir. Hasil ini sama dengan penelitian Misrun, Haerani, Fikry, dan Budianita (2023) yang menunjukkan bahwa kata-kata seperti "bohong", "gagal", dan "buruk" mendominasi komentar negatif pada platform YouTube[16]. Kata yang paling banyak diungkapkan dalam ulasan J&T adalah "paket". Secara keseluruhan, kumpulan kata ini mencerminkan pendapat negatif pengguna terhadap layanan pengiriman.



4. KESIMPULAN DAN SARAN

4) Hasil penelitian ini sesuai dengan studi Rahman et al. (2023) yang membandingkan algoritma Naïve Bayes dan SVM dalam analisis sentimen terhadap layanan J&T Express. Mereka menemukan bahwa SVM mencapai akurasi sebesar 82,40%, sedangkan Naïve Bayes hanya 72,80%. Seperti penelitian ini, ketidakseimbangan data menjadi masalah utama, sehingga disarankan menggunakan metode SMOTE. Temuan ini mendukung hasil penelitian saat ini yang berhasil mencapai akurasi 91,58%, meski nilai precision, recall, dan f1-score untuk kelas positif berada di angka 0 karena dominasi data negatif.

- 1) Klasifikasi ulasan pengguna terhadap aplikasi J&T Express telah berhasil dilaksanakan menggunakan metode Naïve Bayes melalui beberapa langkah penting. terdiri dari pengumpulan data, pembersihan serta normalisasi teks, transformasi fitur menggunakan metode TF-IDF, dan juga pelatihan serta pengujian model guna menilai kinerjanya.

2) Berdasarkan hasil dapat disimpulkan bahwa model klasifikasi sentimen yang dikembangkan menunjukkan bahwa akurasi model mencapai 91,58%, yang mencerminkan performa klasifikasi yang cukup baik. Namun, akurasi tinggi belum mencerminkan kinerja sebenarnya karena data negatif jauh lebih banyak dari data positif.

- 1) Tingkat akurasi model cukup memuaskan, model tersebut belum mampu mengidentifikasi ulasan positif secara akurat. Oleh sebab itu, disarankan untuk menjajaki metode selain Naïve Bayes, pendekatan lain seperti SVM atau Random Forest juga berpotensi memberikan hasil yang kompetitif., untuk memperoleh hasil klasifikasi yang lebih seimbang.

2) Selain itu, jumlah ulasan negatif jauh lebih tinggi dibandingkan dengan ulasan positif, disarankan untuk menyeimbangkan jumlah data. Salah satu caranya dengan menambahkan ulasan positif supaya model dapat lebih efektif dalam mengenali kedua tipe ulasan itu.

Daftar Pustaka:

- [1] Indarwati, K. D. (2023). Analisis Sentimen Terhadap Kualitas Pelayanan Aplikasi Gojek Menggunakan Metode Naive Bayes Classifier. *Jatisi (Jurnal Teknik Informatika Dan Sistem Informasi)*, 10(1).
- [2] Safira, A., & Hasan, F. N. (2023). Analisis Sentimen Masyarakat Terhadap Paylater Menggunakan Metode Naive Bayes Classifier. *ZONAsi: Jurnal Sistem Informasi*, 5(1), 59-70.
- [3] Zulcharnain, R., Abdurrahman, G., & Daryanto, D. (2025). Analisis Sentimen Ulasan Duolingo dengan Metode Algoritma Multinomial Naïve Bayes. *Jurnal Informatika dan Teknologi Pendidikan*, 5(1), 1-15.
- [4] Gunawana, Y. A. V., ERa, N. A. S., Mahendraa, I. B. M., & Made, I. (2022). Analisis Sentimen Ulasan Aplikasi Transportasi Online Menggunakan Multinomial Naive Bayes dan

