

PENERAPAN NLP (NATURAL LANGUAGE PROCESSING) DALAM ANALISIS SENTIMEN PENGGUNA TELEGRAM DI PLAYSTORE

Nurwanda¹, Nana Suarna², Willy Prihartono³

^{1,2} Teknik Informatika, STMIK IKMI Cirebon

³ Komputerisasi Akuntansi, STMIK IKMI Cirebon

Jalan Perjuangan No. 10 B Cirebon, Indonesia

¹wanur2039@gmail.com

ABSTRAK

Pendapat umum pengguna tentang aplikasi Telegram di Play Store umumnya positif, negatif, atau netral. Penelitian ini bertujuan untuk menilai sentimen pengguna terhadap fitur-fitur aplikasi Telegram, yang telah menjadi salah satu aplikasi pesan instan paling populer. Analisis sentimen ulasan pengguna di Play Store dapat memberikan wawasan tentang kepuasan pengguna dan masalah yang dihadapi. Permasalahan penelitian ini adalah menemukan sentimen pengguna terhadap Telegram di Play Store. Metode pemrosesan bahasa alami (NLP) digunakan untuk menganalisis sentimen pengguna Telegram di Play Store. Ulasan dibagi ke dalam tiga kategori: positif, negatif, dan netral. Penelitian ini juga mengidentifikasi tren umum dan masalah yang sering muncul dalam ulasan. Model NLP digunakan untuk mengklasifikasikan sentimen dan menganalisis kata kunci yang sering muncul. Hasil evaluasi menunjukkan bahwa analisis sentimen ulasan pengguna Telegram di Play Store dengan metode NLP memiliki tingkat akurasi 85,31%. Rasio pembagian data untuk evaluasi adalah 80:20, dengan nilai presisi 93%, recall 76%, dan F1-Score 88%. Penelitian ini memberikan wawasan untuk pengembang dalam meningkatkan aplikasi dan memastikan kepuasan pengguna yang berkelanjutan.

Kata kunci: Analisis Sentimen, Telegram, Play Store, Natural Language Processing-

1. PENDAHULUAN

Peningkatan penggunaan aplikasi pesan instan seperti Telegram telah menciptakan kebutuhan untuk memahami pandangan dan pengalaman pengguna dengan lebih baik. Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap ulasan pengguna Telegram di Play Store menggunakan metode Natural Language Processing (NLP). NLP merupakan pendekatan inovatif yang memungkinkan komputer untuk memahami dan menganalisis bahasa manusia secara otomatis, membuka peluang untuk mendapatkan wawasan yang lebih dalam terkait dengan persepsi dan respons pengguna terhadap aplikasi tersebut. Dengan melibatkan teknologi NLP, penelitian ini bertujuan untuk mengatasi tantangan dalam memproses bahasa alami, seperti ejaan yang salah, variasi gaya penulisan, dan kemampuan model untuk menangkap sentimen subjektif dengan lebih akurat. Dengan demikian, hasil penelitian ini diharapkan dapat memberikan kontribusi yang signifikan untuk pemahaman lebih lanjut tentang penggunaan Telegram berdasarkan sudut pandang pengguna yang tercermin dalam ulasan mereka di Play Store [1][2].

Ulasan pengguna sering kali mengandung banyak omong kosong seperti: ejaan, singkatan, dan bahasa informal yang tidak tepat. Gaya penulisan yang berbeda mungkin menyulitkan perolehan hasil analisis yang akurat. Meskipun sering kali dihilangkan selama prapemrosesan, menghilangkan kata-kata berhenti atau kata-kata pendek dapat menghilangkan konteks penting dari teks. Beberapa kalimat dan ulasan mungkin mengandung sentimen yang bermakna meskipun hanya terdiri dari kata-kata berhenti.

Beberapa ulasan menyiratkan perasaan subjektif yang sulit ditentukan. Misalnya, kalimat yang mengandung metafora atau ironi sulit ditentukan nada sebenarnya. Suasana bervariasi tergantung pada situasi dan budaya. Pemahaman suatu kata atau frasa dalam suatu budaya mungkin berbeda dengan penafsirannya dalam budaya lain. Hal ini dapat menyebabkan kesalahan dalam analisis sentimen[3].

