

ANALISIS SENTIMEN PADA ULASAN PENGGUNA APLIKASI JAMSOSTEK *MOBILE* DENGAN MENGUNAKAN NAIVE BAYES DAN LOGISTIC REGRESSION

O'neal Efrata Madao¹, Akhmad Irsyad², Muhammad Rivani Ibrahim³

1,2,3) Sistem Informasi, Fakultas Teknik, Universitas Mulawarman, Indonesia

Article Info	ABSTRACT
Article history: Received: 10 Juni 2025 Revised: 20 Juni 2025 Accepted: 29 Juni 2025	Abstrak Aplikasi Jamsostek Mobile (JMO) merupakan transformasi digital dari BPJS Ketenagakerjaan yang mulai digunakan sejak September 2021 sebagai upaya mempermudah pelayanan kepada peserta. Meskipun dirancang untuk memperluas akses dan meningkatkan mutu layanan, masih ditemukan sejumlah ulasan negatif terkait fungsi dan kinerja aplikasi. Ulasan ini menjadi sumber penting dalam memahami pengalaman pengguna. Penelitian ini bertujuan melakukan analisis sentimen terhadap ulasan pengguna aplikasi JMO dengan menerapkan dua metode klasifikasi: Naïve Bayes dan Logistic Regression. Dataset diperoleh dari Google Playstore dengan total 1.500 ulasan berbahasa Indonesia, yang kemudian dilabeli secara manual sebagai sentimen positif atau negatif. Proses analisis dilakukan dengan bahasa Python melalui platform Google Colab. Evaluasi kinerja model dilakukan menggunakan metrik akurasi, presisi, recall, dan F1-score. Hasil evaluasi menunjukkan bahwa kedua algoritma mencapai akurasi sebesar 92,67%. Naïve Bayes unggul dalam presisi sebesar 97,58%, sedangkan Logistic Regression lebih baik dalam recall (93,30%) dan F1-score (93,82%). Walaupun keduanya menunjukkan kinerja optimal, Logistic Regression dinilai lebih seimbang dalam melakukan klasifikasi. Pemilihan metode terbaik dapat disesuaikan dengan fokus analisis, apakah lebih menekankan pada ketepatan klasifikasi atau kelengkapan deteksi sentimen. Kata Kunci: Jamsostek <i>Mobile</i> (JMO), Analisis Sentimen, <i>Naïve Bayes</i>, <i>Logistic Regression</i> Abstract The Jamsostek Mobile (JMO) application is a digital transformation initiative by BPJS Ketenagakerjaan that has been in use since September 2021 to simplify service delivery to participants. Although designed to improve accessibility and service quality, numerous negative reviews have emerged regarding the application's functionality and performance. These user reviews serve as a valuable source for understanding the user experience. This study aims to conduct sentiment analysis on user reviews of the JMO application using two classification methods: Naïve Bayes and Logistic Regression. The dataset was obtained from the Google Play Store, consisting of 1,500 user reviews written in Indonesian, which were manually labeled as either positive or negative sentiments. The analysis was carried out using Python on the Google Colab platform. Model performance was evaluated using accuracy, precision, recall, and F1-score as metrics. The evaluation results showed that both algorithms achieved an accuracy of 92.67%. Naïve Bayes outperformed in terms of precision, reaching 97.58%, while Logistic Regression performed better in recall (93.30%) and F1-score (93.82%). Although both algorithms delivered optimal performance, Logistic Regression was found to be more balanced in classification. The choice of the most suitable method may depend on the focus of the analysis, whether prioritizing classification accuracy or completeness in sentiment detection. Keywords: Jamsostek <i>Mobile</i> (JMO), Sentiment Analysis, <i>Naïve Bayes</i>, <i>Logistic Regression</i>.

Djtechno: Jurnal Teknologi Informasi oleh Universitas Dharmawangsa Artikel ini bersifat open access yang didistribusikan di bawah syarat dan ketentuan dengan Lisensi Internasional Creative Commons Attribution NonCommercialL ShareAlike 4.0 ([CC-BY-NC-SA](https://creativecommons.org/licenses/by-nc-sa/4.0/)).



Corresponding Author:

E-mail : onealefrata123@gmail.com

1. PENDAHULUAN

Badan Penyelenggara Jaminan Sosial Ketenagakerjaan (BPJAMSOSTEK) merupakan lembaga penyelenggara jaminan sosial yang bertanggung jawab memberikan perlindungan bagi tenaga kerja atas risiko ekonomi dan sosial, termasuk risiko kecelakaan kerja, kematian, dan masa pensiun. Untuk mempermudah akses terhadap layanan, BPJS Ketenagakerjaan meluncurkan Jamsostek Mobile (JMO) sebagai bentuk inovasi berbasis digital. Aplikasi ini menyediakan berbagai fitur seperti layanan klaim Jaminan Hari Tua (JHT), akses investasi, informasi dana siaga, fasilitas perumahan pekerja, e-wallet, hingga promosi dan layanan streaming[1].

