

IMPLEMENTASI ALGORITMA NAIVE BAYES UNTUK ANALISIS SENTIMEN PENGGUNA PADA APLIKASI MOBILE JKN

¹⁾Ikhwan, ²⁾Merlin Diana, ³⁾Bagus Setya

^{1,2,3)} Program Studi Sistem Informasi, Fakultas Ilmu Komputer, Universitas Pamulang

¹⁾dosen03358@unpam.ac.id, ²⁾dosen03044@unpam.ac.id, ³⁾dosen03357@unpam.ac.id

INFO ARTIKEL

Riwayat Artikel :

Diterima : 15 Juni 2026

Disetujui : 10 Agustus 2026

Kata Kunci :

Analisis Sentimen, *Multinomial Naive Bayes*, Ulasan Pengguna, Mobile JKN, *Confusion Matrix*.

ABSTRAK

Mobile JKN merupakan aplikasi layanan kesehatan untuk masyarakat berbasis perangkat mobile yang diakses secara *online*. Tingginya pemanfaatan layanan digital tersebut terlihat dari *Google Play Store* yang diakses pada Mei 2026 dengan jumlah unduhan 50 juta lebih, total 4,4 rating dari 5, dan 1,02 juta lebih ulasan. Meskipun partisipasi masyarakat cukup tinggi, masih terdapat keluhan dari masyarakat, baik terkait fitur aplikasi maupun sistem pelayanannya. Tujuan penelitian ialah untuk menganalisis masalah yang dialami *user*, mengklasifikasi sentimen ulasan, serta mengevaluasi performa algoritma. Algoritma klasifikasi yang digunakan yaitu Naive Bayes. Data diambil dengan *scraping* ulasan terbaru di *Google Play Store*. Diperoleh data sebanyak 10.000 ulasan dari rentang Maret sampai Mei 2026. Pemrosesan data meliputi pembersihan, *case folding*, *stopword removal*, tokenisasi, dan *stemming*, sehingga menghasilkan 9.938 data. Pelabelan data menghasilkan ulasan positif sebanyak 6.021, ulasan negatif 3.610, dan ulasan netral 307. Proses klasifikasi hanya menggunakan kelas positif dan negatif, sedangkan netral tidak disertakan. Pembobotan kata menggunakan TF-IDF, dan evaluasi menggunakan *confusion matrix* akurasi, *recall*, presisi, F1-Score serta visualisasi data *word cloud*. Analisis terhadap keluhan yang paling dominan terdapat pada proses *login* dan autentikasi pengguna, seperti gagal *login*, masalah kode OTP, dan gagal verifikasi akun, serta keluhan terkait *error* aplikasi. Performa yang dihasilkan sudah cukup baik, pada rasio 80:20 pembagian data training dan testing didapat hasil akurasi 92,84%, presisi 96,84%, *recall* 91,55%, dan F1-score 94,12%. Harapan penelitian yaitu dapat menjadi bahan evaluasi teruntuk pembuat aplikasi agar lebih berkualitas dan meningkatkan mutu pelayanan kepada masyarakat.

ARTICLE INFO

Article History :

Received : Jun 15, 2026

Accepted : Aug 10, 2026

Keywords:

Sentiment Analysis, *Multinomial Naive Bayes*, User Reviews, Mobile JKN, *Confusion Matrix*.

ABSTRACT

Mobile JKN is a mobile-based healthcare service application for the public that is accessed online. The high level of usage for the digital service is evident from its *Google Play Store* listing as of May 2026, showing over 50 million downloads, a 4.4 out of 5 rating, and more than 1.02 million reviews. Despite the relatively high level of public participation, there are still complaints from the public regarding both the application's features and the service system. The objective of the research is to analyze the problems experienced by users, classify review sentiments, and evaluate algorithm performance. The classification algorithm used is Naive Bayes. Data was collected by *scraping* recent reviews from the *Google Play Store*, resulting in a dataset of 10,000 reviews spanning the period from March to May 2026.

Data preprocessing included data cleaning, case folding, stopword removal, tokenization, and stemming, resulting in 9,938 data points. Data labeling yielded 6,021 positive reviews, 3,610 negative reviews, and 307 neutral reviews. The classification process uses only positive and negative classes, while the neutral class is excluded. Word weighting was performed using TF-IDF; evaluation utilized a confusion matrix—covering accuracy, recall, precision, and F1-score—alongside word cloud data visualization. An analysis of the most prevalent complaints reveals issues centered on the user login and authentication processes—such as login failures, OTP code problems, and account verification failures—as well as complaints regarding application errors. The resulting performance is quite good; with an 80:20 split between training and testing data, the results achieved were 92.84% accuracy, 96.84% precision, 91.55% recall, and a 94.12% F1-score. It is hoped that this research will serve as a basis for evaluation for application developers, enabling them to enhance the quality of the application and improve service delivery to the public.

1. PENDAHULUAN

Perkembangan teknologi informasi di era digitalisasi sekarang ini telah memberikan perubahan dan dampak yang cukup signifikan di berbagai bidang. Salah satunya adalah pemanfaatan teknologi pada bidang kesehatan berbasis perangkat mobile yang dapat diakses secara daring. Mobile JKN ialah satu dari aplikasi milik pemerintah yang digunakan dalam memudahkan layanan kesehatan pada masyarakat. Aplikasi Mobile JKN dapat di *download* pada *Google Play Store*. Menu yang ada di antaranya layanan pendaftaran peserta, perubahan data peserta, cek iuran, antrean *online* dan layanan kesehatan lainnya.

Tingginya pemanfaatan layanan digital tersebut terlihat dari *Google Play Store* yang diakses pada Mei 2026 dengan jumlah terunduh 50 juta lebih, 4,4 rating dari 5, dan 1,02 juta lebih ulasan. Meskipun partisipasi masyarakat cukup tinggi, masih terdapat juga keluhan dari masyarakat, baik dari fitur aplikasi maupun dari sistem pelayanannya.

Berdasarkan hasil studi literatur serta observasi pada ulasan pengguna mobile JKN, masih ditemukan keluhan seperti proses autentikasi, kegagalan login dan verifikasi akun. Keluhan yang paling dominan yaitu terdapat pada proses *login*, autentikasi pengguna, gagal *login*, masalah kode OTP, gagal verifikasi akun, serta keluhan terkait *error* aplikasi. Bahkan ada juga pengguna yang memberikan nilai rating positif, akan tetapi ulasan yang diberikan merujuk pada arah negatif. Dengan kondisi tersebut, dirasa akan terasa sulit untuk menentukan penilaian kepuasan pengguna. Untuk itu,

perlu dilakukan analisis sentimen, klasifikasi, dan evaluasi model menggunakan *confusion matrix*.

Analisis sentimen merupakan adalah dalam memahami data, ekstrak data, dan mengolah data untuk mendapatkan informasi ulasan secara otomatis, di mana tujuannya ialah untuk dapat mengetahui kepuasan *user* (Maulana et al., 2024). Analisis sentimen menjadi suatu hal yang sangat penting guna mencari tahu dan memahami pendapat masyarakat terkait aplikasi atau layanan serta topik tertentu (Fatra et al., 2025).