Studi ini menggunakan metode Natural Language Processing (NLP) dan Algoritma Naïve Bayes untuk mempelajari pandangan dan perasaan pengguna tentang pengalaman mereka menggunakan Telegram di Play Store. Evaluasi Pengalaman Pengguna menggunakan analisis sentimen untuk menilai kepuasan dan ketidakpuasan pengguna Telegram melalui ulasan di Play Store. Ini memungkinkan mereka untuk menunjukkan area mana aplikasi dapat ditingkatkan atau diperbaiki. Meningkatkan Pengembangan Aplikasi Meningkatkan pemahaman pengembang Telegram tentang kebutuhan dan preferensi pengguna untuk membantu pengembangan aplikasi dan peningkatan kualitasnya. Memberi Informasi untuk Keputusan Bisnis: Beri penyedia layanan Telegram dan pihak terkait informasi berharga untuk membantu membuat keputusan bisnis yang lebih baik, seperti strategi pemasaran, peningkatan fitur, dan upaya retensi pengguna[4]. Metode analisis sentimen yang menggunakan metode pemrosesan bahasa natural (NLP) di Play Store mencakup pengumpulan data ulasan pengguna, preprocessing teks untuk menghilangkan suara, ekstraksi fitur menggunakan teknik NLP seperti TF-IDF atau pengembedan kata, penerapan algoritma klasifikasi seperti Support Vector Machine atau Naive Bayes, dan

evaluasi hasil untuk mengevaluasi kinerja model dalam memahami sentimen pengguna. Metode ini memungkinkan pemahaman yang lebih mendalam tentang pandangan dan persepsi pengguna terhadap aplikasi Telegram dengan menggunakan teknologi pemrosesan bahasa alami (NLP).

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Investigasi akademis sebelumnya oleh Andre Saputra dkk., (2023), telah mencakup segudang pendekatan beragam yang diterapkan untuk tujuan analisis sentimen di berbagai konteks. Contoh ilustratif dari hal ini dapat dilihat dalam upaya penelitian awal, di mana teknik Naïve Bayes digunakan untuk meneliti emosi yang terkait dengan penggunaan obrolan grup Telegram dalam kaitannya dengan lowongan pekerjaan. Penelitian ini mencapai tingkat akurasi yang terpuji, mencapai tingkat yang mengesankan sebesar 91,05% [5].

Sebuah studi yang dilakukan oleh Muktafin dkk. (2020), menjelaskan lebih lanjut tentang pentingnya ulasan produk dalam pasar Shopee, menekankan signifikansinya sebagai sumber informasi yang berharga bagi pembeli potensial dan sebagai sarana untuk memberikan umpan balik kepada penjual. Studi ini menemukan bahwa analisis sentimen yang dilakukan pada ulasan produk ini menunjukkan keefektifannya dalam menawarkan pemahaman komprehensif tentang peringkat pembeli, sehingga menggambarkan bahwa penghargaan bintang saja mungkin tidak cukup menangkap esensi dan substansi keseluruhan dari ulasan [6].

2.2. Analisis Sentimen

Analisis sentimen, juga dikenal sebagai penambangan opini, adalah proses komputasi rumit yang digunakan untuk memastikan sentimen atau emosi yang berlaku yang dirangkum dalam badan teks tertentu. Metode multifaset ini biasanya digunakan untuk membedakan dan memahami opini publik kolektif, sentimen, dan emosi yang berkaitan dengan produk, layanan, atau materi pelajaran tertentu. Contoh klasik pemanfaatan analisis sentimen terletak pada penerapannya untuk menganalisis dan mengevaluasi dengan cermat banyaknya komentar, komentar, dan ekspresi publik mengenai kebijakan peraturan yang melarang pemanfaatan sirup dalam batas-batas geografis Indonesia [7].

2.3. NLP

Natural Language Processing (NLP), sebuah subbidang kecerdasan buatan, didedikasikan untuk mengeksplorasi hubungan rumit antara komputer dan bahasa manusia. Tujuan utamanya terletak pada memahami, memeriksa, dan menghasilkan informasi tekstual dengan cara yang meniru kemampuan manusia. Contoh penerapan NLP terbukti dalam upaya penelitian, di mana ia memfasilitasi analisis sentimen masyarakat yang berkaitan dengan layanan Kesehatan

BPJS. Dengan menggunakan teknik NLP, peneliti dapat menyelidiki persepsi dan sikap individu yang bernuansa terhadap layanan kesehatan yang telah disebutkan, sehingga meningkatkan pemahaman mereka tentang opini publik [8].