Penggunaan aplikasi mobile menjadi sarana utama dalam meningkatkan layanan BPJS Ketenagakerjaan bagi seluruh tenaga kerja di Indonesia. Aplikasi JMO dirilis pada September 2021 oleh BPJS Ketenagakerjaan. Meskipun aplikasi ini dirancang untuk meningkatkan pelayanan dan mempermudah pengguna dalam mengakses layanan BPJS Ketenagakerjaan, terdapat pengguna yang menyatakan ketidakpuasan terkait fitur serta kualitas layanan yang disediakan [2]. Ulasan pengguna di Play Store memperlihatkan beberapa keluhan, seperti “sering lemot”, “tidak bisa cek saldo”, dan “setelah diperbarui justru error”, yang mengindikasikan tingkat ketidakpuasan pengguna terhadap kinerja aplikasi.

Ulasan pengguna berperan penting untuk melihat pandangan dan pikiran users terhadap produk, serta memberikan masukan untuk instansi dapat melakukan perbaikan. Oleh karena itu, diperlukan analisis mendalam untuk mencari tahu sentimen dari ulasan pengguna. Analisis sentimen adalah pendekatan yang digunakan untuk menggali informasi mengenai opini seseorang terhadap suatu peristiwa atau isu melalui data berbasis teks. Teknik ini berguna dalam mengidentifikasi kecenderungan sikap masyarakat secara umum, baik dalam bentuk positif, negatif, maupun netral [3]. Pada permasalahan ini, dua algoritma yang umum digunakan dalam penerapan analisis

sentimen adalah Naïve Bayes serta Logistic Regression, karena keduanya memiliki keunggulan dalam pengelompokan teks.

Naïve Bayes adalah algoritma berbasis peluang yang sangat efektif untuk memproses kumpulan data besar dan menganalisis teks. Algoritma ini memiliki keunggulan pada kemampuan dalam menangani gangguan komunikasi data serta tetap memberikan hasil yang baik meskipun dengan data yang relatif kecil. Algoritma ini juga memiliki efisiensi komputasi yang tinggi, sehingga cocok untuk data skala besar. Hal ini menunjukkan efektivitasnya naïve bayes dalam berbagai kasus analisis data [4].

Algoritma Logistic Regression dikenal dengan algoritma yang sangat terinterpretasi, terutama dalam klasifikasi biner, seperti pengelompokan sentimen positif dan negatif. Algoritma ini unggul dalam menghasilkan model yang mudah dipahami karena modelnya bersifat linier dan memprediksi kemungkinan data masuk ke dalam salah satu kelas positif atau negatif. Logistic Regression juga sering digunakan dalam analisis statistik dan *mechine learning* karena keandalannya dalam menangani klasifikasi biner [5].

Jika dibandingkan dengan algoritma lain, seperti Support Vector Machines (SVM), algoritma ini mampu memberikan akurasi yang tinggi, namun memerlukan waktu komputasi yang lebih lama, terutama pada dataset berukuran besar. Hal yang sama berlaku untuk Random Forest, yang cenderung lebih kompleks dalam interpretasi dan memiliki risiko overfitting jika tidak ditangani dengan baik, meskipun dapat menghasilkan prediksi yang akurat [6].

Oleh karena itu, penelitian ini menggunakan logistic regression yang memberikan keseimbangan antar performa, efisien komputasi, dan dapat dipahami dengan jelas [7]. Sementara itu, Naïve Bayes memberikan Solusi yang cepat dan efisien untuk pengolahan data dalam jumlah besar, dan Logistic Regression memberikan memberikan hasil yang lebih jelas terhadap hasil analisis sentimen [8].

Penelitian sebelumnya, "Perbandingan Algoritma Naïve Bayes dan KNN dalam Analisis Sentimen Ulasan Pengguna Aplikasi Capcut" menunjukkan hasil algoritma Naïve Bayes memberikan performa yang lebih baik dalam hal akurasi 79,41% dibandingkan KNN yang hanya 75,63%, terutama pada rasio data 80:20 [9]. Sementara itu, dalam penelitian "Perbandingan Performa Algoritma SVM dan Logistic Regression untuk Analisis Sentimen Pengguna Aplikasi Retail di Android," menunjukkan hasil SVM serta