Terdapat algoritma klasifikasi yang dapat digunakan untuk analisis sentimen, misalnya Naive Bayes, *support vector machine* (SVM), dan K-Nearest Neighbor (KNN) atau algoritma yang lainnya. Namun, algoritma yang dipilih adalah *Multinomial* Naive Bayes serta TF-IDF digunakan sebagai pembobotan katanya. Evaluasi performa dengan akurasi, presisi, *recall*, dan *F1-score*. Pemilihan Naive Bayes dikarenakan proses klasifikasi dan kesederhanaannya dalam proses perhitungan, serta kemampuannya menangani data teks dalam jumlah yang cukup besar. Metode ini menggunakan teknik *teks mining* yang memiliki tahapan seperti *data scraping*, *cleaning*, *preprocessing*, dan *labelling*. Pada *preprocessing* di dalamnya terdapat proses seperti *transformation*, *tokenization* dan *filtering* (Darmawan et al., 2023). Implementasi algoritma ini dapat digunakan dengan bahasa pemrograman Python dan berbagai library di dalamnya (Ghozali et al., 2023; Ramadhani et al., 2024).

Selain itu, beberapa artikel menyatakan bahwa algoritma Naive Bayes mendapatkan akurasi yang

cukup tinggi dan juga terbilang cepat dalam melakukan analisis sentimen dan tidak membutuhkan komputasi tinggi, dibandingkan dengan algoritma klasifikasi lain.

Penelitian terdahulu telah dilakukan pada sistem *digital Cortax* dengan menggunakan Naive Bayes. Data ulasan diambil dari sistem X dengan 1.009 data *tweet*. Setelah *data cleaning*, dihasilkan 899 data. Pelabelan dibagi menjadi kelas positif, negatif dan netral. Beberapa tahapan dilakukan, mulai dari *crawling*, pemrosesan data meliputi tokenisasi, *stopword removal*, dan *stemming*, serta TF-IDF. Akurasi yang didapat 77% dengan performa yang baik terdapat pada sentimen negatif dan netral sedangkan dengan sentiment positif masih kurang. Ini disebabkan karena adanya perbedaan jumlah data yang mana kelas positif lebih sedikit dibanding dengan kategori lainnya, sehingga model tidak belajar secara maksimal dalam mendeteksi sentimen kelas positif (Fathoni et al., 2025).

Penelitian oleh Yanah et al. (2025) dilakukan pada aplikasi Komisi Pemilihan Umum (KPU) pada bulan November 2024. Aplikasi menggunakan Naive Bayes *algorithm*, dengan tujuan untuk analisis data ulasan atau dari *user*. *Dataset* didapat dari *Google Play Store* sebanyak 200 ulasan. Pelabelan dilakukan manual, yaitu sikap positif, negatif dan netral. Pada pembagian data, terdapat 40 data *testing* dan 160 data *training*. Pemrosesan data dengan *case folding*, *filtering*, *stemming* dan tokenisasi. Hasil akurasi sentimen positif 82,5% dan 75% untuk sentimen negatif. Penelitian ini memiliki kekhawatiran terhadap kegagalan teknologi dan keamanan data.

Penelitian yang dilakukan oleh Mandasari dan Faizin (2025) Membahas terkait dengan analisis sentimen dengan objek penelitiannya yaitu aplikasi DANA di *Google Play Store*. Tujuannya adalah untuk analisis sentimen menggunakan algoritma Naive Bayes dan menguji efektivitasnya. Total akurasi didapat 88% dengan hasil pada kategori positif yaitu presisi 0,92 dengan 0,94 *recall*, serta F1-score 0,93, pada kelas negatif yaitu presisi 0,76, recall 0,83, dan F1-score 0,79. Pendistribusian data yang dihasilkan ialah sebanyak 72,2% kelas positif dan 22,9% pada kelas negatif, dan kelas netral hanya 4,4%. Pada kelas positif dan negatif algoritma ini bisa dikatakan efektif, akan tetapi pada kelas netral masih dikatakan kurang, ini dikarenakan adanya ketidakseimbangan data.

Penelitian oleh Bunawan (2026) Bertujuan untuk mengklasifikasikan emosi publik di media sosial terkait wilayah Palembang. Kategori menggunakan kelas,

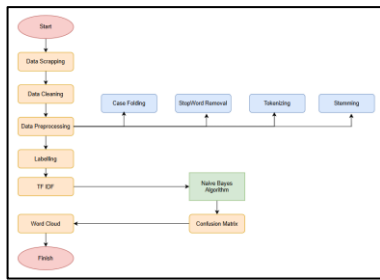
marah, senang, sedih, jijik, takut, terkejut dan netral. *Dataset* diperoleh dengan *scraping* sebanyak 9.125 komentar. Pelabelan menggunakan metode *rule-based* leksikon untuk mempercepat anotasi, walaupun ada keterbatasan dalam nuansa konteks dalam bahasa yang kompleks. Ekstrak data menggunakan TF-IDF yang digunakan pada *Support Vector Machine* atau Naive Bayes. Akurasi yang didapat sebesar 65% dan persentase F1-score sebanyak 58%.

Hasil analisis sentimen ulasan sudah cukup banyak dilakukan pada algoritma klasifikasi, termasuk Naive Bayes, *Support Vector Machine (SVM)* atau yang lainnya. Sebagian besar fokus perbandingan algoritma klasifikasi, dan performa dari nilai akurasi, presisi, *recall* dan F1-score. Sedangkan pembahasan terkait pola keluhan pengguna masih terbatas, terutama dalam melihat kemampuan model mengenali sentimen negatif yang berhubungan dengan keluhan layanan, kendala sistem, dan pengalaman *user*.

Tujuan penelitian ini yaitu untuk analisis pada aplikasi Mobile JKN yang terdapat pada ulasan *Google Play Store* terhadap ulasannya. Evaluasi dilakukan pada ulasan terkait dengan aplikasi maupun pelayanannya. Klasifikasi sentimen ulasan, evaluasi performa algoritma *Multinomial* Naive Bayes dengan akurasi, presisi, *recall*, dan F1-Score, dan juga evaluasi menggunakan *confusion matrix*, serta visualisasi *word cloud* untuk menentukan sentimen positif dan negatif yang dominan. Penggunaan data terbaru pada *Google Play Store* merupakan kontribusi pada penelitian ini. Selain itu, bukan hanya terfokus pada hasil evaluasi model dan akurasi, tetapi juga pada analisis jenis keluhan pengguna yang paling dominan sebagai bahan evaluasi untuk kualitas pelayanan. Dengan demikian, penelitian ini diharapkan membantu dalam penentuan perbaikan dan evaluasi bagi pengembang aplikasi dan pelayanan kesehatan publik.

2. METODE

Beberapa tahapan dilakukan, di antaranya *data scraping*, *cleaning*, *preprocessing*, *labeling*, pembobotan kata dengan TF-IDF, pemodelan dengan Naive Bayes dan evaluasi menggunakan *confusion matrix* serta visualisasi *word cloud*. Untuk tahapan pada penelitian dapat dilihat pada gambar 1 berikut ini.



Gambar 1. Tahapan atau langkah penelitian

2.1. Data Scraping

Teknik web *scraping* ulasan dilakukan pada web *Google Play Store* dengan memanfaatkan *Google Colaboratory*. *Scraping* dilakukan pada tanggal 18 Mei 2026 dengan rentang data bulan Maret hingga Mei 2026 menggunakan *library* dari *google-play-scraper* versi 1.2.7 dan *sastrawi* pada versi 1.0.1. Selain itu, proses tersebut dengan pemrograman Python versi 3.12.13 dan *scikit-learn* versi 1.6.1.

