

IMPLEMENTASI MODEL ANALISIS SENTIMEN TERHADAP GRUP MUSIK BTS MENGGUNAKAN METODE NAÏVE BAYES

Siti Aiwestopa Riyandona^{1*}, Nining Rahaningsih², Raditya Danar Dana³, Mulyawan⁴

^{1,2,3,4} STMIK IKMI Cirebon, Jl. Perjuangan No.10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Jawa Barat 45135, Telp: (0231) 490480

Received: 25 Desember 2024
Accepted: 14 Januari 2025
Published: 20 Januari 2025

Keywords:

Naïve bayes, Analisis Sentimen, Grup Musik BTS, Hiatus, Twitter.

Correspondent Email:

riandonasitiwastopa@gmail.com

Abstrak. BTS saat ini sedang menjalani masa hiatus karena beberapa anggota memenuhi kewajiban wajib militer di Korea Selatan. Meski tidak aktif secara grup, pencapaian individu dan kolaborasi para anggota tetap menarik perhatian. Namun, isu negatif yang beredar di media sosial berpotensi memengaruhi pandangan publik terhadap grup ini. Penelitian ini bertujuan menganalisis sentimen pengguna *Twitter* terhadap BTS selama masa hiatus dengan algoritma *Naïve Bayes*, yang efektif untuk analisis sentimen teks. Data dikumpulkan menggunakan teknik *crawling* pada *tweet* terkait BTS selama Mei–Oktober 2024, lalu diproses melalui pembersihan data, normalisasi, tokenisasi, dan pembobotan menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF). Model klasifikasi menghasilkan akurasi 78,33%, Presisi 79,25%, *Recall* 78,33%, dan *F1-Score* 78,49% dengan sentimen positif dominan, mencerminkan dukungan penggemar yang kuat, meski sentimen negatif dan netral juga muncul. Penelitian ini memberikan wawasan tentang reaksi penggemar dan membuktikan efektivitas analisis sentimen dalam memahami interaksi di media sosial.

Abstract. *BTS is currently on hiatus as some members fulfill their mandatory military service in South Korea. Although they are not active as a group, the individual achievements and collaborations of the members continue to attract attention. However, negative issues circulating on social media have the potential to affect the public's perception of this group. This study aims to analyze Twitter users' sentiment towards BTS during their hiatus using the Naïve Bayes algorithm, which is effective for text sentiment analysis. Data was collected using crawling techniques on tweets related to BTS during May–October 2024, then processed through data cleaning, normalization, tokenization, and weighting using Term Frequency-Inverse Document Frequency (TF-IDF). The classification model produced an accuracy of 78.33%, Precision of 79.25%, Recall of 78.33%, and an F1-Score of 78.49% with a dominant positive sentiment, reflecting strong fan support, although negative and neutral sentiments also appeared. This research provides insights into fan reactions and demonstrates the effectiveness of sentiment analysis in understanding interactions on social media.*

1. PENDAHULUAN

Bangtan Sonyeondan (BTS), grup musik asal Korea Selatan yang terdiri dari tujuh anggota, yaitu RM, Jin, Suga, J Hope, Jimin, V,

dan JK telah menjadi ikon global dalam industri K-Pop[1]. Popularitas mereka didukung oleh penggemar setia, ARMY, yang tersebar di seluruh dunia[2]. Salah satu *platform* yang banyak digunakan oleh masyarakat

Indonesia sebagai sarana menyalurkan pendapat yaitu Twitter. Dengan kemudahan ini, para pengguna media sosial lebih memilih Twitter untuk diskusi dan debat terkait isu-isu terkini[3]. Namun, saat ini BTS sedang menjalani masa hiatus karena beberapa anggotanya melaksanakan wajib militer, sebuah kewajiban bagi warga negara Korea Selatan. Meski grup tidak aktif dalam kegiatan peluncuran musik baru, prestasi individu anggotanya tetap menjadi sorotan.

Di sisi lain, masa hiatus ini juga diwarnai dengan isu-isu negatif yang beredar di media sosial, termasuk tuduhan tidak berdasar dari beberapa pihak, yang dilansir dari akun @dispatchsns. Hal ini memunculkan pertanyaan tentang bagaimana penggemar merespons pencapaian BTS selama hiatus dan sejauh mana popularitas mereka dapat bertahan di tengah tantangan tersebut. Sentimen publik yang diekspresikan di media sosial memiliki peran penting dalam mencerminkan opini dan interaksi, sehingga analisis terhadap data ini menjadi sangat relevan.

Penelitian ini bertujuan untuk menganalisis sentimen pengguna Twitter terhadap BTS menggunakan algoritma Naïve Bayes, yang telah terbukti efektif dalam klasifikasi teks. Dengan fokus pada sentimen positif, negatif, dan netral, penelitian ini diharapkan memberikan wawasan baru mengenai reaksi publik terhadap BTS selama masa hiatus, sekaligus membuktikan efektivitas algoritma Naïve Bayes dalam analisis sentimen berbasis data media sosial.

2. TINJAUAN PUSTAKA

Penelitian yang berjudul Implementasi model BERT pada analisis sentimen Pengguna *twitter* terhadap aksi boikot produk Israel[4]. Penelitian ini menggunakan *IndoBERT* untuk menganalisis sentimen publik di *Twitter* terkait aksi boikot produk Israel. Hasilnya menunjukkan dominasi sentimen positif, yang mencerminkan dukungan masyarakat terhadap aksi tersebut. Dengan akurasi, presisi, *recall*, dan *F1-score* sebesar 85%, penelitian ini menegaskan efektivitas

IndoBERT dalam klasifikasi sentimen media sosial serta memberikan wawasan mendalam tentang opini publik terhadap isu sosial ini.

Penelitian lainnya yang berjudul mengevaluasi sentimen pengguna aplikasi SIREKAP di Play Store menggunakan algoritma *Random Forest Classifier*. Dengan pendekatan *Knowledge Discovery in Database (KDD)*[5], analisis dilakukan terhadap 5000 ulasan, menghasilkan akurasi 74%, presisi 75%, *recall* 74%, dan *F1-score* 74%. Mayoritas ulasan bernada negatif, sebanyak 4002, dibandingkan 762 ulasan positif dan 234 netral. Hasil ini menunjukkan perlunya peningkatan kualitas aplikasi untuk memenuhi kebutuhan pengguna. Studi ini memberikan wawasan berharga bagi pengembang dalam memperbaiki aplikasi demi mendukung proses demokrasi yang lebih efektif di era digital.

2.1 Analisis Sentimen

Analisis sentimen adalah proses mengidentifikasi dan mengklasifikasikan opini atau emosi dalam sebuah teks menjadi kategori tertentu, seperti positif, negatif, atau netral. Teknik ini memanfaatkan Natural Language Processing (NLP) untuk mengolah data teks tidak terstruktur, sehingga dapat digunakan untuk memahami pola sentimen publik terhadap suatu topik tertentu[6][7]

2.1 Algoritma Naïve Bayes

Naïve Bayes adalah algoritma berbasis probabilitas yang sederhana namun efektif untuk klasifikasi data teks. Algoritma ini bekerja dengan menghitung probabilitas suatu data termasuk dalam kategori tertentu berdasarkan fitur yang ada, menggunakan asumsi independensi antar fitur. Dalam konteks analisis sentimen, *Naïve Bayes* sering digunakan karena efisien dalam mengolah data teks besar dan memberikan hasil yang cukup akurat, terutama ketika data telah diproses dengan baik[8].