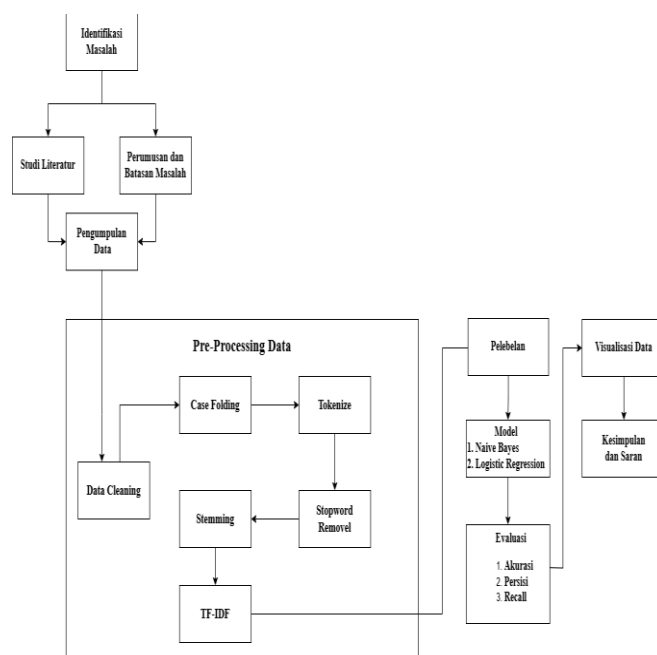
regresi logistik mempunyai nilai akurasi yang sama, tetapi Logistic Regression terbukti lebih unggul dalam hal kemampuan identifikasi skor F1 0,89, lebih sedikit misklasifikasi (110), serta hasil *True Positive* dan *True Negative* yang lebih unggul [10]. Dapat disimpulkan, algoritma Logistic Regression dan Naïve Bayes lebih efektif untuk menganalisis sentimen terutama untuk ulasan aplikasi.

Solusi dari permasalahan dalam meningkatkan aplikasi Jamsostek Mobile (JMO) adalah perlu adanya analisis pada ulasan pengguna aplikasi JMO, dengan menerapkan model Naïve Bayes dan Logistic Regression, yang memungkinkan BPJS Ketenagakerjaan untuk memprioritaskan penanganan ulasan dengan sentimen negatif yang paling banyak di temukan. Hasil analisis dapat menjadi masukan yang berharga dalam upaya peningkatan mutu layanan serta memperbaiki masalah yang ada pada aplikasi JMO [11].

Berdasarkan penjelasan paparan tersebut, penelitian ini mengangkat judul “Analisis Sentimen Pada Ulasan Pengguna Aplikasi Jamsostek *Mobile* (JMO) Menggunakan Naïve Bayes dan Logistic Regression”, yang bertujuan untuk mengetahui ulasan pengguna pada aplikasi Jamsostek *Mobile* (JMO), Dengan penerapan metode klasifikasi Naïve Bayes dan Logistic Regression.

2. METODE PENELITIAN

Proses penelitian diawali dengan perumusan permasalahan, kemudian dilanjutkan dengan preprocessing data melalui enam tahap: *Cleaning*, *Case Folding*, *Tokenizing*, *Stopword Removal*, *Stemming*, dan TF-IDF. Metode klasifikasi yang digunakan adalah Naïve Bayes dan Logistic Regression. Seluruh proses dijalankan di Google Colab menggunakan Python. Alur metode Seperti yang disajikan dalam Gambar 1.



Gambar 1. Gambar Diagram Alur Metode Penelitian

2.1 Crawling Data

Data ulasan pengguna aplikasi JMO dilakukan secara otomatis menggunakan library Google Play Scraper melalui teknik web scraping pada *platform* Google Colab dengan bahasa pemrograman Python. Proses ini bertujuan memperoleh data seperti rating, komentar, dan tanggal ulasan secara otomatis untuk keperluan analisis sentimen.

2.2 Text Preprocessing

Review pengguna yang diperoleh melalui Google Play Store diolah melalui tahapan preprocessing untuk mengurangi noise dan menyiapkan data bagi analisis sentimen. Tahapan yang dilakukan mencakup proses pembersihan data (*cleaning*), *case folding*, *tokenization*, *stopword removal*, *stemming*, serta transformasi TF-IDF, guna memastikan data lebih akurat dan siap digunakan oleh model machine learning [12].

1. Data Cleaning

Pada tahapan *data cleaning* ini penulis bertujuan untuk menghapus seluruh karakter selain huruf yang tidak digunakan dalam proses analisis sentimen. Penghapusan karakter selain huruf tersebut akan memperkecil ukuran data, sehingga kinerja model akan lebih efektif dan efisien.

2. Case Folding

Tahap *Case Folding* berfungsi untuk menyamaratakan semua huruf yang ada ke dalam *lowercase*. Hasil dari proses ini akan membuat data lebih mudah di analisis.

3. Stopword Removal

Stopword Removal digunakan untuk menghilangkan kalimat umum seperti konjungsi dan preposisi yang Bersifat kosong atau tidak bermakna secara konteks signifikan dalam analisis sentimen, misalnya "*the*", "*is*", dan sebagainya, agar teks lebih relevan untuk diproses oleh model.