2.4. Naive Bayes

Naive Bayes adalah algoritma klasifikasi berbasis probabilitas yang berdasarkan pada teorema Bayes. Algoritma ini digunakan dalam berbagai aplikasi, termasuk analisis sentimen. Naive Bayes merupakan teknik prediksi berbasis probabilitas sederhana yang berdasarkan pada penerapan teorema Bayes (atau aturan Bayes) dengan asumsi independensi (ketidaktergantungan) yang kuat (naïf) [9]. Menggunakan rumus:

$$P(H|X) = \frac{P(X|H)P(H)}{P(X)}$$

Keterangan:

X : Data dengan class yang belum diketahui

H : Hipotesis data X yang merupakan suatu class spesifik

P(H|X) : Probabilitas hipotesis H berdasarkan kondisi X (posteriori probabilitas)

P(H) : Probabilitas hipotesis H (prior probabilitas)

P(X|H) : Probabilitas X berdasarkan kondisi pada hipotesis H

P(X) : Probabilitas X

2.5. TF-IDF

TF-IDF (Term Frequency-Inverse Document Frequency) adalah metode statistik yang digunakan dalam pengolahan teks untuk mencerminkan seberapa penting sebuah kata bagi dokumen dalam kumpulan dokumen. Misalnya, TF-IDF digunakan dalam penelitian untuk menganalisis sentimen review film [10].

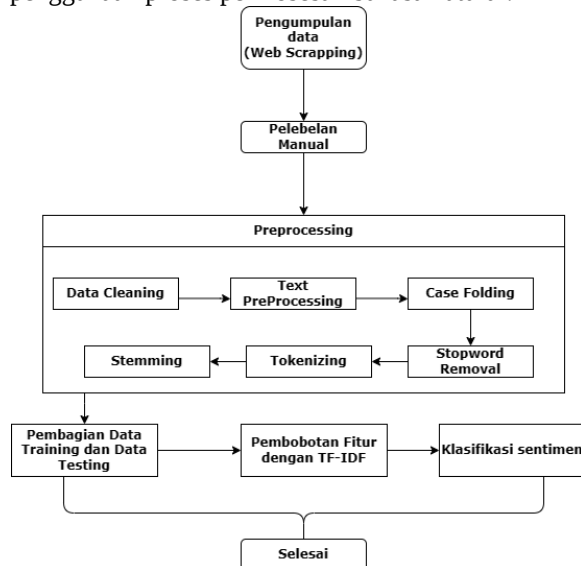
2.6. Confusion Matrix

Confusion Matrix adalah tabel yang digunakan untuk mengukur kinerja model klasifikasi. Tabel ini menunjukkan jumlah prediksi yang benar dan salah yang dibuat oleh model, dibagi menjadi empat kategori: true positives (TP), false positives (FP), true negatives (TN), dan false negatives (FN). Misalnya, dalam penelitian tentang efisiensi analisis Intrusion Detection System Alert, confusion matrix digunakan untuk mengukur metrik seperti akurasi, recall, presisi, dan F1-Score [11].

3. METODE PENELITIAN

Penelitian ini menggunakan algoritma pengklasifikasi Naive Bayes untuk menganalisis sentimen karena keandalannya dalam klasifikasi teks berdasarkan aturan dan probabilitas Bayesian. Pertama, data Shopee dikumpulkan untuk digunakan sebagai dataset selama berbelanja. Sebelum preprocessing dilakukan, metode pemrosesan bahasa natural (NLP) digunakan untuk membersihkan dan

menyusun data. Pra-pemrosesan mencakup pelipatan huruf, penghapusan kata berhenti, dan normalisasi kata. Selanjutnya, dataset dibagi menjadi dua bagian. Data uji digunakan untuk mengukur kinerja model Naive Bayesian, dan data latih digunakan untuk melatih model tersebut. Setiap langkah dijelaskan dalam tahap ini, yang memberikan gambaran umum yang lengkap. Metode yang digunakan, termasuk penggunaan proses pemrosesan bahasa natural.



Gambar 1. Metodologi Penelitian

3.1. Pengumpulan Data

Data yang digunakan untuk penelitian ini berasal dari 1000 ulasan Telegram yang dikumpulkan dari aplikasi Telegram di Google Play Store. Data ini dikumpulkan menggunakan teknik scrapping web, yang merupakan teknik untuk mengekstrak data dalam jumlah besar dari web dan kemudian disimpan dalam bentuk file CSV.