Teknik ini dilakukan guna mendukung efisiensi waktu pada saat pengambilan data dengan jumlah data yang banyak, sehingga memungkinkan pengambilan data lebih cepat. Diperoleh data ulasan terbaru yang berasal dari wilayah negara Indonesia, dengan jumlah data yang diambil sebanyak 10.000.

Parameter yang digunakan saat *scraping* di antaranya parameter *app.bpjs.mobile* yang dipakai untuk mengetahui identitas aplikasi sebagai bagian dari objek *scraping*, *lang='id'* yaitu parameter yang berfungsi untuk mengambil ulasan dengan bahasa Indonesia, parameter *country='id'* digunakan untuk mengambil data ulasan untuk wilayah Indonesia, serta parameter *sort=Sort.NEWEST* digunakan untuk *sorting* atau mengurutkan data hasil *scraping* berdasarkan data terbaru. Serta, parameter *count=10000* digunakan untuk menentukan jumlah maksimum ulasan yang akan diambil.

2.2. Data Cleaning

Pembersihan data adalah proses yang dilakukan untuk menghapus data kosong (*missing value*), data duplikat, karakter yang tidak sesuai, dan ulasan yang tidak memenuhi kriteria dalam analisis.

Komentar pengguna yang didapat dari web *Google Play Store* ada kemungkinan masih berisi tanda baca, angka, simbol, link, dan yang

lainnya, sehingga perlu dibersihkan. Kondisi tersebut bisa menyebabkan masalah saat analisis data, sehingga data menjadi tidak berkualitas. Setelah *data cleaning*, dari total 10.000 data ulasan, dihasilkan data sebanyak 9.938.

2.3. Data Preprocessing

Langkah ini bertujuan untuk mengurangi adanya ketidaksesuaian data yang menyebabkan terganggunya perhitungan saat analisis dan klasifikasi data. *Preprocessing* dilakukan dengan teknik *case folding*, *stopword removal*, *tokenizing* dan *stemming* (Hayatin et al., 2020). Bagan berikut menunjukkan Tahapan *Data Preprocessing*.



Gambar 2. Data Preprocessing

Case folding merupakan proses untuk menjadikan format teks menjadi seragam. Misal mengubah huruf pada teks agar menjadi format yang seragam contohnya huruf kapital diubah menjadi huruf kecil. Contohnya kata Mobile, MOBILE, dan mobile akan diubah menjadi kata mobile.

Stopword removal berfungsi untuk menghapus kata-kata umum yang dirasa tidak memiliki makna yang dianggap penting. Proses ini bertujuan untuk menyederhanakan teks, sehingga proses analisis menjadi lebih efektif dan efisien. Dengan demikian, algoritma dapat bekerja lebih fokus pada kata yang memiliki pengaruh pada sentimen. Kata pada *stopword* antara lain “ke”, “yang”, “dan”, “di”, “itu”, “adalah”, serta kata lainnya yang sering digunakan dalam kalimat sehari-hari dan bersifat umum.

Tokenizing yaitu proses menyederhanakan teks dari pola teks panjang menjadi bagian-bagian teks yang lebih kecil. Setiap kalimat ulasan akan dipecah menjadi kumpulan kata, sehingga mempermudah proses analisis sentimen dan pembobotan kata. Misalnya, kalimat “platform aplikasi mobile jkn sangat bagus” akan diubah menjadi token-token ([“platform”, “aplikasi”, “mobile”, “jkn”, “sangat”, “bagus”]).

Stemming merupakan proses yang dilakukan dengan cara mengganti kata yang berimbuhan kedalam bentuk kata-kata dasar. Misalnya, kata (membantu, dibantu, dan bantuan) diubah jadi kata “bantu”.

2.4. Data Labeling

Labeling adalah tahapan memberikan kategori tertentu pada data, yang tujuannya adalah untuk memberikan informasi jenis sentimen yang terkandung dalam setiap data ulasan. Pemberian label dibentuk dari kombinasi peringkat atau rating dengan skala 1 sampai 5 (Jelita, Saad & Wahyuni 2025).

Proses ini dilakukan dengan cara membagi kelompok data ke beberapa kategori, contoh kategori positif, negatif, dan netral yang berdasarkan nilai *rating*. Penelitian ini kelas positif dan negatif yang digunakan, sedangkan kelas netral tidak digunakan.

2.4. TF-IDF

TF-IDF atau *Term Frequency-Inverse Document Frequency* merupakan tahapan yang digunakan dalam pemberian bobot pada kata atau proses untuk mengubah data menjadi bentuk angka, sehingga dapat diproses oleh algoritma klasifikasi. Proses ini bertujuan untuk mengetahui tingkat kepentingan pada kata dalam dokumen berdasarkan frekuensi kemunculannya. Akan memiliki bobot lebih tinggi dan dianggap lebih penting jika kata yang muncul pada suatu dokumen dan jarang muncul pada dokumen lain.

Proses ini dilakukan dengan *TfidfVectorizer* pada *library scikit-learn*. TF-IDF menggunakan parameter *max_features=5000*, sehingga membatasi fitur sebanyak 5000 kata dengan bobot yang sesuai, *ngram_range=(1,1)*, yaitu fitur *unigram* yang digunakan, *min_df=2*, yaitu kata yang muncul minimal pada dua dokumen yang dipertahankan. Selain itu, jika yang muncul hanya satu dokumen, maka kata akan diabaikan. Pemberian bobot (TF-IDF) pada kata berdasarkan frekuensi pada suatu dokumen (*Term Frequency*). Dan juga seberapa jarang kata akan muncul dalam dokumen (*Inverse Document Frequency*). Bobot tinggi akan diberikan pada kata yang sering muncul pada suatu dokumen (Punusingun et al., 2025).

$$TF(t, d) = \frac{\text{Jumlah kemunculan kata}}{\text{total kata}} \quad (1)$$

$$IDF(t, D) = \log\left(\frac{N}{N(t)}\right) \quad (2)$$

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

N : total jumlah pada kumpulan dokumen.
N(t) : jumlah dokumen yang berisi kata t.

TF(t,d) : frekuensi kata t pada dokumen d.
IDF(t,D) : *invers* frekuensi dokumen untuk kata pada kumpulan dokumen

2.5. Naive Bayes Algorithm

Algoritma kalsifikasi yang digunakan yaitu *Multinomial Naive Bayes*. Algoritma ini digunakan pada klasifikasi pemodelan analisis sentimen dan mampu menghasilkan akurasi yang cukup baik pada *dataset* yang digunakan (Somantri et al., 2024). Algoritma dipakai untuk klasifikasi data pada kategori tertentu dengan menghitung kemungkinan kemunculan suatu kelas. Naive Bayes mengklasifikasikan sentimen ulasan pengguna ke dalam kategori positif dan negatif. Klasifikasi dilakukan berdasarkan hasil TF-IDF yang dilakukan pada tahap sebelumnya. *Dataset* pada penelitian akan dibagi menjadi dua bagian, yaitu data uji dan data latih, rasio perbandingan yaitu 80:20 yang artinya 80% data training dan sisanya data uji (Enjeli dan Wijaya, 2024).

2.6. Confusion Matrix

Tahapan ini dilakukan guna evaluasi dan mengukur performa algoritma atau model klasifikasi, akurasi, dan presisi (Arminda et al., 2023). Metode ini digunakan untuk mengetahui performa dari pemodelan klasifikasi sentimen ulasan pengguna aplikasi Mobile JKN dalam kategori positif ataupun negatif.