2.2 Pra-pemrosesan Data Teks

Tahapan pra-pemrosesan data meliputi pembersihan teks (*cleansing*), normalisasi, tokenisasi, penghapusan kata-kata umum

(stopword removal), dan stemming. Langkah-langkah ini bertujuan untuk menghilangkan elemen-elemen yang tidak relevan dan menyederhanakan teks agar dapat diolah lebih mudah oleh algoritma machine learning. Selain itu, metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan untuk memberikan bobot pada kata-kata penting dalam teks[9].

2.3 Twitter sebagai Sumber Data

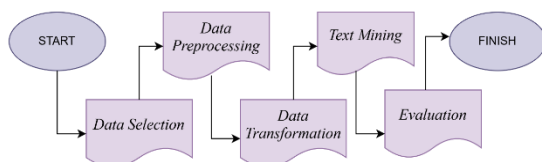
Twitter adalah platform media sosial yang sering digunakan untuk mengungkapkan opini dan diskusi mengenai berbagai isu. Dengan karakteristik teks yang singkat, data dari Twitter membutuhkan pengolahan khusus untuk menangani singkatan, hashtag, atau simbol yang sering muncul. Data yang diperoleh melalui teknik crawling ini sangat relevan untuk analisis sentimen [10].

2.4 Evaluasi Model

Kinerja algoritma Naïve Bayes dalam penelitian ini diukur menggunakan metrik evaluasi seperti akurasi, presisi, recall, dan F1-score. Metrik ini digunakan untuk mengevaluasi kemampuan model dalam mengklasifikasikan data dengan benar dan memberikan hasil yang konsisten[11].

3. METODE PENELITIAN

3.1 Rancangan Penelitian



Gambar 1 Metode Penelitian

Penelitian ini dirancang untuk menganalisis sentimen pengguna Twitter terhadap grup musik BTS selama masa hiatus dengan menggunakan algoritma *Naïve Bayes*. Model penelitian menggunakan pendekatan *Knowledge Discovery in Database* (KDD), yang terdiri dari tahapan-tahapan berikut:

Tabel 1 Keterangan tahapan Metode penelitian

Tahapan	Aktivitas
<i>Data Selection</i>	Menggunakan teknik crawling untuk mengunduh data dari Twitter.
<i>Data Preprocessing</i>	Menerapkan tahapan Pemrosesan data. Pemrosesan ini melalui beberapa langkah yaitu: <i>Cleansing, case folding Tokenization, Normalization, Stopword Removal, dan Stemming.</i>
<i>Data Transformation</i>	Mengubah data teks menjadi representasi numerik menggunakan metode TF-IDF.
<i>Text Mining</i>	Menggunakan algoritma <i>Naïve Bayes</i> untuk mengklasifikasikan sentimen.
<i>Evaluation</i>	Mengukur performa model dengan metrik evaluasi.

3.2 Teknik Pengumpulan Data

Data dikumpulkan dari Twitter menggunakan pustaka *Python Tweepy*, yang memungkinkan akses ke API Twitter. Kata kunci seperti “BTS,” “Bangtan,” “BTS ARMY,” dan nama anggota BTS digunakan untuk menyaring *tweet* yang relevan. Data dikumpulkan untuk periode Mei hingga Oktober 2024, menghasilkan total 1.623 *tweet*. Data yang diperoleh mencakup teks *tweet*, waktu posting, dan metadata tambahan seperti jumlah *retweet* dan *likes*.

3.3 Sumber Data

Data yang digunakan dalam penelitian ini adalah data sekunder yang bersumber dari Twitter. Platform ini dipilih karena merupakan salah satu media sosial utama yang digunakan oleh penggemar BTS untuk berdiskusi dan mengekspresikan opini. Data yang diambil

adalah tweet publik yang tersedia secara bebas dan tunduk pada kebijakan penggunaan *Twitter* API.

3.4 Teknik Analisis Data

1) Pra-pemrosesan Data:

Data yang terkumpul melalui proses *crawling* diproses dengan langkah-langkah berikut:

- Cleansing*: Menghapus simbol, URL, hashtag, dan elemen teks yang tidak relevan.
- Case Folding*: Mengubah semua teks menjadi huruf kecil.
- Tokenization*: Memecah teks menjadi kata-kata individual.
- Normalization*: Mengubah kata tidak baku menjadi kata baku.
- Stopword Removal*: Menghapus kata umum yang tidak signifikan seperti “di,” “ke,” dan “yang.”
- Stemming*: Mengubah kata berimbuhan menjadi kata dasar.

2) Transformasi Data:

Data yang telah diproses diubah menjadi representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Teknik ini memberikan bobot pada kata-kata yang sering muncul dalam dokumen tertentu tetapi jarang muncul di dokumen lain, sehingga lebih relevan untuk analisis sentimen.

3) Pembangunan Model:

Model klasifikasi sentimen dibangun menggunakan algoritma *Naïve Bayes*. Algoritma ini bekerja dengan menghitung probabilitas setiap tweet termasuk dalam kategori positif, negatif, atau netral.

4) Evaluasi Model:

Evaluasi dilakukan menggunakan metrik seperti akurasi, presisi, *recall*, dan *F1-score*. *Confusion Matrix* digunakan untuk menilai kinerja model berdasarkan prediksi yang benar dan salah terhadap data uji.

3.5 Implementasi dan Reproduksi

Seluruh proses penelitian dilakukan menggunakan bahasa pemrograman *Python* di *Google Colab* untuk memanfaatkan komputasi berbasis *cloud*. Skrip yang digunakan mencakup fungsi-fungsi untuk *crawling* data, pra-pemrosesan, transformasi, pembangunan model, dan evaluasi. Dengan penjelasan rinci pada setiap tahap, penelitian ini dapat diulang oleh peneliti lain untuk analisis sentimen pada data serupa.

4. HASIL DAN PEMBAHASAN

4.1 Hasil Penelitian

4.1.1 Data Selection

Pada penelitian ini, dilakukan pengumpulan data *twitter* dengan proses *crawling* data menggunakan *python*. Data yang dikumpulkan mencakup informasi berupa cuitan (*tweets*), metadata (misalnya, waktu posting, jumlah *retweet*, jumlah *reply*, jumlah *quote* atau komentar, dan jumlah *likes*), serta *username* atau informasi akun pengguna. Proses pengumpulan data menggunakan metode *web scraping* atau *crawling* dengan menggunakan pustaka pemrograman khusus yang sudah terkoneksi dengan kode *API twitter*, dan dirancang untuk mengakses, mengambil, dan menyimpan data dari *Twitter* secara sistematis. Tampilan data hasil *crawling* dapat dilihat pada gambar berikut :

tweet_id	tweet_text	lang	quote_count	reply_count	retweet_count	username
1	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
2	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
3	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
4	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
...	...	...	...	...	...	...
1618	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
1619	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
1620	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
1621	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts
1622	Siapa yang bilang BTS bukan band? Mereka...	id	0.0	0.0	0.0	ufficialbts

Gambar 2 Dataset Hasil Crawling

Setelah menyelesaikan proses pengambilan data dengan proses *crawling*, hasilnya akan disimpan dalam format *csv*. Data yang berhasil didapatkan adalah 1.623 data, yang selanjutnya dilakukan tahap pelabelan data untuk menentukan kategori sentimen setiap cuitan. Proses pelabelan dilakukan secara manual oleh seorang guru Bahasa Indonesia di SMAN 7 Cirebon, guna memastikan akurasi dalam menentukan apakah cuitan tersebut berisi sentimen positif, negatif, atau netral. Berikut tampilan data setelah pelabelan :