3.2. Preprocessing

Pelabelan dilakukan dengan menandai ulasan yang bersifat komentar positif atau negatif secara manual dengan mengacu pada Kamus Besar Bahasa Indonesia (KBBI). Merupakan tahapan awal yang dilakukan dalam memproses teks, dimana terdapat beberapa proses didalamnya yaitu data cleaning untuk membersihkan data dengan menghilangkan simbol-simbol tertentu, text processing untuk menganalisis dan menyortir data teks yang tidak terstruktur guna mendapatkan wawasan berharga, case folding merupakan proses mengganti semua huruf besar menjadi huruf kecil pada data, stopwords removal yaitu menghilangkan kata yang dianggap tidak berpengaruh dalam kalimat, tokenization memisahkan kalimat menjadi kata per kata, dan stemming menghapus awalan dan akhiran kata untuk menghasilkan kata dasar. Kemudian dilakukan pembagian dataset menjadi 2 jenis yaitu data latih dan data uji. Data latih (training) adalah data yang diolah dan hasilnya sebagai prediksi untuk data uji (testing). Penelitian ini menggunakan rasio perbandingan data 80:20.

3.3. TF-IDF

Fungsi yang menghitung nilai setiap kata dengan TF (Term Frequency) dan IDF (Inverse Document Frequency) adalah bagian dari metode TF-IDF. Metode ini digunakan dalam berbagai aplikasi pemrosesan teks dan analisis sentimen.

3.4. Modeling Naive Bayes

Metode Naive Bayes digunakan untuk mengklasifikasikan ulasan bersifat komentar positif atau negatif. Pada penelitian ini menggunakan metode Naive Bayes untuk analisis sentimen, data ulasan dibagi menjadi dua kelas, yaitu kelas positif dan negatif, dan dilakukan klasifikasi untuk mengidentifikasi ulasan yang termasuk dalam masing-masing kelas. Metode Naive Bayes telah terbukti efektif dalam mengklasifikasikan sentimen berdasarkan ulasan pengguna, seperti ulasan film, ulasan aplikasi, dan komentar di media sosial seperti Twitter dan YouTube.

3.5. Evaluasi

Evaluasi model pembelajaran mesin bertujuan untuk mengetahui performa terbaik dari model yang dikembangkan. Hasil performa model disajikan dalam bentuk heatmap, tabel, dan langkah perhitungan matematis terkait berdasarkan pengujian model yang diimplementasikan dan dilatih. Performa model dievaluasi menggunakan matriks konfusi, yang memberikan informasi tentang akurasi prediksi, presisi, dan skor F1.

4. HASIL DAN PEMBAHASAN

Memuat hasil, Pengujian dan pembahasan tentang skripsi yang telah dilakukan.

4.1. Pengumpulan Data

Dalam tahap pengambilan data, peneliti menggunakan teknik web scraping dengan Google Colab dan library Python, `google_play_scraper`. Google Colab memungkinkan peneliti menjalankan kode Python melalui browser, sementara `google_play_scraper` memfasilitasi pengambilan data review dari Google Play Store. Peneliti menjalankan script Python di Google Colab untuk mengambil data review aplikasi Telegram. Fungsi dari `google_play_scraper` dipanggil dengan parameter tertentu, seperti identifikasi unik aplikasi, bahasa review, metode pengurutan review, dan jumlah review yang diambil. Hasilnya adalah data review dan `continuation_token`, yang berisi informasi seperti nama pengguna, waktu ulasan, skor, dan konten ulasan. Data review disimpan dalam `DataFrame`, struktur data Python yang memudahkan manipulasi data, diurutkan berdasarkan skor ulasan, dan disimpan sebagai file CSV untuk analisis lebih lanjut. Dengan demikian, peneliti dapat mengumpulkan data review dari Google Play Store dan menyimpannya dalam format yang mudah diakses dan dianalisis.