4. Tokenize

Tokenizing merupakan langkah untuk membagi teks panjang menjadi potongan-potongan kecil berupa token atau kata individual.

5. Stemming

Stemming berfungsi Untuk menghapus awalan dan akhiran kata secara tepat, dengan memperhatikan *prefix* termasuk sufiks yang sering melekat pada kata dalam bentuk dasarnya, seperti "*me-*", "*ber-*", "*-kan*", dan lainnya, Agar kata yang mengandung arti serupa dapat disatukan ke Pada satu Jenis dasar.

6. Term Frequency –Inverse Document Frequency

TF-IDF digunakan untuk menilai relevansi suatu istilah dalam dokumen melalui pemberian bobot kata. TF menghitung frekuensi kemunculan kata, sementara IDF memperkecil bobot istilah yang paling dominan dalam kemunculannya di banyak data. Fungsi TF-IDF memperhitungkan tiga faktor penting sebagai berikut [13]:

a. Term Frequency (TF)

TF dihitung untuk mengetahui seberapa sering suatu kata muncul dalam sebuah dokumen. Perhitungannya ditunjukkan pada Persamaan (1).

$$TF(d, t) = f(d, t) \quad (1)$$

Diketahui :

f = Frekuensi kemunculan

t = term (Frasa/kata)

d = dokumen

b. *Inverse Document Frequency* (IDF)

IDF dihitung untuk mengetahui seberapa penting suatu kata dalam konteks koleksi dokumen yang lebih besar. Perhitungannya ditunjukkan pada Persamaan (2).

$$IDF = \text{LOG}\left(\frac{D}{d f t}\right) \quad (2)$$

Diketahui :

D = jumlah dokumen yang memuat *keyword* (kata kunci)

d f (t) = banyaknya dokumen yang mengandung istilah atau term (t)

c. *TF dan IDF*

etelah proses perhitungan nilai TF-IDF dilakukan, maka dilakukan perhitungan bobot skor pada dokumen yang diperoleh dari hasil perhitungan TF dan IDF. Perhitungannya ditunjukkan pada Persamaan (3).

$$W(d,t) = TF \times IDF \quad (3)$$

Diketahui :

W = bobot setiap dokumen yang memuat *keyword* (kata kunci)

TF = *Term Frequency*

IDF = *Inverse Document Frequency*

2.3 *Labeling Data*

Proses pelabelan dilakukan secara manual dan telah divalidasi oleh Drs. 5. Rusydi Ahmad, M.Hum., dosen pakar linguistik dari Universitas Mulawarman, guna memastikan bahwa pelabelan tersebut sesuai dengan kaidah kebahasaan.

2.4 *Model*

Data yang telah diberi label sentimen kemudian dipisahkan dengan proporsi 80 persen untuk pelatihan dan 20 persen untuk pengujian. Klasifikasi Diterapkan melalui penggunaan algoritma Naïve Bayes dan Logistic Regression untuk menganalisis sentimen dan membandingkan kinerja kedua model [14].

1. Naïve Bayes

Naïve Bayes merupakan salah satu algoritma dalam kelompok supervised learning yang diaplikasikan setelah data diberi label dan dibagi sesuai kebutuhan pelatihan dan pengujian. Dengan mengasumsikan bahwa fitur-fitur tidak saling bergantung, algoritma ini terbukti efektif dalam mengelola klasifikasi berbasis teks [15], termasuk analisis sentimen ulasan pengguna aplikasi JMO. Perhitungannya ditunjukkan pada Persamaan (4).

$$p(y|x) = \frac{p(x|y)p(y)}{p(x)} \quad (4)$$

Diketahui :

(P) : Probabilitas

(Y) : Pernyataan bahwa kelas spesifik

(X) : Pernyataan bahwa kelas belum diketahui

(Y|X) : Probabilitas dari kelas yang berasal dari hipotesis sebelumnya

P(X) : Probabilitas dari Y

P(Y|X) : Hasil perkalian antara *likelihood* dan *prior* dibagi *evidence* *Likelihood* merupakan Probabilitas atribut X pada kelas Y, *prior* merupakan Probabilitas kelas Y dari total data set.