Selanjutnya, model akan diukur menggunakan tabel *Confusion Matrix* (Chaithra, 2019). Matriks konfusi adalah tabel yang berisi banyaknya baris data uji berdasarkan prediksi benar atau salah menurut model klasifikasi (Amanda dan Negara, 2020).

Pada *confusion matrix* terdapat 4 komponen, diantaranya *True Positive*, *True Negative*, *False Positive*, dan *False Negative*.

True Positive ialah jumlah data ulasan positif yang berhasil diprediksi dengan benar sebagai positif, *True Negative* ialah data ulasan negatif yang berhasil diprediksi dengan benar sebagai negatif, *False Positive* merupakan ulasan negatif yang diprediksi sebagai positif, dan *False Negative* merupakan data ulasan positif yang diprediksi sebagai negatif. Metrik pengukuran yang digunakan pada evaluasi model algoritma yaitu *accuracy*, *precision*, *recall*, dan *F1-score*.

Accuracy secara keseluruhan digunakan untuk mengukur tingkat tepatnya suatu model dalam melakukan klasifikasi. Seberapa banyak data yang berhasil diklasifikasikan dengan benar oleh model dibandingkan dengan total jumlah data yang diuji. Semakin tinggi nilai akurasi, semakin baik performa model yang dihasilkan. Persamaan *accuracy* dapat dilihat sebagai berikut:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (4)$$

Presisi untuk menghitung tingkat tepatnya suatu prediksi pada ulasan positif. Semakin tinggi nilainya, semakin baik model tersebut menghindari kesalahan prediksi pada kategori positif. Persamaan *precision* dapat dilihat sebagai berikut:

$$Precision = \frac{TP}{TP+FP} \quad (5)$$

Recall dilakukan untuk mengukur kemampuan algoritma agar dapat menemukan semua data pada setiap kategori. *Recall* dikenal sebagai *True Positive Rate*, yaitu mengukur sampai sejauh mana algoritma dapat mengidentifikasi semua kriteria positif. Persamaan *recall* dapat dilihat pada bagan berikut:

$$Recall = \frac{TP}{TP+FN} \quad (6)$$

F1-score berfungsi untuk menyeimbangkan atau menggabungkan nilai presisi dan *recall* menjadi satu nilai tunggal. Sehingga menghasilkan evaluasi model yang lebih optimal. Persamaan *F1-Score* dapat dilihat pada bagan berikut ini:

$$F1 \text{ Score} = 2 \frac{Presisi \times recall}{Presisi+Recall} \quad (7)$$

Keterangan :

TP : nilai sebenarnya positif dan juga nilai prediksi positif.

TN : nilai sebenarnya negative dan juga nilai prediksi negatif.

FP : nilai sebenarnya negatif dan juga nilai prediksi positif

FN : nilai sebenarnya positif dan juga nilai prediksi negatif.

2.7. Word Cloud

Tahap ini menggunakan visualisasi dalam penggambaran dominan baik ulasan positif dan negatif dari seluruh ulasan. *Word cloud* menunjukkan data dominan pada ulasan pengguna berdasarkan kelas sentimen.

3. HASIL DAN PEMBAHASAN

3.1. Dataset penelitian

Scraping di *Google Play Store* menggunakan bahasa pemrograman Python. Gambar berikut adalah kode program untuk melakukan *scraping*.

```
# Install library google-play-scraper
!pip install google-play-scraper

# Import library
from google_play_scraper import reviews, sort
import pandas as pd

# Scraping ulasan aplikasi Mobile MA
result = reviews(app_id='com.spotify.music', sort='newest')

# Data ulasan Mobile MA
data = pd.DataFrame(result)

# Mengambil kolom penting
df = df[['name', 'score', 'text', 'timestamp']]

# Mengganti nama kolom
df.columns = ['nama_app', 'rating', 'ulasan', 'waktu']

# Menampilkan 5 data pertama
print(df.head())

# Menyimpan hasil ke file CSV
df.to_csv('ulasan_mobile_ma.csv', index=False)
```

Gambar 3. *Scraping* dengan python

Pada *data scraping* yang didapat akan dilakukan pemilihan kolom. Pemilihan bertujuan untuk menyaring data yang memang benar-benar digunakan. Kolom yang digunakan pada penelitian ini yaitu kolom ulasan dan *rating*.

	B	D
1	Rating	Ulasan
2	5	simple, praktis dan sangat membantu
3	5	mudah2an bermanfaat bagi seluruh masyarakat Indonesia... amiiii...
4	3	parah juga ini aplikasi. tidak ada fitur pembayaran lunas yang bisa dilakukan secara mandiri/virtual
5	5	sangat membantu
6	5	Luar biasa
7		Sampah, tiba tiba log out gilirin login malah kaga bisa. Kalo mau ada perbaikan minimal ada info atau apa
8	1	kek. Orang mau berobat jadi kesusahannya. Jangn karena bps gua yang dari pemerintah!
9	5	semuanya jadi mudah, asal peringatnya mantap
10	5	good job

Gambar 4. Dataset hasil pemilihan kolom

Dari Gambar 4 di atas, kolom ulasan dan kolom rating akan dilanjutkan ke proses selanjutnya.

3.2. Pembersihan Data

Proses ini dilakukan untuk menghapus karakter yang sifatnya khusus, ketidaksesuaian teks dan link yang tidak dibutuhkan.

```

# Proses Data Cleaning
def cleaning_text(text):
    # Hapus karakter string
    text = str(text)

    # Hapus URL
    text = re.sub('https?://www.?https?', '', text)

    # Hapus mention dan hashtag
    text = re.sub('@\w+', '', text)

    # Hapus angka
    text = re.sub('\d+', '', text)

    # Hapus karakter khusus dan tanda baca
    text = re.sub('[^a-zA-Z]', '', text)

    # Hapus spasi berlebih
    text = re.sub(' ', ' ', text).strip()

    return text

# Bersihkan cleaning pada file ulasan
df['ulasan'] = df['ulasan'].apply(cleaning_text)

# Hapus data kosong
df = df[df['ulasan'] != '']

df.to_csv('3_cleaning_ulasan_mobile_30.csv', index=False)
    
```

Gambar 5 data cleaning

Proses pembersihan data menggunakan fungsi *def cleaning_text(text)*. Dengan parameter yang dibersihkan: URL, mention, numerik, karakter dan spasi. Setelah proses pembersihan data, jumlah data ulasan yang awalnya 10.000 data, berkurang menjadi 9.938 data.

```

Data Sebelum Cleaning
0      simple, praktis dan sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah juga ini aplikasi tidak ada fitur pushnot...  3
3      sangat membantu  5
4      luar biasa  5
...
95     kenapa sekarang perpanjang rujukan pasien hrs ...  1
96     koplollllll nunggu verifikasi nomor hp mesti se...  1
97     gaje banget mau daftar aja dipersulit gini.kah...  1
98     lagi lihr kah . koo gangguan mlu mau login ...  3
99     sangat membantu  5

[100 rows x 2 columns]

Data Setelah Cleaning
0      simple praktis dan sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah juga ini aplikasi tidak ada fitur pushnot...  3
3      sangat membantu  5
4      luar biasa  5
...
96     koplollllll nunggu verifikasi nomor hp mesti se...  1
97     gaje banget mau daftar aja dipersulit gini.kah...  1
98     lagi lihr kah . koo gangguan mlu mau login ...  3
99     sangat membantu  5
100    Verifikasi wajah gagal terus tanpa ada alasan ...  5
    
```

Gambar 6. Sebelum dan sesudah pembersihan

Hasil tersebut menunjukkan bahwa proses ini diharapkan dapat membuat data yang digunakan menjadi lebih akurat, bebas dari kesalahan atau ketidaksesuaian data serta data menjadi lebih konsisten.