Jumlah Keseluruhan Data: 1000

	userName	at	score	
381	Olivia Sri Ananda	2023-11-05 08:01:07	5	
444	Nazwa Awa	2023-11-08 09:39:28	5	
851	Arhas Saheri	2023-10-27 15:02:03	5	
450	Nyna Septalia	2023-11-07 07:07:39	5	
452	Imang Rachman	2023-10-29 13:55:07	5	
453	Irfan Hakim	2023-11-04 04:15:42	5	
454	Sumani	2023-11-07 05:13:58	5	
455	Oki Diah Indah Lestari	2023-10-20 11:16:09	5	
456	Muhamad Jaenudin	2023-10-26 05:28:01	5	
457	kurdi seiha	2023-11-14 08:30:56	5	
content				
381	Aplikasi nya sangat bagus dan mudah untuk digu...			
444	27/11/20 Untuk telegram tolong di perbarui lag...			
851	Aplikasi luar biasa, dengan chat yang mudah di...			
450	Buat penggemar drakor. aplikasi ini sangat mem...			
452	Bisa buat nonton, nyimpen dokumen pribadi, tek...			
453	VERY GOOD banget apk nya simpel dan bikin nyam...			
454	Telegram aplikasi keren menyediakan anime-anim...			
455	VERY GOOD,, Berkat adanya telegram, saya bisa ...			
456	Aplikasi bagus digunakan bisa untuk menambah t...			
457	Telegram ini. aplikasi permesejan terpanas di...			

Gambar 2. Hasil Scrapping Data

4.2. Preprocessing

Tahap awal pada proses ini adalah melakukan pelabelan data secara manual. Proses ini memakan waktu tiga hari dua malam dan melibatkan pelabelan manual data ulasan oleh penulis. Kumpulan data ini dibagi menjadi kategori negatif dan positif. Materi informasi positif menunjukkan hal-hal positif tentang aplikasi Telegram, sedangkan materi negatif menunjukkan hal-hal negatif. Hasil dari proses ini dapat dilihat pada Tabel 1.

Table 1. Hasil Labeling

No	Content	Sentimen
1	Verry Good	Positif
2	Bagus Banget	Positif
3	Tidak bisa di simpan di galeri	Negatif
4	Udh susah dipakai lagi. ga bisa login	Negatif
5	aplikasi nya bagus	Positif

Selanjutnya melakukan beberapa tahapan berikut ini:

Berikut adalah perbaikan deskripsi dan penjelasan untuk setiap proses dalam text preprocessing:

a. Data Cleaning

Proses membersihkan data dengan menghilangkan karakter-karakter yang tidak diperlukan seperti tanda baca, angka, link, hashtag, mention, emoji, dan lain-lain. Tujuannya adalah untuk menghilangkan noise dan elemen yang tidak relevan dalam data teks. Dapat dilihat pada gambar 3.

id	content	score	sentiment
0	Apakah nge nge bug telponan sinyal bagus...	1	Negatif
1	Apakah nge nge bug telponan sinyal bagus...	1	Negatif
2	Apakah nge nge bug telponan sinyal bagus...	2	Positif
3	Apakah nge nge bug telponan sinyal bagus...	4	Positif
4	Apakah nge nge bug telponan sinyal bagus...	1	Negatif
5	Apakah nge nge bug telponan sinyal bagus...	1	Negatif
6	Apakah nge nge bug telponan sinyal bagus...	3	Positif
7	Apakah nge nge bug telponan sinyal bagus...	2	Positif
8	Apakah nge nge bug telponan sinyal bagus...	2	Positif
9	Apakah nge nge bug telponan sinyal bagus...	2	Positif

Gambar 3 Hasil Cleaning Data

b. Case Folding

Proses mengubah semua huruf dalam teks menjadi huruf kecil (lowercase) agar konsisten. Contoh: "Data SCIENCE" menjadi "data science". Tujuannya adalah agar tidak ada perbedaan antara kata yang ditulis dengan huruf besar dan kecil. Dapat dilihat pada Tabel 2.

Table 2. Case Folding

Sebelum	Sesudah
Aplikasi Suka Nge Bug Telponan sinyal Bagus Harap Diperbaiki	aplikasi suka nge bug telponan sinyal bagus harap diperbaiki

c. Stop Word Removal

Proses menghapus kata-kata umum yang sering muncul tetapi tidak memberikan banyak informasi seperti "yang", "di", "itu", dan lain-lain. Tujuannya adalah untuk meningkatkan akurasi analisis dengan menghilangkan kata yang dianggap noise.

Table 3. Stop Word Removal

Sebelum	Sesudah
aplikasi suka nge bug telponan sinyal bagus harap diperbaiki	aplikasi nge bug telponan sinyal harap perbaiki

d. Tokenizing

Proses memecah teks menjadi kata-kata individu atau token. Contoh: "Saya belajar data science" dipecah menjadi token ["Saya", "belajar", "data", "science"]. Tujuannya agar teks lebih mudah diproses pada tahap berikutnya.