2. Logistic Regression

Logistic Regression digunakan untuk memprediksi nilai dependen biner (positif atau negatif) berdasarkan kombinasi linier variabel input, yang dikonversi menjadi probabilitas melalui fungsi sigmoid [16]. Perhitungannya ditunjukkan pada Persamaan (5).

$$P = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \varepsilon_1 \quad (5)$$

Diketahui :

P : Probabilitas kejadian suatu peristiwa atau suatu kelas (contoh, Y = 1)

β_0 : Intercept atau konstanta model

$\beta_1, \beta_2, \dots, \beta_n$: Koefisien regresi yang menunjukkan pengaruh masing-masing variabel prediktor terhadap probabilitas kejadian

$x_1, x_2, \dots, x_n$: Variabel prediktor (independen) yang digunakan untuk memprediksi probabilitas kejadian

ε_1 : kesalahan acak, $\varepsilon_i \sim N(0, \sigma^2)$ yang tidak berkorelasi

2.5 Visualisasi Data

Tahap selanjutnya melakukan visualisasi data dalam bentuk *wordcloud* untuk mengetahui kata apa saja yang paling banyak terdapat pada ulasan pengguna aplikasi Jamsostek Mobile (JMO) dan berapa frekuensi kata yang muncul dalam sebuah ulasan.

2.6 Evaluasi

Pada tahap ini melakukan pemodelan menggunakan dua model yaitu Naïve Bayes dan Logistic Regression maka tahapan selanjutnya adalah melakukan evaluasi berdasarkan nilai *confusion matrix* yaitu dengan menentukan tiga macam performa machine learning yaitu tingkat akurasi, ketepatan (presisi), sensitivitas (recall), serta nilai F1-Score pada ulasan pengguna aplikasi Jamsostek *Mobile* (JMO).

3. HASIL DAN PEMBAHASAN

3.1 Pengambilan Data

Hasil dari proses scraping menghasilkan total sebanyak 1.500 ulasan yang berhasil dikumpulkan. Hasil pengambilan data ditampilkan pada Tabel 1.

Tabel 1. Hasil Pengumpulan Data

No	Content
1	sangat membantu
2	abis diperbaharui malah eror
3	ok informatif
4	baguuuuuuussss 👍
5	baru di download langsung minta di update.. minimal tunggu beberapa bulan keg.. aplikasi sampah, ngerugiin org aja
...	...
1449	udh susah sekarang untuk Klim JHT-nya
1500	halo ini pembaruan aplikasi pun masi banyak yg bug terutama dibagian Live detection face. tolong segera diupdate agar memudahkan terimakasih

3.1 Text-Preprocessing

Data ulasan pengguna JMO yang beragam memerlukan *text preprocessing* untuk membersihkan dan menstandarisasi teks sebelum analisis sentimen. Tahapan

preprocessing mencakup *case folding*, penghapusan tanda baca dan angka, tokenisasi, *stopword removal*, dan *stemming* guna meningkatkan kualitas dan akurasi data untuk klasifikasi. Data berikut ini telah diproses melalui serangkaian tahapan *text-processing* dapat dilihat pada Tabel 2.

Tabel 2. Data Hasil Tahapan Pra-Pemrosesan

No	Content
1	dulu lancar tapi telah di baruh kok malah mental terus mohon di baik
2	sangat bantu
3	sering lemot
4	bagus
5	informasi nya falid
...	...
1449	tiap mau check saldo jmo pasti suruh update mulu heran sama jm
1500	japuk apnya

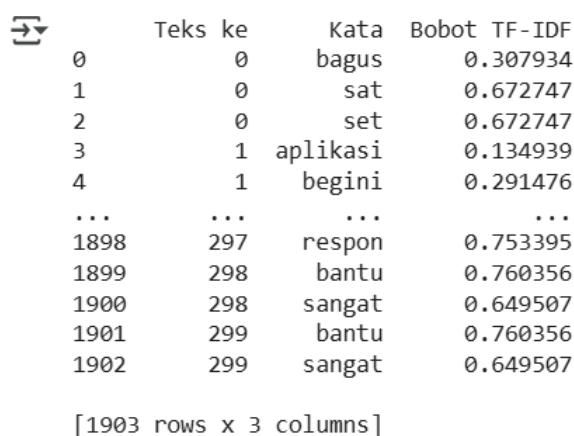
Berikutnya melakukan pembobotan kata menggunakan TF-IDF untuk mengukur frekuensi relatif dan tingkat kepentingan suatu kata dalam dokumen. Semakin sering kata muncul di banyak dokumen, bobotnya akan semakin rendah karena dianggap kurang informatif. Hasil perhitungan bobot kata menggunakan TF-IDF Dapat dilihat pada ilustrasi Gambar 2 dan Gambar 3.