3.3. Pemrosesan Data

3.3.1 Case Folding

Proses ini digunakan untuk mengganti huruf kapital menjadi huruf kecil.

```

Sebelum Case Folding
0      simple praktis dan sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah juga ini aplikasi tidak ada fitur pushnot...  3
3      sangat membantu  5
4      luar biasa  5
...
95     kenapa sekarang perpanjang rujukan pasien hrs ...  1
96     koplollllll nunggu verifikasi nomor hp mesti se...  1
97     gaje banget mau daftar aja dipersulit gini.kah...  1
98     lagi lihr kah . koo gangguan mlu mau login ...  3
99     sangat membantu  5

[100 rows x 2 columns]

Setelah Case Folding
0      simple praktis dan sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah juga ini aplikasi tidak ada fitur pushnot...  3
3      sangat membantu  5
4      luar biasa  5
...
95     kenapa sekarang perpanjang rujukan pasien hrs ...  1
96     koplollllll nunggu verifikasi nomor hp mesti se...  1
97     gaje banget mau daftar aja dipersulit gini.kah...  1
98     lagi lihr kah . koo gangguan mlu mau login ...  3
99     sangat membantu  5
    
```

Gambar 7. Hasil proses case folding

Tahapan ini memungkinkan teks ulasan yang dikelola menjadi lebih konsisten.

3.3.2. Stopword Removal

Ulasan pengguna biasanya menggunakan bahasa yang tidak formal, singkatan, atau bahkan istilah percakapan sehari-hari. Oleh karena itu, diperlukan normalisasi pada kata singkatan. Berikut ini sampel normalisasi kata.

Table 2. Normalisasi kata

Sebelum Normalisasi	Setelah Normalisasi
yg	yang
dgn	dengan
utk	untuk
tdk	tidak
blm	belum
bgt	banget

Sementara itu, bentuk bahasa percakapan seperti kata gak, ga, nggak, dan ngga juga dinormalisasi menjadi kata “tidak”. Proses ini dilakukan guna mengurangi berbagai macam penulisan kata yang memiliki makna yang sama, sehingga meningkatkan konsistensi representasi teks pada tahap pembobotan kata atau TF-IDF. Hasilnya dapat dilihat pada Gambar 8 berikut ini.

```

Sebelum Stopword Removal
0      simple praktis dan sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah juga ini aplikasi tidak ada fitur pushnot...  3
3      sangat membantu  5
4      luar biasa  5
...
95     kenapa sekarang perpanjang rujukan pasien hrs ...  1
96     koplollllll nunggu verifikasi nomor hp mesti se...  1
97     gaje banget mau daftar aja dipersulit gini.kah...  1
98     lagi lihr kah . koo gangguan mlu mau login ...  3
99     verifikasi wajah gagal terus tanpa ada alasan ...  5

[100 rows x 2 columns]

Setelah Stopword Removal
0      simple praktis sangat membantu  5
1      mudah dan bermanfaat bagi seluruh masyarakat ind...  5
2      parah aplikasi fitur pushnotan lran dilakuk...  3
3      sangat membantu  5
4      luar biasa  5
...
95     koplollllll nunggu verifikasi nomor hp mesti se...  1
96     gaje banget mau daftar aja dipersulit gini.kah...  1
97     lihr koo gangguan mlu mau login ...  3
98     verifikasi wajah gagal terus tanpa ada alasan ...  5
    
```

Gambar 8. Hasil Stopword Removal

Kata-kata seperti (di, ke, bisa, telah) ditiadakan, kata-kata tersebut merupakan *stopword* yang umum.

3.3.3. Tokenizing

Proses *tokenizing* digunakan untuk membagi kalimat menjadi kata yang lebih kecil, untuk memudahkan dalam pengelolaan dan agar kata

tersebut dapat mudah dimengerti oleh system komputer.

Sebelum Tokenizing			
	Ulasan	Rating	
0	simple praktis sangat membantu	5	
1	mudahan bermanfaat seluruh masyarakat indonesi...	5	
2	parah aplikasi fitur pembayaran iuran dilakuka...	3	
3	sangat membantu	5	
4	luar biasa	5	
Setelah Tokenizing			
	Ulasan	Rating	
0	[simple, praktis, sangat, membantu]	5	
1	[mudah, bermanfaat, seluruh, masyarakat, ind...	5	
2	[parah, aplikasi, fitur, pembayaran, iuran, di...	3	
3	[sangat, membantu]	5	
4	[luar, biasa]	5	

Gambar 9. Tokenizing

Gambar 9 menunjukkan hasil pembentukan token dari teks ulasan, di mana setiap katanya akan diapit oleh tanda kutip dan koma untuk dapat menjadi sebuah token dalam satu kalimat.

3.3.4. Stemming

Kata dasar pada tahapan tokenisasi akan diolah pada proses stemming yaitu pembentukan kata dasar. Proses melibatkan *pustaka Sastrawi* dalam pemrosesan untuk merubah kata – kata menjadi bentuk dasar, sehingga awalan dan akhiran dari kata dasarnya akan terhapus. *Library Sastrawi* digunakan untuk pemrosesan teks. Hasil *stemming* dapat dilihat pada Gambar 10 di bawah ini.

Sebelum Stemming			
	Ulasan	Rating	
0	['simple', 'praktis', 'sangat', 'membantu']	5	
1	['mudah', 'bermanfaat', 'seluruh', 'masyarakat...']	5	
2	['parah', 'aplikasi', 'fitur', 'pembayaran', '...']	3	
3	['sangat', 'membantu']	5	
4	['luar', 'biasa']	5	
Setelah Stemming			
	Ulasan	Rating	
0	simple praktis sangat bantu	5	
1	mudah manfaat seluruh masyarakat indonesia ami...	5	
2	parah aplikasi fitur bayar iur laku mandirivir...	3	
3	sangat bantu	5	
4	luar biasa	5	

Gambar 10. Hasil data *stemming*

Setelah semua tahapan dilakukan, mulai dari pembersihan data, *case folding*, *stopword removal*, tokenisasi dan pembentukan *stemming*. Rangkaian tahap keseluruhan dapat dilihat pada Gambar 11 berikut.