Table 4. Tokenizing

Sebelum	Sesudah
aplikasi nge bug telponan sinyal harap perbaiki	[aplikasi, nge, bug, telponan, sinyal, harap, perbaiki]

e. Stemming

Proses mengubah kata-kata ke bentuk kata dasarnya dengan memotong imbuhan. Contoh: "Pembelajaran" diubah ke kata dasar: "Belajar". Tujuannya agar kata-kata yang semakna dianggap sama. Dapat dilihat pada Tabel 5.

Table 5. Stemming

Sebelum	Sesudah
[aplikasi, nge, bug, telponan, sinyal, harap, perbaiki]	[aplikasi, suka, nge, bug, telpon, sinyal, bagus, harap, baik]

f. Pembagian Data

Proses membagi data yang sudah dipreprocess menjadi 2 bagian: data latih (training) dan data uji (testing). Data latih digunakan untuk melatih model, dan data uji digunakan untuk

menguji performa model. Dengan membagi data latih dan data uji dengan perbandingan 80:20, diperoleh 768 data latih dan 192 data uji dari keseluruhan dataset. Tujuannya agar performa model bisa dievaluasi dengan data yang belum pernah dilihat sebelumnya.

4.3. TF-IDF

Proses TF-IDF (Term Frequency-Inverse Document Frequency) adalah pendekatan statistik yang digunakan untuk menentukan pentingnya kata dalam dokumen dalam kumpulan dokumen atau korpus. Proses ini melibatkan berbagai langkah kompleks yang bertujuan menganalisis data tekstual dan menghasilkan representasi numerik yang dapat digunakan untuk memahami dan mengekstrak fitur dari dokumen. Salah satu langkah utama dalam proses ini adalah mengimpor kelas `CountVectorizer` dari modul `sklearn`, yang digunakan untuk menghitung frekuensi kata dalam dokumen. Setelah itu, contoh kelas `CountVectorizer` dibuat dengan tujuan menerapkan metode yang sesuai ke objek. Metode ini digunakan untuk menghasilkan kosakata, yang merupakan kumpulan kata atau karakter unik yang kemudian digunakan untuk membangun vektor fitur. Selanjutnya, data teks untuk pelatihan dan pengujian diubah menjadi representasi numerik berdasarkan frekuensi kata menggunakan objek vektorisasi yang sama. Dengan cara ini, kumpulan data, baik untuk pelatihan maupun pengujian, dapat diubah menjadi representasi numerik berdasarkan frekuensi kata menggunakan kosakata yang dibangun dari data pelatihan. Proses TF-IDF sangat berguna dalam analisis teks karena memungkinkan ekstraksi fitur-fitur penting dan memahami pentingnya kata-kata dalam dokumen.

	apiknya	apl	apli	aplication	aplikanya	aplikasi	aplikasigak	\
0	0.0	0.0	0.0	0.0	0.0	0.464938	0.0	
1	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
2	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
3	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
4	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
...	...	...	...	...	...	...	...	
895	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
896	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
897	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
898	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	
899	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	

Gambar 4. Hasil TF-IDF

4.4. Naive Bayes

Multinomial Naive Bayes adalah pendekatan pembelajaran probabilistik yang umumnya digunakan untuk mengklasifikasikan teks atau data kategoris. Algoritma ini beroperasi dengan memanfaatkan teorema Bayes, prinsip dasar dalam teori probabilitas, yang memungkinkan perhitungan kemungkinan suatu peristiwa terjadi mengingat pengetahuan sebelumnya tentang kondisi yang terkait dengan peristiwa itu. Dalam bidang pemrosesan bahasa alami (NLP), Multinomial Naive Bayes berfungsi sebagai alat yang berharga untuk memprediksi tag atau kategori teks

tertentu berdasarkan fitur khas yang ada dalam teks tertentu. Dengan memanfaatkan fitur-fitur ini, Multinomial Naive Bayes memungkinkan klasifikasi data tekstual yang akurat dan efisien, sehingga memfasilitasi berbagai aplikasi di bidang NLP.