TF-IDF Matrix:

	abdet	acc	ada	adalah	adm	admin	administrasi	\
0	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1	0.000000	0.000000	0.231334	0.000000	0.000000	0.000000	0.000000	0.000000
2	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
3	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
4	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
...	...	...	...	...	...	...	...	...
1195	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1196	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1197	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1198	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1199	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000

	agak	agar	aju	...	yakarna	yakin	yang	yasaya	\
0	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
1	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
2	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
3	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
4	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
...	...	...	...	...	...	...	...	...	...
1195	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
1196	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
1197	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
1198	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000
1199	0.000000	0.000000	0.000000	...	0.000000	0.000000	0.000000	0.000000	0.000000

Gambar 2. Hasil Pembobotan Kata TF-IDF



	Teks ke	Kata	Bobot TF-IDF
0	0	bagus	0.307934
1	0	sat	0.672747
2	0	set	0.672747
3	1	aplikasi	0.134939
4	1	begini	0.291476
...	...	...	...
1898	297	respon	0.753395
1899	298	bantu	0.760356
1900	298	sangat	0.649507
1901	299	bantu	0.760356
1902	299	sangat	0.649507

[1903 rows x 3 columns]

Gambar 3. Output Bobot TF-IDF Dalam *Non-Zero*

Gambar 3 memperlihatkan matriks TF-IDF dengan sebagian besar nilai nol, dan nilai non-nol menunjukkan bobot kata berdasarkan frekuensi kemunculan. Gambar 4 perubahan dari sparse matrix TF-IDF ke dalam bentuk tabel yang lebih terstruktur agar lebih mudah dibaca dan dianalisis, di mana setiap baris menampilkan kata-kata yang memiliki bobot TF-IDF *non-zero* terhadap dokumen tertentu.

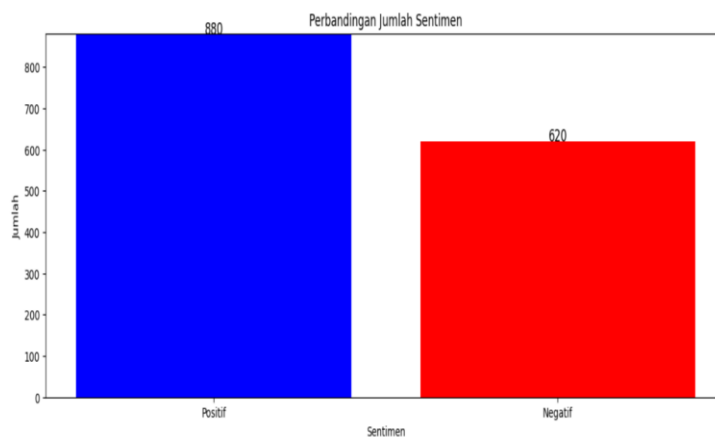
3.2 Pelabelan Data

Proses pelabelan dilakukan untuk mengelompokkan 1.500 ulasan aplikasi Menjadi dua klasifikasi, yaitu kalimat positif (1) dan negatif (0). Berikut data yang telah dilakukan pelabelan dapat dilihat pada Tabel 3.

Tabel 3. Pelabelan Sentimen

No	Content	Label
1	dulu lancar tapi telah di baruh kok malah mental terus mohon di baik	Negatif
2	sangat bantu	Positif
3	sering lemot	Negatif
4	bagus	Positif
5	informasi nya falid	Positif
...	...	...
1449	tiap mau check saldo jmo pasti suruh update mulu heran sama jm	Negatif
1500	japuk apknya	Negatif

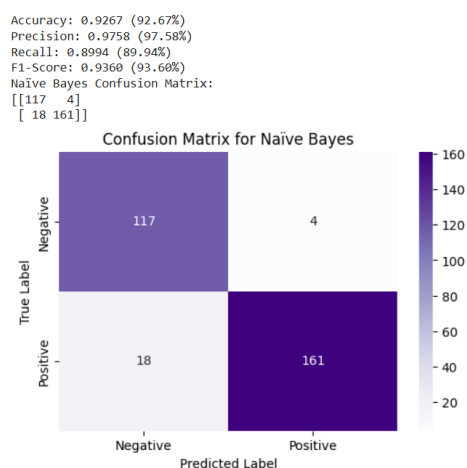
Selanjutnya perbandingan sentimen ulasan pengguna aplikasi JMO menunjukkan dominasi sentimen positif sebanyak 880 ulasan dibandingkan 620 ulasan negatif. Visualisasi perbandingan sentimen ini disajikan *bar chart* pada Gambar 4.



Gambar 4. Jumlah Sentimen Positif dan Negatif

3.3 Model

Hasil algoritma Naïve Bayes ditampilkan pada *confusion matrix* dengan 117 *true negative*, 161 *true positive*, 4 *false positive*, dan 18 *false negative*. Model ini mencapai akurasi 92,67%, presisi 97,58%, *recall* 89,94%, dan F1-score 93,60%. Visualisasi ditampilkan dalam Gambar 5.



Gambar 5. Visualisasi *confusion matrix* Algoritma Naive Bayes

Hasil algoritma Logistic Regression ditampilkan pada *confusion matrix* dengan 111 *true negative*, 167 *true positive*, 10 *false positive*, dan 12 *false negative*. Model ini menunjukkan kinerja yang baik dengan akurasi sebesar 92,67%, presisi 94,35%, *recall* 93,30%, dan F1-score 93,82%. Visualisasi ditampilkan dalam Gambar 6.