Ulasan Asli	Cleaning	Case Folding	Stopword Removal	Tokenizing	Stemming
simple, praktis dan sangat membantu	simple praktis dan sangat membantu	simple praktis dan sangat membantu	simple praktis sangat membantu	['simple', 'praktis', 'sangat', 'membantu']	simple praktis sangat bantu
mudahan bermanfaat bagi seluruh masyarakat indonesia...	mudahan bermanfaat bagi seluruh masyarakat indonesia amin	mudahan bermanfaat bagi seluruh masyarakat indonesia amin	mudahan bermanfaat seluruh masyarakat indonesia amin	['mudah', 'bermanfaat', 'seluruh', 'masyarakat', 'indonesia', 'amin']	mudah manfaat seluruh masyarakat indonesia amin
parah juga ini aplikasi tidak ada fitur pembayaran iuran yang bisa dilakukan secara mandirivirtual account	parah juga ini aplikasi tidak ada fitur pembayaran iuran yang bisa dilakukan secara mandirivirtual account	parah juga ini aplikasi tidak ada fitur pembayaran iuran yang bisa dilakukan secara mandirivirtual account	parah aplikasi fitur pembayaran iuran dilakukan mandirivirtual account	['parah', 'aplikasi', 'fitur', 'pembayaran', 'iuran', 'dilakukan', 'mandirivirtual', 'account']	parah aplikasi fitur bayar iur laku mandirivirtual account
sangat membantu	sangat membantu	sangat membantu	sangat membantu	['sangat', 'membantu']	sangat bantu
Luar biasa	Luar biasa	luar biasa	luar biasa	['luar', 'biasa']	luar biasa

Gambar 11. Sebelum dan sesudah *Data preprocessing*

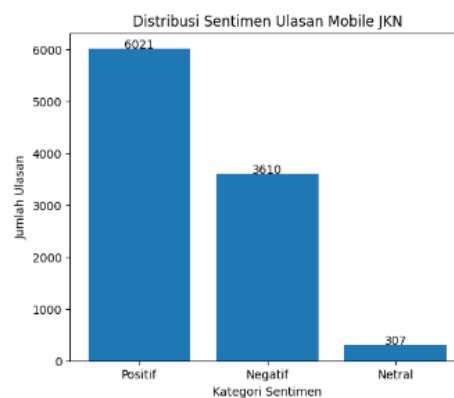
3.4. Labelling

Pemberian label dapat membantu sistem untuk memahami bagaimana *user* merasakan dan mengevaluasi aplikasi serta layanan. Setelah pemrosesan data, selanjutnya akan diberi label. Sebagai berikut:

Bagus untuk para penjual online & UMKM	5	bagus untuk para penjual online umum	bagus penjual online umum	['bagus', 'penjual', 'online', 'umum']	bagus jual online umum	Positif
Kapan kami bisa berjalan lagi di titik sedi...	1	kapan kami bisa berjalan lagi di titik sedi...	berjalan titik sedi...	['kapan', 'kami', 'berjalan', 'titik', 'sedi', 'banger', 'ga...']	jual titik sedi banger ga jual titik	Negatif
Apk ngga jelas akun buat apa apa engga di hapu...	1	apk ngga jelas akun buat apa apa engga di hapu...	hapu itu saldonya merugi...	['apk', 'ngga', 'jelas', 'akun', 'buat', 'apa', 'apa', 'engga', 'di', 'hapu', 'itu', 'saldonya', 'merugi...']	apk ngga akun engga hapu itu saldo ngg guna...	Negatif
tolong titik shop dibuka lagi kak...	5	tolong titik shop dibuka lagi kak...	tolong titik shop dibuka lagi kak...	['tolong', 'titik', 'shop', 'dibuka', 'lagi', 'kak...']	tolong titik shop buka kak	Positif
sangat membantu	5	sangat membantu	membantu	['membantu']	bantu	Positif
Manganya jgn susah susah kalo tulup pement...	1	manganya jgn susah susah kalo tulup pement...	manganya jgn susah susah kalo tulup pement...	['manganya', 'jgn', 'susah', 'susah', 'kalo', 'tulup', 'pement...']	manganya jgn susah susah kalo tulup pement...	Negatif

Gambar 12. Hasil Pelabelan Data

Penentuan kategori sentimen akan diberikan berdasarkan nilai rating yang dipilih oleh *user*.



Gambar 13. Hasil Distribusi Data

Pada proses pemberian label, data dibagi menjadi 3 kategori atau disusun dalam kombinasi rating, yaitu sentimen positif, negatif, dan netral. Namun, untuk menghindari ambiguitas data dan kurangnya akurasi, pada penelitian ini sentimen netral dihilangkan atau dikeluarkan dan yang digunakan hanya sentimen positif dan negatif. Penghapusan data netral bertujuan untuk memperoleh pemisahan kelas yang lebih jelas sehingga model dapat mempelajari karakteristik sentimen positif dan negatif secara lebih optimal.

Ulasan pengguna positif memiliki rating antara 4-5, sedangkan yang negatif berasal dari rating 1-2. Dari Gambar 10 ditunjukkan bahwa hasil pembersihan data dengan total 9.938, terdapat 6.021 sentimen positif, 3.610 sentimen negatif dan 307 sentimen netral. Pada penelitian ini, data yang diproses kategori sentimen positif dan kategori sentimen negatif, dengan demikian data yang digunakan sebanyak 9.631 data.

3.5. TF-IDF

Proses TF-IDF hanya di-fit pada data latih saja melalui fungsi *fit_transform()*, sedangkan data uji hanya ditransformasikan menggunakan fungsi *transform()* berdasarkan *vocabulary* dan bobot IDF yang diperoleh dari data latih. Setelah dilakukan proses TF-IDF, akan memiliki bobot pada angka di setiap baris kata. Bobot setiap *term* dihitung untuk semua ulasan, yang bertujuan untuk menghitung setiap kata-kata yang akan muncul pada ulasan.

Semakin tinggi bobot atau nilai *term*, maka semakin sering kata tersebut muncul, 10 sampel data dapat dilihat dari TF-IDF pada Tabel 3 di bawah.

Table 3. Top 10 TF-IDF

Index	Term	Average TF-IDF
5060	sangat	0.072351
542	bantu	0.060050
437	bagus	0.059732
3379	mantap	0.039744
3757	mudah	0.039114
304	aplikasi	0.036888
471	baik	0.036105
3834	nan	0.030087
1198	daftar	0.021782
3092	layan	0.019599

Bobot dihitung pada setiap kata berdasarkan frekuensi kemunculannya pada dokumen tertentu, dan tingkat kelangkaannya pada semua dokumen. Pada 10 data uji terdapat nilai rata-rata sekitar 0.072351 pada kemunculan kata sangat, sehingga ada kemungkinan kata tersebut memiliki tingkat keseringan muncul paling banyak dibandingkan dengan yang lain. Dengan demikian, kata-kata yang terlalu umum akan mendapatkan bobot yang lebih kecil.

Untuk proses klasifikasi sentimen, TF-IDF dapat menghasilkan fitur yang lebih representatif. TF-IDF lebih sederhana, efisien, dan memiliki kompatibilitas yang baik dengan algoritma Multinomial Naive Bayes, sehingga sesuai untuk penelitian ini.

Selanjutnya, penggunaan *unigram* dipilih karena struktur ulasan pengguna Mobile JKN umumnya pendek dan mengandung kata-kata sentimen yang dapat berdiri sendiri, seperti (bagus, lambat, error, mudah, dan sulit). Selain

itu, penggunaan *unigram* menghasilkan jumlah fitur yang lebih terkendali dengan demikian dapat mengurangi kerumitan kemunculan data dan pada saat proses klasifikasi dengan algoritma akan menjaga efisiensinya.

```
Tfidf_code = TfidfVectorizer(
    ngram_range=(1,1),
    max_features=5000,
    min_df=2 )
```

Berikutnya pada parameter *min_df=2* berfungsi untuk melewati atau mengabaikan kata yang muncul hanya pada satu dokumen, sehingga fitur yang sangat jarang dapat dikurangi. Selain itu, parameter *max_features=5000* digunakan untuk membatasi jumlah fitur maksimum yang digunakan dalam proses klasifikasi. Untuk pembentukan *n-gram*, digunakan *ngram_range=(1,1)* yang menunjukkan penggunaan *unigram*, di mana setiap kata diperlakukan sebagai satu fitur.