4.5. Evaluasi

Hasil evaluasi dari model klasifikasi menggunakan metode evaluasi confusion matrix disajikan dalam bentuk tabel yang berisi nilai akurasi, presisi, recall, dan f1-score. Confusion matrix adalah pengukuran performa untuk masalah klasifikasi machine learning yang menampilkan empat kombinasi berbeda dari nilai prediksi dan nilai aktual, yaitu True Positif (TP), False Positif (FP), False Negatif (FN), dan True Negatif (TN). Nilai-nilai ini digunakan untuk mengukur kinerja model klasifikasi. F1-score merupakan metrik yang menggabungkan akurasi dan recall, dan digunakan untuk mengevaluasi kinerja model klasifikasi. Proses evaluasi menggunakan confusion matrix berguna untuk mengukur sejauh mana model klasifikasi mampu memprediksi dengan benar dan mengidentifikasi kesalahan yang dilakukan oleh model. Berdasarkan pengujian model yang telah dilakukan, nilai confusion matrix hasil pengujian untuk dilakukan evaluasi dapat dilihat pada tabel 6.

Table 6. Confusion Matrix

Prediksi Aktual \	Negatif(0)	Positif (1)
Negatif (0)	47(TN)	7(FP)
Positif (1)	19(FN)	104(TP)

Beberapa informasi dapat disimpulkan dari Tabel 6: True Negative (TN): Data dengan judul negatif yang diprediksi dengan tepat sebagai negatif adalah 47. True Positive (TP): Data dengan tanda positif yang diprediksi dengan tepat sebagai tanda positif adalah 104. False negative (FN): data dengan tanda positif yang salah diprediksi sebagai tanda negatif, yaitu 19. False Positive (FP): data dengan tanda negatif yang salah diprediksi sebagai tanda positif yaitu 7.

Selanjutnya, nilai-nilai dari confusion matrix akan digunakan untuk menghitung kinerja model, yang mencakup akurasi, presisi, recall, dan skor F1. Proses perhitungannya adalah sebagai berikut:

$$1) \text{ Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \times 100\%$$

$$\text{Accuracy} = \frac{104 + 47}{104 + 47 + 7 + 19} \times 100\%$$

$$\text{Accuracy} = \frac{151}{177} \times 100\% = 85.31\%$$

$$2) \text{ Presisi} = \frac{TP}{TP+FP}$$

$$\text{Presisi} = \frac{104}{104+7} = 0.93$$

$$3) \text{ Recall} = \frac{TP}{TP+FN}$$

$$\text{Recall} = \frac{104}{104 + 19} = 0.84$$

$$4) \text{ F1 - Score} = 2x \frac{\text{Recall} \times \text{Presisi}}{\text{Recall} + \text{Presisi}}$$

$$\text{F1 - Score} = 2x \frac{0.84 \times 0.93}{0.84 + 0.93} = \frac{1.5624}{1.77} = 0.88$$

Lebih mudahnya dapat dilihat pada Tabel 7.

Table 7. Hasil Evaluasi

Matrix	Perbandingan 80:20
Akurasi (%)	85.31%
Presisi (%)	93%
Recall (%)	76%
F1-Score (%)	88%

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian, pengguna aplikasi Telegram umumnya memberikan ulasan positif terkait kemudahan penggunaan, keuntungan, layanan pelanggan yang responsif, dan aspek keamanan, meskipun ada juga ulasan negatif yang mencakup kesulitan penggunaan, kerugian ekonomis, layanan pelanggan yang kurang responsif, dan ketidaknyamanan dalam aspek keamanan. Analisis sentimen menggunakan algoritma Naive Bayes menunjukkan tingkat akurasi yang tinggi sebesar 85.31%, dengan rasio pembagian data untuk evaluasi adalah 80:20, nilai presisi 93%, recall 76%, dan F1-Score 88%, menunjukkan efektivitas model dalam mengklasifikasikan sentimen ulasan pengguna. Meskipun terdapat beberapa kelemahan, kelebihan aplikasi lebih dominan, sehingga disarankan agar pihak Telegram terus meningkatkan kualitas aplikasinya dengan mempermudah penggunaan, menambahkan fitur baru, meningkatkan kualitas sistem, menangani gangguan layanan secara efisien, serta memperbaiki masalah login yang sering dialami pengguna untuk meningkatkan kepuasan pengguna dan daya saingnya di pasar.