Sebagai tahap akhir setelah proses pemodelan sentimen dilakukan, representasi *wordcloud* dimanfaatkan untuk menggambarkan Kata-kata dengan tingkat kemunculan paling dominan dalam teks pada masing-masing label sentimen yang terbagi ke dalam dua jenis, yakni positif dan negatif. Visualisasi tidak digunakan untuk tahap eksplorasi awal, melainkan sebagai pendukung interpretasi hasil klasifikasi dari data *test*.



Tabel 4. a. *Visualisasi Wordcloud* Ulasan Positif dan b. Ulasan Negatif *Visualisasi Wordcloud*

Visualisasi *Wordcloud* dibuat Secara individual untuk setiap kategori kelompok sentimen guna memberikan gambaran umum terhadap fokus perhatian dan persepsi pengguna. Ulasan dengan sentimen positif sebagian besar terdiri dari kata-kata seperti “bagus,” “bantu,” “baik,” dan “mudah” yang menunjukkan kepuasan terhadap fitur aplikasi. Sementara itu, pada ulasan negatif, kata-kata seperti “tidak,” “bisa,” “login,” dan “error” paling sering muncul, menunjukkan adanya keluhan terkait kendala teknis dan akses layanan. Hasil visualisasi ini ditampilkan pada Tabel 4.

3.5 Evaluasi

Pada tahap evaluasi, dilakukan perbandingan hasil klasifikasi antara algoritma Naïve Bayes dan Logistic Regression. Keduanya menunjukkan akurasi sama sebesar 92,67%. Naïve Bayes unggul dalam *precision* (97,58%), sedangkan Logistic Regression lebih tinggi pada *recall* (93,30%) dan *F1-score* (93,82%). Perbandingan lengkap ditampilkan pada Tabel 5.

Tabel 5. Hasil perbandingan Naïve Bayes dan Logistic Regression

Metode	Accuracy	Precision	Recall	F1-Score
Naïve Bayes	92.67%	97.58%	89.94%	93.60%
Logistic Regression	92.67%	94.35%	93.30%	93.82%

4. SIMPULAN

Berdasarkan hasil evaluasi, dapat disimpulkan bahwa baik Naïve Bayes maupun Logistic Regression Efektif dalam melakukan klasifikasi terhadap data sentimen ulasan Pemakai aplikasi JMO Menunjukkan performa akurasi yang setara. Namun, Naïve Bayes lebih tepat dalam memprediksi sentimen positif, sedangkan Logistic Regression lebih baik dalam mendeteksi seluruh data yang benar-benar positif. Hal ini tercermin dari nilai *precision* tertinggi pada Naïve Bayes dan *recall* tertinggi pada Logistic Regression. Logistic Regression melihatkan performa yang lebih seimbang di antara nilai presisi dan sensitivitas (*recall*), dilihat dari skor F1 yang menunjukkan peningkatan tipis. Temuan ini menjawab pertanyaan penelitian mengenai efektivitas kedua algoritma, dengan menegaskan bahwa pemilihan metode terbaik bergantung pada fokus analisis, apakah lebih mengutamakan ketepatan atau kelengkapan deteksi.