3.6. Algoritma Naive Bayes

Setelah tahapan TF-IDF dilakukan, selanjutnya dilakukan implementasi model *Multinomial Naive Bayes* menggunakan bahasa pemrograman Python.

```
X_tfidf = tfidf.fit_transform(X)
# ratio 80 : 20

X_train, X_test, y_train, y_test = train_test_split(
    X_tfidf,
    y,
    test_size=0.20,
    random_state=42
)

# MODEL MULTINOMIAL NAIVE BAYES
model = MultinomialNB()

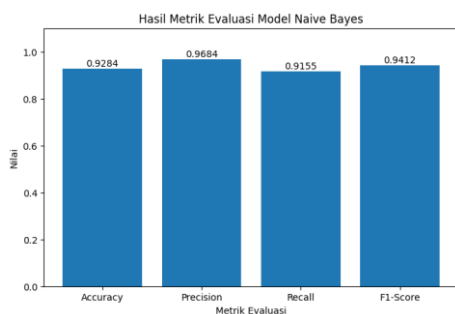
# Training model
model.fit(X_train, y_train)

# Prediksi data testing
y_pred = model.predict(X_test)
```

Gambar 14. *Multinomial Naive Bayes*

Pada proses pelatihan digunakan teknik *train_test_split*, dimana membagi data jadi dua yaitu data pelatihan atau *training* dengan dengan data *testing* atau pengujian. Pada rasio 80:20, data pelatihan berjumlah 80% dan pengujian 20%. Karena kelas netral tidak disertakan, *dataset* yang tersedia yaitu sebanyak 9.631 data ulasan. Jumlah ini adalah hasil dari pemilihan kelas positif dan negatif dari proses pelabelan. Dari data tersebut akan didapat 80% x 9.631, yaitu 7.704 data *training*, dan 20% x 9.631, yaitu 1.927 data *testing*. Selanjutnya, hasil data testing yaitu 1.927 data, disajikan ke dalam bentuk

confusion matrix. Kemudian model dievaluasi dengan akurasi, presisi, *recall*, dan *F1-score* berdasarkan kinerja algoritma pada *Confusion Matrix*. Metrik evaluasi dengan Naive Bayes dapat dilihat pada grafik berikut.

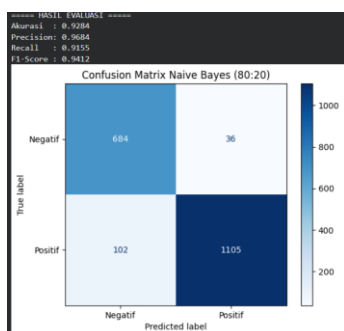


Gambar 15. Evaluasi model Naive Bayes

Data evaluasi menunjukkan *accuracy* 0,9284, *precision* 0,9684, *recall* 0,9155 dan *F1-score* 0,9412. Nilai *accuracy* sebesar 0,9284 atau 92,84% menunjukkan akurasi dalam prediksi sudah cukup baik. Nilai *precision* yang diperoleh sebesar 0,9684 atau 96,84%, Angka ini menunjukkan bahwa prediksi yang dihasilkan adalah sebagian besar prediksi dengan benar. *Recall* sebesar 0,9155 atau 91,55% menunjukkan bahwa model sebagian besar mampu mengidentifikasi data yang benar dari kelas positif. Selain itu, *F1-score* sebanyak 0,9412 atau 94,12% menunjukkan adanya keseimbangan baik antara *precision* dan *recall*.

3.7 Confusion matrix

Model dengan rasio 80:20 akan menghasilkan visual sebagai berikut.



Gambar 16. Confusion Matrix rasio 80:20

Dari gambar tersebut dapat ditunjukkan bahwa *True Negative*, label negatif dengan data yang memiliki prediksi benar sebanyak 684 data sebagai label negatif. *True Positive*, label positif yang memiliki prediksi benar sebanyak 1.105 data

pada label positif. *False Negative*, label positif yang memiliki prediksi salah sebanyak 102 data sebagai label negatif. *False Positive*, label negatif yang memiliki prediksi salah sebanyak 36 data sebagai label positif. Hasil tersebut menunjukkan bahwa model secara umum sudah mampu mengklasifikasikan data uji dengan baik.

3.8. Visualisasi Word Cloud

Visualisasi sentimen positif dan negatif yang dominan menggunakan *word cloud* untuk identifikasi kata yang kerap muncul pada jenis sentimen.

Dari analisis yang dilakukan, adanya kata yang kerap muncul pada ulasan positif, misalnya kata mantap, bantu, sangat, mudah, dan aplikasi. Selain itu, diperoleh keluhan pengguna yang paling dominan pada ulasan negatif, di antaranya yang berkaitan dengan proses *login* dan autentikasi pengguna, seperti gagal *login*, masalah kode OTP, gagal verifikasi akun serta *error* aplikasi. Gambar berikut memperlihatkan visualisasi dengan *word cloud*.



Gambar 17. Word Cloud Positif



Gambar 18. Word Cloud Negatif

Berdasarkan visualisasi tersebut, didapat hasil perhitungan frekuensi kemunculan kata tersebut, Tabel frekuensi kemunculan sentimen dapat dilihat sebagai berikut.

Table 4. Frekuensi sentimen positif

No	Kata	Frekuensi
1	sangat	1621
2	bantu	1179
3	mudah	951
4	bagus	853
5	aplikasi	575
6	baik	509

7	mantap	458
8	layan	427
9	jkn	317
10	cepat	302

Table 5. Frekuensi sentimen negatif

No	Kata	Frekuensi
1	aplikasi	1633
2	daftar	1059
3	mau	1034
4	ga	880
5	masuk	851
6	nya	771
7	gak	761
8	kode	757
9	otp	744
10	susah	658

Pada Tabel 4, sepuluh kata yang dominan pada sentimen positif menunjukkan bahwa pengguna memberikan apresiasi dan penilaian positif pada aplikasi mobile JKN dan pemanfaatannya. Kualitas layanan yang diberikan dan manfaat aplikasi dalam mendukung akses layanan kesehatan. Sentimen tersebut mengindikasikan bahwa aplikasi dapat memberikan kemudahan bagi masyarakat. Sedangkan pada sentimen negatif, pengguna menyampaikan kendala teknis dan kemunculan kata-kata negatif. Temuan ini menunjukkan bahwa sisi aplikasi dan mekanisme pelayanan masih perlu ditingkatkan.

4. PENUTUP

4.1. Kesimpulan

Hasil penelitian dengan implementasi algoritma Naive Bayes secara keseluruhan mendapatkan performa yang cukup baik. Total *dataset* yang digunakan adalah 10.000 data yang melewati tahapan *data cleaning* dan *data preprocessing* yang menghasilkan 9.938 data. Tahapan *data processing* meliputi tahapan *case folding*, *stopword removal*, tokenisasi, dan *stemming*. *Labeling* dengan membaginya menjadi 3 kategori yaitu positif, negatif dan netral. Namun, hanya kelas positif dan negatif yang dipakai pada proses klasifikasi, sedangkan kelas netral dihilangkan, sehingga *dataset* ulasan menjadi 9.631 data. Pembobotan kata menggunakan TF-IDF dan pemodelan dengan Naive Bayes. Evaluasi model juga digunakan

dengan *confusion matrix* dan visualisasi *word cloud*.

Dari metrik evaluasi model didapat hasil bahwa nilai *accuracy* 0,9284 atau 92,84%, *precision* 0,9684 atau 96,84%, *recall* 0,9155 atau 91,55% dan *F1-score* 0,9412 atau 94,12%. Model memiliki kemampuan dalam klasifikasi analisis sentimen yang cukup baik. Hasil *confusion matrix* menunjukkan bahwa *True Negative* sebanyak 684, *True Positive* 1.105, *False Negative* 102, dan *False Positive* 36.