DAFTAR PUSTAKA

- [1] D. D. Putri, G. F. Nama, And W. E. Sulistiono, "Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (Dpr) Pada Twitter Menggunakan Metode Naive Bayes Classifier," *J. Inform. Dan Tek. Elektro Terap.*, Vol. 10, No. 1, Pp. 34–40, 2022.
- [2] A. Syafianto, "Perbandingan Algoritma Naive Bayes Dan Decision Tree Pada Sentimen Analisis," *Indones. J. Comput. Sci. Res.*, Vol. 1, Pp. 1–15, 2022, [Online]. Available: <https://Subset.Id/Index.Php/Ijcsr/Article/View/11%0ahttps://Subset.Id/Index.Php/Ijcsr/Article/Download/11/4>
- [3] A. I. Tanggraeni And M. N. N. Sitokdana, "Analisis Sentimen Aplikasi E-Government Pada Google Play Menggunakan Algoritma Naive Bayes," *Jatisi (Jurnal Tek. Inform. Dan Sist. Informasi)*, Vol. 9, No. 2, Pp. 785–795, 2022, Doi: 10.35957/Jatisi.V9i2.1835.
- [4] A. A. Permana, W. A. Noviyanto, And D. A. Kristiyanti, "Sentimen Analisis Opini Masyarakat Terhadap Umkm Pada Media Sosial Twitter Dengan Metode Naive Bayes Classifier," *J. Minfo Polgan*, Vol. 12, No. 1, Pp. 163–170, 2023, Doi: 10.33395/Jmp.V12i1.12337.
- [5] N. Andre Saputra, J. Alexandra, And I. Budi Trisno, "Analisis Sentimen Pemanfaatan Obrolan Grup Telegram Berbagi Informasi Lowongan Kerja Menggunakan Metode Naive Bayes Classifier," *Jati (Jurnal Mhs. Tek. Inform.)*, Vol. 7, No. 2, Pp. 1321–1327, 2023, Doi: 10.36040/Jati.V7i2.6693.
- [6] N. Agustina, D. H. Citra, W. Purnama, C. Nisa, And A. R. Kurnia, "Implementasi Algoritma Naive Bayes Untuk Analisis Sentimen Ulasan Shopee Pada Google Play Store," *Malcom Indones. J. Mach. Learn. Comput. Sci.*, Vol. 2, No. 1, Pp. 47–54, 2022, Doi: 10.57152/Malcom.V2i1.195.
- [7] Fitri Wulandari, Elin Haerani, Muhammad Fikry, And Elvia Budianita, "Analisis Sentimen Larangan Penggunaan Obat Sirup Menggunakan Algoritma Naive Bayes Classifier," *J. Coscitech (Computer Sci. Inf. Technol.)*, Vol. 4, No. 1, Pp. 88–96, 2023, Doi: 10.37859/Coscitech.V4i1.4781.
- [8] D. Lihardo Girsang, A. Sidiq, And T. Salsabila Elenaputri, "Analisis Sentimen Masyarakat Terhadap Layanan Bpjs Kesehatan Dan Faktor-Faktor Pendukung Opini Dengan Pemodelan Natural Language Processing (Nlp)," *Emerg. Stat. Data Sci. J.*, Vol. 1, No. 2, Pp. 238–249, 2023.
- [9] D. M. Efendi, S. Mintoro, . S., S. H. Lubis, And S. Lestari, "Klasifikasi Kinerja Pembayaran Angsuran Dengan Algoritma Naive Bayes (Studi Kasus : Data Nasabah Koperasi Simpan Pinjam Pembiayaan Syariah Bina Bersama)," *J. Inf. Dan Komput.*, Vol. 10, No. 1, Pp. 57–61, 2022, Doi: 10.35959/Jik.V10i1.305.
- [10] O. I. Gifari, M. Adha, F. Freddy, And F. F. S. Durrand, "Analisis Sentimen Review Film Menggunakan Tf-Idf Dan Support Vector Machine," *J. Inf. Technol.*, Vol. 2, No. 1, Pp. 36–40, 2022, Doi: 10.46229/Jifotech.V2i1.330.
- [11] M. K. Suryadewiansyah And T. E. E. Tju, "Naive Bayes Dan Confusion Matrix Untuk Efisiensi Analisa Intrusion Detection System Alert," *J. Nas. Teknol. Dan Sist. Inf.*, Vol. 8, No. 2, Pp. 81–88, 2022, Doi: 10.25077/Teknosi.V8i2.2022.81-88.