REFERENCES

- [1]. Melvern Pradana. (2021, June 1). Apa Itu BPJS Ketenagakerjaan? Investbro.Id. <https://investbro.id/Apa-Itu-Bpjs-Ketenagakerjaan/>
- [2]. Arisoemaryo, B. S., & Prasetio, R. T. (2022). Analisis Tingkat Kepuasan Pengguna Aplikasi Jamsostek Mobile Menggunakan Metode End User Computing Satisfaction. *Jurnal Responsif*, 4(1), 110–117. <https://ejurnal.ars.ac.id/index.php/jti>
- [3]. Prayoga Siswono, A., Fauzi, S., Zalva Surya Hermawan, A., Riyandi, A., Panjaitan No, J. DI, Kidul, P., Purwokerto Selatan, K., & Banyumas, K. (2024). Analisis Sentimen Pelantikan Presiden Indonesia 2024 Menggunakan Model Klasifikasi Dan Algoritma Naive Bayes (Vol. 4, Issue 1). <https://doi.org/https://conferences.itelkom-pwt.ac.id/index.php/centive/article/view/351>
- [4]. Krisna Perdana Jaya Sitompul, Adi Rizky Pratama, & Kiki Ahmad Baihaqi. (2023). Komparasi Algoritma Naive Bayes, Support Vector Machine, Dan Logistic Regression Pada Analisis Sentimen Pengguna Aplikasi Transportasi Online. *Kumpulan Jurnal Ilmu Komputer (KLIK)*, 10. <https://doi.org/https://klik.ulm.ac.id/index.php/klik/article/view/6>
- [5]. Hanif, I. F., Affandi, I. R., Hasan, F. N., Sinduningrum, E., & Halim, Z. (2022). Analisis Sentimen Opini Masyarakat Terkait Penyelenggaraan Sistem Elektronik Menggunakan Metode Logistic Regression. *Jurnal Linguistik Komputasional*, 5(2), 77–84. <https://doi.org/https://simakip.uhamka.ac.id/download?type=jurnal&id=2976>
- [6]. Barreñada, L., Dhiman, P., Timmerman, D., Boulesteix, A.-L., & Van Calster, B. (2024). Understanding Overfitting In Random Forest For Probability Estimation: A Visualization And Simulation Study. *Diagnostic And Prognostic Research*, 8(1), 14. <https://doi.org/10.1186/S41512-024-00177-1>
- [7]. Bahtiar, S. A. H., Dewa, C. K., & Luthfi, A. (2023). Comparison Of Naive Bayes And Logistic Regression In Sentiment Analysis On Marketplace Reviews Using Rating-Based Labeling. *Journal Of Information Systems And Informatics*, 5(3), 915–927. <https://doi.org/10.51519/journalisi.v5i3.539>
- [8]. Kisma, A. J. N., Widiawati, C. R. A., & Suliswaningsih, S. (2023). Analysis Of Applications In Playstore Based On Rating And Type Using Naive Bayes And Logistic Regression. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, 10(2), 174–184. <https://doi.org/https://jurnal.mdp.ac.id/index.php/jatisi/article/view/4784>
- [9]. Muslim, S. N. S., Nurdiansyah, F., & Rahman, A. Y. (2024). Perbandingan Algoritma Naive Bayes Dan Knn Dalam Analisis Sentimen Ulasan Pengguna Aplikasi Capcut. *Jurnal Informatika Dan Teknik Elektro Terapan*, 12(3S1). <https://doi.org/10.23960/jitet.v12i3s1.5156>
- [10]. Budianto, A. G., Rusilawati, R., Suryo, A. T. E., Cahyono, G. R., Zulkarnain, A. F., & Martunus, M. (2024). Perbandingan Performa Algoritma Support Vector Machine (SVM) Dan Logistic Regression Untuk Analisis Sentimen Pengguna Aplikasi Retail Di Android. *Jurnal Sains Dan Informatika*, 10(2). <https://doi.org/10.34128/jsi.v10i2.911>
- [11]. Umat, Y. N. K., Nafsyi, D. R., Kusumaningsih, D., & Hakim, L. (2024). Analisis Faktor Yang Mempengaruhi Pemilihan Gubernur Daerah Khusus Jakarta Menggunakan Algoritma Naive Bayes Dan Regresi Logistik. *Rabit: Jurnal Teknologi Dan Sistem Informasi Univrab*, 9(2), 211–224. <https://doi.org/https://pdfs.semanticscholar.org/7ddf/3331e9c37b54b913b4d4a99b667841c806ac.pdf>
- [12]. Firdlous, D. A., Andrian, R., & Widodo, S. (2023). Sentiment Analysis Public Twitter On 2024 Election Using The Long Short Term Memory Model. *SISTEMASI*, 12(1), 52. <https://doi.org/10.32520/stmsi.v12i1.2145>
- [13]. Septiani, D., & Isabela, I. (2022). SINTESIA: Jurnal Sistem Dan Teknologi Informasi Indonesia Analisis Term Frequency Inverse Document Frequency (TF-IDF) DALAM TEMU KEMBALI INFORMASI PADA DOKUMEN TEKS. *Jurnal Sistem Dan Teknologi Informasi Indonesia*, 1. <https://doi.org/https://journal.unj.ac.id/unj/index.php/SINTESIA/article/view/39364>
- [14]. Atimi, R. L., & Enda Esyudha Pratama. (2022). Implementasi Model Klasifikasi Sentimen Pada Review Produk Lazada Indonesia. *Jurnal Sains Dan Informatika*, 8(1), 88–96. <https://doi.org/10.34128/jsi.v8i1.419>
- [15]. Azka, F. W., & Romadhony, A. H. (2022). Sentiment Analysis Of University Social Media Using Support Vector Machine And Logistic Regression Methods. <https://doi.org/10.34818/indojc.2022.7.2.638>
- [16]. Aldren Marpaung, J., & Devega, M. (2024). Analisis Sentimen Kepuasan Pengguna Aplikasi BCA Mobile Menggunakan Metode Naive Bayes Dan Support Vector Machine (SVM). In *Prosiding-Seminar Nasional Teknologi Informasi & Ilmu Komputer (SEMATER)* (Vol. 3, Issue 1). <https://doi.org/https://pustaka-psm.unilak.ac.id/index.php/semaster/article/view/23980>