Gambar 3 Dataset Hasil Pelabelan

Dari 1.623 data tersebut memiliki 9 kolom atribut, dimana Setiap kolom merepresentasikan karakteristik atau informasi tertentu dari data yang digunakan dalam analisis. Untuk mengetahui informasi lebih rinci, bisa dilihat pada tabel dibawah ini:

Tabel 2 Keterangan setiap atribut

Atribut	Keterangan
<i>Created_at</i>	Disebut sebagai atribut waktu. Dan biasanya digunakan untuk analisis temporal, misalnya memantau perubahan sentimen berdasarkan waktu tertentu.
<i>Full_text</i>	Disebut sebagai teks utama atau data input utama, kolom ini berisi teks cuitan yang akan dianalisis untuk menentukan sentimen (positif, negatif, atau netral).Teks ini menjadi fokus utama dalam proses praproses data dan representasi teks (contoh: <i>TF-IDF</i> atau <i>Bag of Words</i>).
Label	Disebut sebagai target atau label sentimen Kolom ini adalah variabel yang digunakan dalam pelatihan model (<i>supervised learning</i>). Setelahnya menjadi keluaran atau hasil yang ingin diprediksi oleh model (positif, negatif, netral).
<i>Favorite_count</i> , <i>Retweet_count</i> ,	Atribut-atribut tersebut disebut sebagai atribut

<i>Quote_count</i> , <i>Reply_count</i>	metrik interaksi, yang memberikan informasi tambahan tentang popularitas atau tingkat interaksi pada cuitan.
<i>Lang</i>	Atribut ini sering disebut sebagai atribut bahasa, berguna untuk menyaring data berdasarkan bahasa tertentu sebelum analisis sentimen dilakukan.
<i>Username</i>	Disebut sebagai atribut identitas pengguna, kolom ini biasanya tidak digunakan langsung dalam analisis sentimen kecuali jika ada kebutuhan untuk menganalisis sentimen berdasarkan profil atau pola tertentu dari pengguna.

Tidak hanya keterangan dari setiap atribut, informasi lengkap dari masing-masing atribut data juga sama pentingnya. Menggunakan salah satu fungsi dalam pustaka pandas pada Python untuk menampilkan informasi ringkas tentang struktur sebuah *DataFrame*. Fungsi ini memberikan gambaran umum mengenai dataset yang sedang dianalisis, seperti jumlah baris, kolom, tipe data, dan jumlah nilai non-kosong (*non-null*) di setiap kolom.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 1623 entries, 0 to 1622
Data columns (total 8 columns):
#   Column              Non-Null Count  Dtype
---  -
0   created_at          1623 non-null   object
1   favorite_count      1623 non-null   int64
2   full_text           1623 non-null   object
3   lang                1462 non-null   object
4   quote_count         1519 non-null   float64
5   reply_count         1516 non-null   float64
6   retweet_count       1516 non-null   float64
7   username            1506 non-null   object
dtypes: float64(3), int64(1), object(4)
memory usage: 101.6+ KB
```

Gambar 4 Informasi Lengkap Data Frame

Langkah selanjutnya adalah mengecek keberadaan data yang kosong atau hilang pada dataset. Karena kehadiran nilai kosong dapat memengaruhi akurasi analisis, jadi pengecekan data kosong wajib dilakukan sebelum

pemrosesan data. Berikut gambar hasil pengecekan nilai kosong:

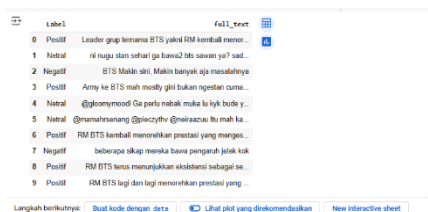


	0
created_at	0
favorite_count	0
Label	0
full_text	0
lang	161
quote_count	104
reply_count	107
retweet_count	107
username	117

dtype: int64

Gambar 5 Hasil pengecekan missing value

Data menunjukkan beberapa kolom dengan nilai yang hilang, seperti *lang* (161 nilai kosong) akibat *tweet* tanpa informasi bahasa, serta *quote_count*, *reply_count*, dan *retweet_count* (104–107 nilai kosong) terkait interaksi sosial. Kolom seperti *created_at*, *favorite_count*, dan *Label* memiliki data lengkap. Proses pembersihan dilakukan untuk menghapus kolom tidak relevan, seperti metadata, bahasa, tanggal, dan nama pengguna, guna memfokuskan analisis pada informasi yang mendukung hasil secara optimal.

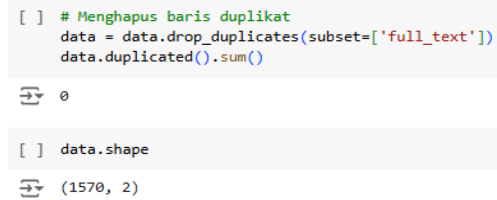


	Label	full_text
0	Positif	Leader grup ternama BTS yakni RM kembali mem...
1	Natal	ri ragu dan refat ga kawat klo server ya? sad...
2	Negatif	BTS Maaf ari. Makin banyak aja masalahnya
3	Positif	Army ke BTS mah mostly gtl bukan ngendon come...
4	Natal	@glennymood Ga perlu rebek muka lu lyk budy y...
5	Natal	@memahmawang @gricyyfr @nelrasca lu mah ka...
6	Positif	RM BTS kembali menorehkan prestasi yang menges...
7	Negatif	beberapa sikap mereka benar pengah? jlek kkk
8	Positif	RM BTS terus menunjukkan eksistensi sebagai se...
9	Positif	RM BTS lagi dan lagi menorehkan prestasi yang ...

Gambar 6 Data hasil Pembersihan kolom

Untuk memastikan bahwa setiap entri dalam dataset adalah unik, langkah selanjutnya adalah menghapus baris data yang memiliki nilai duplikat. Karena model cenderung belajar pola yang sama berulang kali, baris duplikat dapat menyebabkan bias dalam analisis. Ini mengurangi kemampuan model untuk generalisasi terhadap data baru. Identifikasi dilakukan dengan memeriksa kolom-kolom utama, seperti teks dan label, untuk memastikan bahwa data tersebut muncul dengan cara yang sama. Dataset keseluruhan, yang terdiri dari

1570 baris duplikat, menjadi lebih bersih, efektif, dan menggambarkan variasi data yang sebenarnya dengan lebih baik.



```
[ ] # Menghapus baris duplikat
data = data.drop_duplicates(subset=['full_text'])
data.duplicated().sum()

0

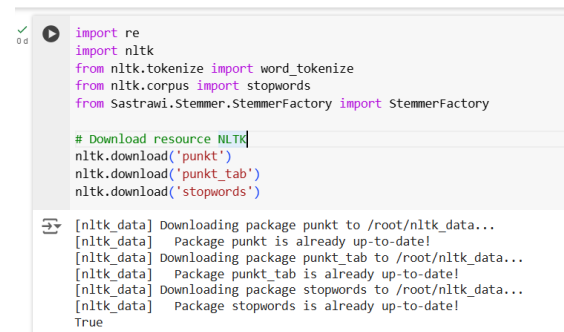
[ ] data.shape

(1570, 2)
```

Gambar 7 Hasil Pembersihan Baris Duplikat

4.1.2 Data Preprocessing

Tahap *preprocessing* adalah langkah awal yang penting dalam analisis data teks untuk memastikan bahwa data yang digunakan dalam model analitik atau *machine learning* berada dalam kondisi optimal. Proses ini bertujuan untuk membersihkan data dari elemen yang tidak relevan dan menyederhanakan teks sehingga lebih mudah diproses oleh algoritma. Gambar dibawah ini adalah beberapa *library* *Natural Language Toolkit* (NLTK) yang merupakan bagian dari tahap *preprocessing* data dalam analisis sentimen.