- Query Expansion Ranking. *J. Elektron. Ilmu Komput. Udayana p-ISSN*, 2301, 5373
- [5] Hasanah, A. N., & Sari, B. N. (2024). Analisis Sentimen Ulasan Pengguna Aplikasi Jasa Ojek Online Maxim Pada Google Play Dengan Metode Naïve Bayes Classifier. *Jurnal Informatika dan Teknik Elektro Terapan*, 12(1).
- [6] Nurraharjo, E. (2023). Analisis Sentimen Dan Klasifikasi Tweet Terkait Naiknya Kasus Omicron Menggunakan Naive Bayes Classifier. *Jurnal Informatika dan Rekayasa Elektronik*, 6(1), 1-8.
- [7] Zakasih, M. I., & Handoko, W. T. (2022). Analisis Sentimen Pengguna Twitter Tentang Nft (Non Fungible Token) Dengan Metode Naive Bayes Classifier. *Jurnal Informatika Dan Rekayasa Elektronik*, 5(2), 221-229
- [8] Prayogo, W. B. (2025). Analisis Sentimen Ulasan Pengguna terhadap Ekspedisi Layanan Logistik dengan Metode Naïve Bayes. *eProceedings of Engineering*, 12(2), 1-9.
- [9] Erfina, A., & Al-shufi, M. F. (2022). Analisis Sentimen Aplikasi Jasa Kurir Di Play Store Menggunakan Algoritma Naive Bayes. *Jurnal Sistem Informasi Dan Informatika (Simika)*, 5(2), 103-110.
- [10] Rahman, H. A., Santoso, R., & Widiharah, T. (2023). Analisis Sentimen Pada Perusahaan Penyedia Jasa Logistik J&T Menggunakan Algoritma Multinomial Naive Bayes dan Support Vector Machine. *Jurnal Gaussian*, 12(2), 242-253.
- [11] Irawansyah, R. S., & Wiriasto, G. W. (2023). ANALISIS SENTIMEN TERHADAP PROGRAM MERDEKA BELAJAR-KAMPUS MERDEKA (MBKM) PADA TWITTER MENGGUNAKAN ALGORITMA NAIVE BAYES CLASSIFIER. *Jurnal Teknologi Informasi, Komputer, dan Aplikasinya (JTika)*, 5(2), 237-244.
- [12] Agustina, N., Citra, D. H., Purnama, W., Nisa, C., & Kurnia, A. R. (2022). Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Shopee pada Google Play Store: The Implementation of Naïve Bayes Algorithm for Sentiment Analysis of Shopee Reviews On Google Play Store. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 2(1), 47-54.
- [13] Salsabila, S. M., Murtopo, A. A., & Fadhillah, N. (2022). Analisis Sentimen Pelanggan Tokopedia Menggunakan Metode Naive Bayes Classifier. *Jurnal Minfo Polgan*, 11(2), 30-35.
- [14] Styawati, S., Nurkholis, A., Aldino, A. A., Samsugi, S., Suryati, E., & Cahyono, R. P. (2022, January). Sentiment analysis on online transportation reviews using Word2Vec text embedding model feature extraction and support vector machine (SVM) algorithm. In *2021 International Seminar on Machine Learning, Optimization, and Data Science (ISMODE)* (pp. 163-167). IEEE.
- [15] Nishfi, D., & Prianto, C. (2023). Analisis Sentimen Layanan Jasa Pengiriman Pada Ulasan Play Store: Systematic Literature Review. *Jurnal Informatika Dan Teknologi Komputer (J-ICOM)*, 4(2), 87-98.
- [16] Misrun, C. A., Haerani, E., Fikry, M., & Budianita, E. (2023). Analisis sentimen komentar youtube terhadap Anies Baswedan sebagai bakal calon presiden 2024 menggunakan metode naive bayes classifier. *jurnal coscitech (computer science and information technology)*, 4(1), 207-215.
- [17] Putri, D. D., Nama, G. F., & Sulistiono, W. E. (2022). Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier. *Jurnal Informatika dan Teknik Elektro Terapan*, 10(1).
- [18] Alam, S., & Sulisty, M. I. (2023). Analisis Sentimen Berdasarkan Ulasan Pengguna Aplikasi Mypertamina Pada Google Playstore Menggunakan Metode Naïve Bayes. *Storage: Jurnal Ilmiah Teknik Dan Ilmu Komputer*, 2(3), 100-108.
- [19] Rizaldi, S. A. R., Alam, S., & Kurniawan, I. (2023). Analisis Sentimen Pengguna Aplikasi JMO (Jamsostek Mobile) Pada Google Play Store Menggunakan Metode Naive Bayes. *STORAGE: Jurnal Ilmiah Teknik dan Ilmu Komputer*, 2(3), 109-117.
- [20] Kevin, K., Enjeli, M., & Wijaya, A. (2024). Analisis Sentimen Penggunaan Aplikasi Kinemaster Menggunakan Metode Naïve Bayes. *Jurnal Ilmiah Computer Science*, 2(2), 89-98.