Pada analisis visualisasi *word cloud*, terdapat kata yang sering muncul pada sentimen positif atau negatif. Seperti kata sangat, bantu, sudah dan aplikasi. Selain itu, terdapat juga pengguna yang paling dominan memberikan ulasan negatif, di antaranya yang berkaitan dengan proses *login* dan autentikasi pengguna, seperti gagal *login*, masalah kode OTP, dan gagal verifikasi akun, serta keluhan pengguna terkait *error* aplikasi.

Hasil dari model dapat memberikan informasi yang berguna bagi pengembang aplikasi dan pelayanan kesehatan publik dalam menganalisis tingkat kepuasan pengguna. Serta, menjadi bahan evaluasi dan identifikasi dari keluhan negatif yang muncul untuk meningkatkan kualitas layanan.

4.2. Saran

Penelitian selanjutnya disarankan untuk membandingkan hasil model klasifikasi lanjutan menggunakan *dataset* yang besar dan algoritma lainnya, misalnya *Random Forest* atau *Support Vector Machine*, agar dapat memperdalam analisis dan mengetahui metode paling efektif.

Penelitian ini memerlukan model yang lebih sensitif terhadap ulasan negatif. Agar model mendapatkan kemampuan dan hasil yang lebih baik dalam menganalisis sentimen. Dengan begitu, pada penelitian ini perlu mendapat perhatian terhadap kekurangan dan keterbatasannya, agar dijadikan bahan pengembangan penelitian lebih lanjut.

Selain itu, disarankan untuk meningkatkan kinerja model analisis sentimen pada area industri agar didapat komparasi dengan penelitian ini pada lingkup analisis sentimen.

5. DAFTAR PUSTAKA

- Amanda, R., & Negara, E. S. (2020). Analysis and implementation machine learning for YouTube data classification by comparing the performance of classification algorithms. *JOIN (Jurnal Online Informatika)*, 5(1), 61–72. <https://doi.org/10.15575/join.v5i1.505>
- Arminda, N. F., Sulistiyowati, N., & Padilah, T. N. (2023). Implementasi algoritma multinomial Naive Bayes pada analisis sentimen terhadap ulasan pengguna aplikasi BRImo. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(3), 1817–1822. <https://doi.org/10.36040/jati.v7i3.7012>
- Bunawan, I. (2026). Perbandingan support vector machine dan Naive Bayes dalam mengklasifikasikan emosi teks bahasa Palembang. *JIKO (JURNAL INFORMATIKA DANKOMPUTER)*, 10(1), 193–201. <https://doi.org/10.26798/jiko.v10i1.2399>
- Chaithra, V. D. (2019). Hybrid approach: Naive Bayes and sentiment VADER for analyzing sentiment of mobile unboxing video comments. *International Journal of Electrical and Computer Engineering (IJECE)*, 9(5), 4452–4459. <https://doi.org/10.11591/ijece.v9i5.pp4452-4459>
- Darmawan, G., Alam, S., & Sulistyo, M. I. (2023). Analisis sentimen berdasarkan ulasan pengguna aplikasi MyPertamina pada Google Play Store menggunakan metode Naive Bayes. *STORAGE – Jurnal Ilmiah Teknik Dan Ilmu Komputer*, 2(3), 100–108. <https://doi.org/10.55123/storage.v2i3.2333>
- Enjeli, M., & Wijaya, A. (2024). Analisis sentimen penggunaan aplikasi KineMaster menggunakan metode Naive Bayes. *JURNAL ILMIAH COMPUTER SCIENCE (JICS)*, 2(2), 89–98. <https://doi.org/10.58602/jics.v2i2.24>
- Fathoni, Ansori, A. F., Ramadhani, I. N., Anissa, C. R., & Putri, S. A. (2025). Analisis sentimen masyarakat Indonesia di Twitter terhadap sistem perpajakan Coretax menggunakan metode Naive Bayes. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(4), 6749–6753. <https://doi.org/10.36040/jati.v9i4.14214>
- Fatra, E., Manguma, T. T. F., & Mastura. (2025). Analisis sentimen media sosial menggunakan metode Naive Bayes. *Culture Education and Technology Research (Cetera)*, 2(3), 456–466. <https://doi.org/10.31004/ctr.v2i2.166>
- Ghozali, M. I., Sugiharto, W. H., & Iskandar, A. F. (2023). Analisis sentimen pinjaman online di media sosial Twitter menggunakan metode Naive Bayes. *Kajian Ilmiah Informatika Dan Komputer*, 3(6), 1340–1348. <https://doi.org/10.30865/klik.v3i6.936>
- Hayatin, N., Marthasari, G. I., & Nuraini, L. (2020). Optimization of sentiment analysis for Indonesian presidential election using Naive Bayes and Particle Swarm Optimization. *JOIN (Jurnal Online Informatika)*, 5(1), 81–88. <https://doi.org/10.15575/join.v5i1.558>
- Jelita, H. P., Saad, M. I., & Wahyuni. (2025). Penerapan algoritma Naive Bayes dalam analisis sentimen masyarakat terhadap STMIK Widya Cipta Dharma. *Bulletin of Information Technology (BIT)*, 6(2), 148–160. <https://doi.org/10.47065/bit.v5i2.2029>
- Mandasari, P. A., & Faizin, A. (2025). Analisis sentimen ulasan aplikasi DANA di Google Play Store menggunakan algoritma Naive Bayes. *Jurnal Sains Student Research*, 3(5), 1082–1093. <https://doi.org/10.61722/jssr.v3i5.6436>
- Maulana, B. A., Fahmi, M. J., Imran, A. M., & Hidayati, N. (2024). Sentiment analysis of Pluang applications with Naive Bayes and support vector machine (SVM) algorithm. *Indonesian Journal of Machine Learning and Computer Science*, 4(April), 375–384. <https://doi.org/10.57152/malcom.v4i2.1206>
- Punusingon, P. V., Yusupa, A., & Tarigan, V. (2025). Analisis sentimen pemilihan remaja teladan wilayah Ratahan menggunakan algoritma Naive Bayes. *Journal Computer and Technology*, 3(1), 65–73. <https://doi.org/10.69916/comtechno.v3i1.337>
- Ramadhani, A. A., Saputra, R. A., & Ningrum, P. (2024). Klasifikasi tingkat kepuasan pengguna Google Classroom dalam

pembelajaran online menggunakan algoritma Naive Bayes. *JIKO (JURNAL INFORMATIKA DANKOMPUTER)*, 8(2), 310–317.

<https://doi.org/10.26798/jiko.v8i2.1221>

Somantri, O., Purwaningrum, S., & Maharrani, R. H. (2024). Sentiment review of coastal assessment using neural network and Naive Bayes. *International Journal of Electrical and Computer Engineering (IJECE)*, 14(1), 681–689.

<https://doi.org/10.11591/ijece.v14i1.pp681-689>

Yanah, S., Purbaratri, W., Paylina, S., Safitri, A. N. I., & Tachjar, N. K. (2025). Analisis sentimen aplikasi Pemilu menggunakan algoritma Naive Bayes. *JURNAL ELEKTRO & INFORMATIKA*, 05(01), 131–139.

<https://doi.org/10.56486/jeis.vol5no1.684>