```
import re
import nltk
from nltk.tokenize import word_tokenize
from nltk.corpus import stopwords
from Sastrawi.Stemmer.StemmerFactory import StemmerFactory

# Download resource NLTK
nltk.download('punkt')
nltk.download('punkt_tab')
nltk.download('stopwords')
```

[nltk_data] Downloading package punkt to /root/nltk_data...

[nltk_data] Package punkt is already up-to-date!

[nltk_data] Downloading package punkt_tab to /root/nltk_data...

[nltk_data] Package punkt_tab is already up-to-date!

[nltk_data] Downloading package stopwords to /root/nltk_data...

[nltk_data] Package stopwords is already up-to-date!

True

Gambar 8 Library pada tahap Preprocessing

1. Cleansing

Cleansing data melibatkan penghapusan elemen-elemen yang tidak diperlukan seperti tanda baca, simbol, angka, dan karakter khusus yang tidak relevan dengan konteks analisis sentimen. Selain itu, *link*, *mention*, dan *hashtag* yang sering ditemukan dalam data media sosial juga dihapus untuk mengurangi gangguan dalam analisis. Pembersihan ini memastikan bahwa teks yang diproses hanya mengandung kata-kata yang memiliki arti penting dalam analisis. Hasilnya adalah teks yang lebih bersih

dan siap untuk diproses lebih lanjut. Gambar berikut merupakan hasil dari tahap *cleaning* :

```

Data Cleansing:
0 Leader grup ternama BTS yakni RM kembali menor...
1 ni nugu stan sehari ga bawa2 bts sawan ya? sadar...
2 BTS Makin sini, Makin banyak aja masalahnya
3 Army ke BTS mah mostly gini bukan ngestan cuma...
4 @gloomymoodl Ga perlu nebak muka lu kyk bude yng jijik kan ...

clean_text
0 Leader grup ternama BTS yakni RM kembali menor...
1 ni nugu stan sehari ga bawa bts sawan ya sadar...
2 BTS Makin sini Makin banyak aja masalahnya
3 Army ke BTS mah mostly gini bukan ngestan cuma...
4 Ga perlu nebak muka lu kyk bude yng jijik kan ...

```

Gambar 9 Data Hasil Cleansing

2. Case Folding

Pada tahap ini, semua teks diubah menjadi huruf kecil (*lowercase*) untuk memastikan konsistensi dalam pengolahan data. Langkah ini penting untuk menghindari perlakuan berbeda terhadap kata yang sama dengan bentuk huruf berbeda, seperti "Data" dan "data". Selain itu, case folding membantu menyederhanakan proses pencocokan kata dalam tahap tokenisasi dan analisis selanjutnya. Proses ini sangat esensial dalam menyatukan struktur data agar lebih seragam. Gambar berikut merupakan hasil dari tahap *Case Folding*:

```

Data Case folding:
0 Leader grup ternama BTS yakni RM kembali menor...
1 ni nugu stan sehari ga bawa bts sawan ya sadar...
2 BTS Makin sini Makin banyak aja masalahnya
3 Army ke BTS mah mostly gini bukan ngestan cuma...
4 Ga perlu nebak muka lu kyk bude yng jijik kan ...

clean_text
0 Leader grup ternama BTS yakni RM kembali menor...
1 ni nugu stan sehari ga bawa bts sawan ya sadar...
2 BTS Makin sini Makin banyak aja masalahnya
3 Army ke BTS mah mostly gini bukan ngestan cuma...
4 Ga perlu nebak muka lu kyk bude yng jijik kan ...

case folding_text
0 leader grup ternama bts yakni rm kembali menor...
1 ni nugu stan sehari ga bawa bts sawan ya sadar...
2 bts makin sini makin banyak aja masalahnya
3 army ke bts mah mostly gini bukan ngestan cuma...
4 ga perlu nebak muka lu kyk bude yng jijik kan ...

```

Gambar 10 Data Hasil Case Folding

3. Tokenizing

Tokenisasi dilakukan untuk memecah teks menjadi unit-unit yang lebih kecil, biasanya berupa kata atau frasa, yang disebut sebagai token. Proses ini membantu dalam mengidentifikasi setiap elemen teks secara individu sehingga dapat dianalisis lebih mendalam. Proses ini menjadi dasar bagi langkah-langkah berikutnya dalam *preprocessing*. Gambar berikut merupakan data hasil dari tahap *tokenizing* :

```

Data Tokenizing:
0 leader grup ternama bts yakni rm kembali menor...
1 ni nugu stan sehari ga bawa bts sawan ya sadar...
2 bts makin sini makin banyak aja masalahnya
3 army ke bts mah mostly gini bukan ngestan cuma...
4 ga perlu nebak muka lu kyk bude yng jijik kan ...

tokens
0 [leader, grup, ternama, bts, yakni, rm, kembal...
1 [ni, nugu, stan, sehari, ga, bawa, bts, sawan,...
2 [bts, makin, sini, makin, banyak, aja, masalah...
3 [army, ke, bts, mah, mostly, gini, bukan, nges...
4 [ga, perlu, nebak, muka, lu, kyk, bude, yng, j...

```

Gambar 11 Data Hasil Tokenizing

4. Normalization

Normalisasi dilakukan untuk menyamakan bentuk kata yang memiliki arti serupa tetapi ditulis secara berbeda. Kata-kata seperti "tdk" diubah menjadi "tidak", atau "gk" menjadi "tidak", agar analisis lebih konsisten. Proses ini sering melibatkan penggunaan kamus normalisasi yang telah dibuat sebelumnya. Normalisasi penting untuk mengurangi variasi bahasa yang tidak relevan dan meningkatkan akurasi pada tahap analisis. Gambar berikut merupakan data hasil *normalization*:

```

Data Normalization:
0 [leader, grup, ternama, bts, yakni, rm, kembal...
1 [ni, nugu, stan, sehari, ga, bawa, bts, sawan,...
2 [bts, makin, sini, makin, banyak, aja, masalah...
3 [army, ke, bts, mah, mostly, gini, bukan, nges...
4 [ga, perlu, nebak, muka, lu, kyk, bude, yng, j...

tokens
0 [leader, grup, ternama, bts, yakni, rm, kembal...
1 [ni, nugu, stan, sehari, ga, bawa, bts, sawan,...
2 [bts, makin, sini, makin, banyak, aja, masalah...
3 [army, ke, bts, mah, mostly, gini, bukan, nges...
4 [ga, perlu, nebak, muka, lu, kyk, bude, yng, j...

normalized_tokens
0 [leader, grup, ternama, bts, yakni, rm, kembal...
1 [ni, nugu, stan, sehari, tidak, bawa, bts, sa...
2 [bts, makin, sini, makin, banyak, saja, masala...
3 [army, ke, bts, cuma, mostly, gini, bukan, sta...
4 [tidak, perlu, menebak, muka, anda, seperti, b...

```

Gambar 12 Data Hasil Normalization

5. Stop word Removal

Penghapusan *stop word* adalah langkah untuk menghilangkan kata-kata umum yang sering muncul tetapi tidak memiliki nilai informasi, seperti "dan", "atau", "di", "ke", dan sebagainya. Daftar stop word biasanya diambil dari pustaka yang tersedia atau disesuaikan dengan kebutuhan analisis. Dengan menghapus kata-kata ini, fokus analisis dapat diarahkan pada kata-kata yang lebih penting dan relevan dengan sentimen yang ingin diidentifikasi. Gambar berikut merupakan data hasil *Stop Word Removal* :

```

Data Stopword Removal:
                                normalized_tokens \
0 [leader, grup, ternama, bts, yakni, rm, kembal...
1 [ini, nugu, stan, sehari, tidak, bawa, bts, sa...
2 [bts, makin, sini, makin, banyak, saja, masala...
3 [army, ke, bts, cuma, mostly, gini, bukan, sta...
4 [tidak, perlu, menebak, muka, anda, seperti, b...

                                tokens_no_stopwords
0 [leader, grup, ternama, bts, rm, menorehkan, p...
1 [nugu, stan, sehari, bawa, bts, sadar, idolamu...
2 [bts]
3 [army, bts, mostly, gini, stan, visualnya, bts...
4 [menebak, muka, bibi, jijik, iri, idola]

```

Gambar 13 Data Hasil Stopword

6. Stemming

Tahap terakhir dalam *preprocessing* adalah *stemming*, yaitu proses mengubah kata-kata menjadi bentuk dasar atau akarnya. Misalnya, kata "berlari", "lari-lari", dan "pelari" akan dikembalikan menjadi "lari". Stemming dilakukan untuk menyatukan variasi kata dengan akar yang sama agar tidak dianggap sebagai entitas yang berbeda. Teknik ini meningkatkan efisiensi dan akurasi dalam analisis sentimen dengan mengurangi redundansi dalam representasi data. Gambar berikut merupakan data hasil proses *stemming* :

```

Data Stemming:
                                tokens_no_stopwords \
0 [leader, grup, ternama, bts, rm, menorehkan, p...
1 [nugu, stan, sehari, bawa, bts, sadar, idolamu...
2 [bts]
3 [army, bts, mostly, gini, stan, visualnya, bts...
4 [menebak, muka, bibi, jijik, iri, idola]
...
1618 [sifatnya, manusia, gitu, melupakan, kebaikan...
1619 [orang, mengaitkan, angka, kesialan, keburukan...
1620 [bts, mengajarkan, army, membalas, keburukan, ...
1621 [terharu, lihat, kenal, kebaikan, keburukan, a...
1622 [bts, mencontohkan, fans nya, keburukan, dibal...

                                stemmed_text
0 leader grup nama bts rm toreh prestasi kesan r...
1 nugu stan hari bawa bts sadar idola prestasi k...
2 bts
3 army bts mostly gin stan visual bts jual music...
4 tebak muka bibi jijik iri idola
...
1618 sifat manusia gitu lupa baik keburu menu pintu...
1619 orang kait angka sial keburu lihat jimin bts n...
1620 bts ajar army balas keburu keburu prayforkanju...
1621 haru lihat kenal baik keburu anti army anti bt...
1622 bts contoh fans nya keburu balas baik

```

[1570 rows x 2 columns]

Gambar 14 Data Hasil Stemming

Setelah proses *stemming*, langkah berikutnya adalah mempersiapkan data yang telah dibersihkan dan melibatkan pemanggilan kolom-kolom yang relevan untuk analisis selanjutnya. Dalam hal ini, kolom *stemmed_text* berisi teks yang telah diproses (melalui normalisasi, penghapusan *stopwords*,

dan *stemming*), sementara kolom *Label* berisi kategori sentimen. Berikut gambar hasil pemanggilan kolom yang akan diproses lebih lanjut:

```

                                stemmed_text Label
0 leader grup nama bts rm toreh prestasi kesan r... Positif
1 nugu stan hari bawa bts sadar idola prestasi k... Netral
2 bts Negatif
3 army bts mostly gin stan visual jual music pre... Positif
4 tebak muka jijik iri idola Netral
...
1618 sifat manusia gitu lupa baik keburu menu pintu... Positif
1619 orang kait angka sial keburu lihat jimin bts n... Positif
1620 bts ajar army balas keburu prayforkanju... Positif
1621 haru lihat kenal baik keburu anti army bts mik... Netral
1622 bts contoh fans nya keburu balas baik Positif

```

[1551 rows x 2 columns]

Gambar 15 Data Hasil Pembersihan

Sebelum melanjutkan ke tahap berikutnya, peneliti melakukan pemeriksaan ulang terhadap dataset untuk memastikan tidak adanya kesalahan dalam proses pra-pemrosesan, seperti label yang hilang. Pemeriksaan ini juga memberikan gambaran awal mengenai komposisi dataset, yang menjadi dasar untuk langkah visualisasi selanjutnya. Langkah ini bertujuan untuk meningkatkan transparansi dalam proses penyediaan data sekaligus membantu memahami karakteristik dataset secara mendalam sebelum digunakan pada model pembelajaran mesin. Informasi ini sangat penting untuk memastikan bahwa dataset dalam kondisi seimbang dan tidak memiliki ketidakseimbangan data yang signifikan, yang dapat memengaruhi kinerja model secara keseluruhan.

```

Label
Positif    779
Netral     595
Negatif    177
Name: count, dtype: int64

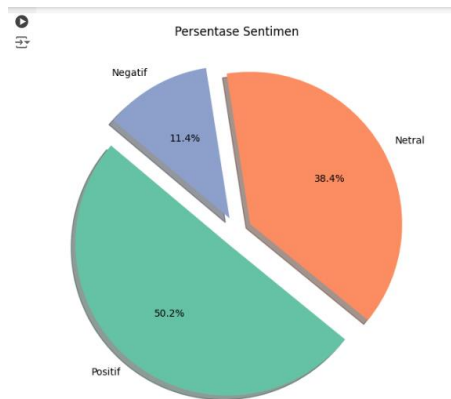
```

Gambar 16 Tampilan jumlah sentimen pada masing-masing Label

1. Visualisasi

Berikut adalah penggunaan diagram pie untuk memvisualisasikan distribusi sentimen dari dataset berdasarkan kolom *Label*. Tiga kategori sentimen yang ditampilkan adalah *Positif*, *Netral*, dan *Negatif*. Diagram pie memberikan representasi grafis dalam bentuk lingkaran, di mana setiap bagian

melambangkan proporsi kategori tertentu terhadap keseluruhan data.

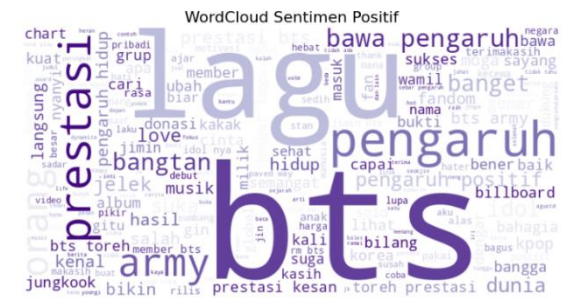


Gambar 17 Diagram Pie Presentase sentimen

Gambar 16 dan 17 menunjukkan distribusi sentimen: label positif mendominasi dengan 779 data (50,2%), netral 595 data (38,4%), dan negatif 177 data (11,4%). Visualisasi ini membantu memahami proporsi sentimen dalam dataset, di mana sentimen positif menjadi yang terbanyak, sementara netral cukup signifikan tetapi lebih kecil. Sentimen negatif memiliki jumlah paling sedikit, mencerminkan opini negatif terhadap BTS yang relatif rendah dibandingkan opini netral dan positif.

Visualisasi Word Cloud digunakan untuk merepresentasikan kata-kata yang sering muncul dalam data teks berdasarkan sentimen: positif, negatif, dan netral. Prosesnya dimulai dengan mengelompokkan teks berdasarkan label sentimen, lalu menggabungkan teks dalam setiap kategori menjadi satu string. Dengan library WordCloud Python, frekuensi kata dihitung dan divisualisasikan secara proporsional. Warna ungu digunakan untuk

sentimen positif, biru untuk negatif, dan hijau untuk netral.



Gambar 18 WordCloud Sentimen Positif

WordCloud sentimen positif menunjukkan kata-kata dominan seperti "bts", "lagu", "prestasi", "pengaruh", dan "army", yang menggambarkan opini positif terhadap BTS. Kata "prestasi" dan "pengaruh" mencerminkan penghargaan atas kesuksesan grup, sementara "army" menunjukkan dukungan penggemar. Kata-kata seperti "bangga", "sukses", dan "dunia" mengilustrasikan apresiasi terhadap dampak global BTS. Visualisasi ini memberikan wawasan tentang persepsi pengguna dalam sentimen positif.



Gambar 19 WordCloud Sentimen Negatif

WordCloud sentimen negatif menunjukkan kata-kata seperti "kecewa", "jelek", dan "salah", yang mencerminkan ketidakpuasan atau kritik. Kata "army" dan "fandom" juga sering muncul, menunjukkan opini negatif yang mungkin terkait dengan komunitas penggemar. Kata "prestasi" dalam konteks ini kemungkinan digunakan dalam kritik. Visualisasi ini menggambarkan tema utama opini negatif dan memberikan wawasan terkait sentimen negatif terhadap BTS.



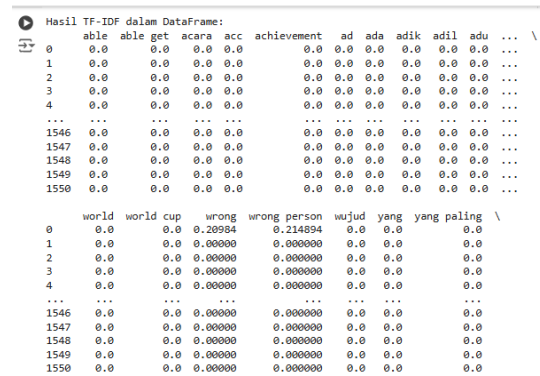
Gambar 20 WordCloud Sentimen Netral

WordCloud sentimen netral menampilkan kata-kata seperti "prestasi", "army", dan "pengaruh" sebagai yang dominan, menunjukkan fokus pada pencapaian dan dampak BTS tanpa ekspresi emosional kuat. Kata seperti "suka", "bilang", dan "fandom" sering muncul, mencerminkan percakapan deskriptif atau informatif. Meski "kecewa" dan "jelek" juga terlihat, konteksnya kemungkinan lebih netral atau bercampur. Visualisasi ini menggambarkan pandangan yang lebih objektif terhadap BTS.

4.1.3 Transformasi Data

1. TF-IDF

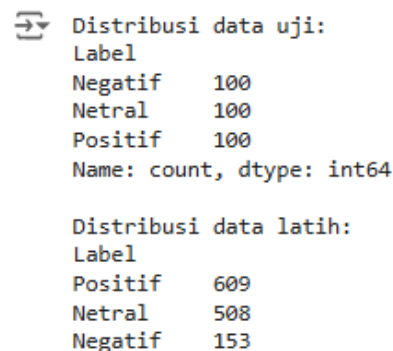
Metode TF-IDF digunakan untuk mengubah teks yang telah dibersihkan menjadi fitur numerik. Teknik ini menghitung pentingnya kata berdasarkan frekuensi kemunculannya dalam satu dokumen dan penyebarannya di seluruh dokumen. Dengan `ngram_range (1, 2)`, dihasilkan pola unigram dan bigram hingga 5000 fitur dengan frekuensi minimal 3. Nilai TF-IDF yang memenuhi kriteria disajikan dalam bentuk `DataFrame`, mempermudah interpretasi data dan membantu algoritma pembelajaran mesin mengenali pola teks penting.



Gambar 21 Hasil Transformasi TF_IDF

2. Split Data

Data dibagi menjadi dua bagian: data latih dan data uji. Untuk menghindari ketidakseimbangan, pembagian dilakukan secara stratifikasi, dengan jumlah data uji yang sama untuk setiap label. Sebanyak 100 data dari masing-masing label (Negatif, Netral, Positif) dipilih secara acak menggunakan metode groupby dan sampel, memastikan distribusi label yang adil. Data yang tersisa digunakan sebagai data latih. Pendekatan ini menjamin evaluasi model yang representatif terhadap dataset keseluruhan.



Gambar 22 Data Hasil Split Data

Memisahkan data uji sebelum menerapkan SMOTE adalah pendekatan yang tepat untuk menghindari data leakage, sehingga evaluasi model lebih akurat dan mencerminkan kinerjanya pada data baru. Data latih yang tidak seimbang, seperti distribusi awal (609 positif, 508 netral, 153 negatif), kemudian diimbangi menggunakan SMOTE, agar model tidak bias terhadap kelas mayoritas. Transformasi TF-IDF dilakukan terlebih dahulu agar fitur yang digunakan konsisten dan relevan. Pendekatan

ini memungkinkan model mengenali pola kelas minoritas dengan lebih baik, sementara evaluasi data uji tetap objektif.

```

Distribusi data latih sebelum SMOTE:
Label
Positif      609
Netral       508
Negatif      153
Name: count, dtype: int64

Distribusi data latih setelah SMOTE:
Label
Netral       609
Positif      609
Negatif      609
Name: count, dtype: int64

```

Gambar 23 Data Hasil Proses Smote

4.1.4 Text Mining

1. Pemodelan Naïve Bayes

Naive Bayes adalah algoritma klasifikasi probabilistik yang mengasumsikan independensi fitur dan menggunakan Teorema Bayes untuk menghitung kemungkinan data masuk ke dalam kelas tertentu. Varian yang sering digunakan untuk data teks adalah *Multinomial Naive Bayes* (MNB), yang menghitung frekuensi kata dalam dokumen. Penelitian ini menggunakan MNB untuk mengklasifikasikan sentimen positif, negatif, dan netral dari data tweet yang telah diproses. Eksperimen dilakukan dengan mencoba berbagai nilai parameter α untuk melakukan smoothing dan mencegah zero probability pada fitur tertentu. Berikut hasil eksperimen untuk setiap nilai α :

1) Model MultinomialNB(alpha=0.1)

Pada parameter α 0.1, model Multinomial Naive Bayes mencapai akurasi 78,33%, dengan presisi rata-rata 79,25%, recall 78,33%, dan *F1-Score* 78,49%. Model menunjukkan kinerja sangat baik pada kelas negatif, dengan recall sebesar 81%, serta performa yang seimbang pada kelas netral (recall 81%) dan positif (recall 73%). Berikut hasil pemodelan MNB($\alpha=0.1$):

```

Akurasi: 0.7833333333333333
Presisi: 0.7925443796675696
Recall: 0.7833333333333333
F1 Score: 0.7849947589612624

Classification Report:
              precision    recall  f1-score   support

Negatif      0.87         0.81         0.84         100
Netral       0.69         0.81         0.74         100
Positif      0.82         0.73         0.77         100

accuracy      0.78         0.78         0.78         300
macro avg     0.79         0.78         0.78         300
weighted avg  0.79         0.78         0.78         300

```

Gambar 24 Hasil Pemodelan MNB (Alpha=0.1)

2) Model MultinomialNB(alpha=0.5)

Pada pemodelan dengan α 0.5, akurasi mencapai 76,67%, dengan presisi rata-rata 77,18%, recall 76,67%, dan *F1-Score* 76,74%. Model tampil baik pada kelas negatif (recall 83%), namun performa pada kelas netral dan positif sedikit menurun, dengan recall masing-masing 76% dan 71%. Berikut hasil pemodelan MNB($\alpha=0.5$):

```

Akurasi: 0.7666666666666667
Presisi: 0.7717965367965367
Recall: 0.7666666666666667
F1 Score: 0.7674334270038807

Classification Report:
              precision    recall  f1-score   support

Negatif      0.83         0.83         0.83         100
Netral       0.68         0.76         0.72         100
Positif      0.81         0.71         0.76         100

accuracy      0.77         0.77         0.77         300
macro avg     0.77         0.77         0.77         300
weighted avg  0.77         0.77         0.77         300

```

Gambar 25 Hasil Pemodelan MNB (Alpha=0.5)

3) Model MultinomialNB(alpha=1.0)

Pada α 1.0, akurasi mencapai 76,00%, dengan presisi 76,43%, recall 76,00%, dan *F1-Score* 76,03%. Model tampil baik pada kelas negatif (recall 84%), tetapi performa pada kelas netral dan positif lebih rendah, masing-masing 74% dan 70%. Berikut hasil pemodelan MNB($\alpha=1.0$):

Akurasi: 0.76
 Presisi: 0.764286318153074
 Recall: 0.76
 F1 Score: 0.7603370710875782

Classification Report:				
	precision	recall	f1-score	support
Negatif	0.82	0.84	0.83	100
Netral	0.67	0.74	0.70	100
Positif	0.80	0.70	0.75	100
accuracy			0.76	300
macro avg	0.76	0.76	0.76	300
weighted avg	0.76	0.76	0.76	300

Gambar 26 Hasil Pemodelan MNB (Alpha=1.0)

4) Model MultinomialNB(alpha=1.5)

Pada α 1.5, akurasi yang dicapai adalah 75,00%, dengan presisi 75,35%, *recall* 75,00%, dan *F1-Score* 74,98%. Model menunjukkan kinerja stabil pada kelas negatif (*recall* 84%), namun performa pada kelas netral dan positif menurun, masing-masing 73% dan 68%. Berikut hasil pemodelan MNB(α =1.5):

Akurasi: 0.75
 Presisi: 0.7534741483536022
 Recall: 0.75
 F1 Score: 0.7497531947074371

Classification Report:				
	precision	recall	f1-score	support
Negatif	0.80	0.84	0.82	100
Netral	0.67	0.73	0.70	100
Positif	0.79	0.68	0.73	100
accuracy			0.75	300
macro avg	0.75	0.75	0.75	300
weighted avg	0.75	0.75	0.75	300

Gambar 27 Hasil Pemodelan MNB (Alpha=1.5)

5) Model MultinomialNB(alpha=2.0)

Pada α 2.0, akurasi yang dicapai adalah 75,00%, dengan presisi 75,35%, *recall* 75,00%, dan *F1-Score* 74,98%. Kinerjanya serupa dengan α 1.5, dengan *recall* 84% pada kelas negatif, namun lebih rendah pada kelas netral (73%) dan positif (68%). Berikut hasil pemodelan MNB(α =2.0):

Akurasi: 0.75
 Presisi: 0.7534741483536022
 Recall: 0.75
 F1 Score: 0.7497531947074371

Classification Report:				
	precision	recall	f1-score	support
Negatif	0.80	0.84	0.82	100
Netral	0.67	0.73	0.70	100
Positif	0.79	0.68	0.73	100
accuracy			0.75	300
macro avg	0.75	0.75	0.75	300
weighted avg	0.75	0.75	0.75	300

Gambar 28 Hasil Pemodelan MNB (Alpha=1.5)

4.1.5 Evaluation

Hasil penelitian menunjukkan bahwa α terbaik untuk model *Multinomial Naive Bayes* adalah 0.1, dengan akurasi 78,33%. Selain akurasi, nilai *F1-Score* dan *recall* pada α 0.1 juga lebih tinggi dibandingkan dengan α lainnya. α 0.1 memungkinkan model menangkap pola spesifik tanpa terlalu banyak smoothing, menghasilkan kinerja yang lebih baik. Sebaliknya, α lebih tinggi seperti 1.5 dan 2.0 mengurangi sensitivitas model, menjadikannya kurang optimal. Oleh karena itu, α 0.1 dipilih sebagai parameter terbaik, memberikan keseimbangan optimal antara sensitivitas dan kemampuan generalisasi. Berikut gambar hasil pemilihan parameter terbaik, dan confusion matrixnya:

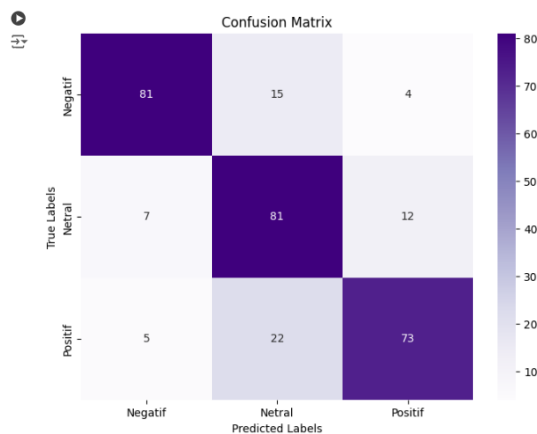
Accuracy Test after tuning: 0.78

Confusion Matrix:

```
[[81 15  4]
 [ 7 81 12]
 [ 5 22 73]]
```

Classification Report:				
	precision	recall	f1-score	support
Negatif	0.87	0.81	0.84	100
Netral	0.69	0.81	0.74	100
Positif	0.82	0.73	0.77	100
accuracy			0.78	300
macro avg	0.79	0.78	0.78	300
weighted avg	0.79	0.78	0.78	300

Gambar 29 Hasil Pemilihan Parameter Terbaik



Gambar 30 Visualisasi Confusion Matrix

Matriks evaluasi terdiri dari *True Positive* (TP), *False Positive* (FP), dan *False Negative* (FN). Pada model *Naive Bayes* dengan α 0.1, akurasi mencapai 78,33%, menunjukkan 78,33% prediksi benar. Model unggul dalam mendeteksi kelas negatif dengan presisi 87% dan *recall* 81%, meskipun terdapat kesalahan, seperti 15 sampel negatif diprediksi netral dan 4 sebagai positif. Untuk kelas netral, *recall* mencapai 81%, namun presisi lebih rendah di 69%, dengan beberapa kesalahan, termasuk 7 sampel netral diprediksi negatif dan 12 sebagai positif. Pada kelas positif, presisi mencapai 82%, tetapi *recall* hanya 73%, menunjukkan beberapa sampel sulit diidentifikasi, seperti 5 sampel positif diprediksi negatif dan 22 sebagai netral. Hasil ini menunjukkan model cukup andal dalam klasifikasi multikelas.

4.2 Pembahasan

4.2.1 Penafsiran dan Analisis Hasil Penelitian

Penelitian ini bertujuan untuk menganalisis sentimen pengguna *Twitter* terhadap grup musik BTS dengan menggunakan algoritma *Naive Bayes*. Model klasifikasi sentimen yang dihasilkan menunjukkan tingkat akurasi sebesar 78,33%, yang dianggap cukup efektif dalam mengkategorikan opini menjadi sentimen positif, negatif, dan netral.

Berdasarkan hasil tersebut, beberapa poin penting dapat diinterpretasikan:

- 1) **Efektivitas *Naive Bayes*:** Algoritma ini terbukti mampu menangani data tidak

terstruktur dari *Twitter*, meskipun akurasi yang diperoleh (78,33%) lebih rendah dibandingkan hasil penelitian Harun & Putri Ananda (2021) [8] yang mencapai 100% dalam konteks berbeda. Hal ini menunjukkan bahwa performa *Naive Bayes* dapat dipengaruhi oleh karakteristik data, seperti perbedaan topik, bahasa, dan kompleksitas dataset.

- 2) **Distribusi Sentimen:** Temuan ini memberikan gambaran tentang opini publik terhadap BTS di media sosial. Sentimen yang dominan dapat mencerminkan popularitas dan reputasi grup musik ini di kalangan pengguna *Twitter*.
- 3) **Peran Teknik Pra-pemrosesan:** Tahapan seperti pembersihan data, normalisasi, dan pemberian bobot menggunakan TF-IDF berperan penting dalam meningkatkan akurasi model. Namun, hasil akurasi menunjukkan bahwa metode ini masih memiliki keterbatasan dalam menangkap konteks linguistik campuran (bahasa Indonesia dan Inggris).

4.2.2 Keterkaitan dengan Literatur Terdahulu

Hasil penelitian ini mendukung dan menambahkan wawasan baru terhadap literatur sebelumnya:

- 1) **Dukungan terhadap *Naive Bayes*:** Penelitian ini sejalan dengan studi Harun & Putri Ananda (2021)[8], yang menyoroti keunggulan algoritma *Naive Bayes* dalam analisis sentimen. Meskipun akurasi penelitian ini lebih rendah, hasil tetap menunjukkan bahwa *Naive Bayes* merupakan algoritma yang efektif untuk dataset besar dan tidak terstruktur.
- 2) **Perbandingan dengan Metode Lain:** Penelitian sebelumnya oleh Rina Noviana & Isram Rasal (2023)[12] menunjukkan bahwa algoritma *Support Vector Machine* memiliki akurasi lebih tinggi dibandingkan *Naive Bayes*. Hal ini menegaskan perlunya

mengeksplorasi metode yang lebih canggih untuk meningkatkan performa analisis sentimen pada data Twitter.

- 3) **Kontribusi dalam Analisis Data Sosial Media:** Safitri et al. (2023)[7] menggunakan data *Twitter* untuk mengevaluasi sentimen publik terhadap BTS sebagai *brand ambassador*. Penelitian ini melanjutkan tradisi tersebut dengan pendekatan berbeda, menggunakan algoritma *Naïve Bayes* dan fokus pada konteks linguistik lokal.

5. KESIMPULAN

1) Kesimpulan

Algoritma *Naïve Bayes* efektif mengklasifikasikan sentimen Twitter tentang BTS menjadi positif, negatif, dan netral, dengan akurasi 78,33%. Proses ini melibatkan pra-pemrosesan dan TF-IDF, memberikan wawasan pola reaksi publik selama hiatus BTS serta berkontribusi pada analisis sentimen bahasa Indonesia.

2) Saran

- a) Optimalkan parameter *Naïve Bayes*, seperti distribusi probabilitas yang digunakan, untuk meningkatkan performa model.
- b) Kembangkan kamus sentimen atau metode khusus untuk menangani data bahasa campuran Indonesia dan Inggris.
- c) Gunakan dataset yang lebih besar dan tambahkan data dari platform lain untuk memperluas cakupan analisis.
- d) Eksplorasi algoritma lebih kompleks, seperti *SVM*, *Random Forest*, atau model berbasis *deep learning*, untuk meningkatkan akurasi.
- e) Terapkan metode ini pada topik lain, seperti analisis sentimen produk atau isu sosial, untuk memperluas manfaatnya.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada pihak-pihak terkait yang telah memberi dukungan terhadap penelitian ini.

DAFTAR PUSTAKA

- [1] L.P.K.F. Aryawan, I.G. Budasi, and D.P. Ramendra, "the Deixis Used in the Song Lyrics Performed By the Most Popular Boy Group Bts," *J. Pendidik. Bhs. Ingg. Indones.*, vol. 10, no. 1, pp. 30–39, 2022, doi: 10.23887/jpbi.v10i1.796.
- [2] N. Fajriah, H. M. Syam, N. Susilawati, R. Rosemary, and Z. Azman, "BTS K-Pop consumption and self-identity in the k-pop army community of Banda Aceh," *J. Geuthèë Penelit. Multidisiplin*, vol. 7, no. 2, p. 105, 2024, doi: 10.52626/jg.v7i2.345.
- [3] D. R. Febryanti, Z. K. Mahmud, S. V. Putri, F. M. Ovalia, and Y. Sekarwangi, "Tipologi Hate Speech di Twitter Terkait Kebijakan Pemerintah Selama Pandemi COVID-19," *J. Komun. Glob.*, vol. 11, no. 2, pp. 274–299, 2022, doi: 10.24815/jkg.v11i2.26733.
- [4] Z. A. Sriyanti, D. S. Y. Kartika, and A. R. E. Najaf, "Implementasi Model Bert Pada Analisis Sentimen Pengguna Twitter Terhadap Aksi Boikot Produk Israel," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, pp. 2335–2342, 2024, doi: 10.23960/jitet.v12i3.4743.
- [5] M. F. Y. Herjanto and C. Carudin, "Analisis Sentimen Ulasan Pengguna Aplikasi Sirekap Pada Play Store Menggunakan Algoritma Random Forest Classifier," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 2, pp. 1204–1210, 2024, doi: 10.23960/jitet.v12i2.4192.
- [6] 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.
- [7] T. Safitri, Y. Umaidah, and I. Maulana, "Analisis Sentimen Pengguna Twitter Terhadap Grup Musik BTS Menggunakan Algoritma Support Vector Machine," *J. Appl. ...*, 2023, [Online]. Available: <https://jurnal.polibatam.ac.id/index.php/JAIC/article/view/5039>
- [8] A. Harun and D. Putri Ananda, "Analisa Sentimen Opini Publik Tentang Vaksinasi Covid-19 di Indonesia Menggunakan Naïve bayes dan Decission Tree," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 1, no. 1, pp. 58–64, 2021, doi: 10.57152/malcom.v1i1.63.
- [9] D. F. Zhafira, B. Rahayudi, and I. Indriati, "Analisis Sentimen Kebijakan Kampus Merdeka Menggunakan Naive Bayes dan Pembobotan TF-IDF Berdasarkan Komentar pada Youtube," *J. Sist. Informasi, Teknol. Informasi, dan Edukasi Sist. Inf.*, vol. 2, no. 1,

- pp. 55–63, 2021, doi: 10.25126/justsi.v2i1.24.
- [10] Muhammad Daffa Al Fahreza, Ardytha Luthfiarta, Muhammad Rafid, and Michael Indrawan, “Analisis Sentimen: Pengaruh Jam Kerja Terhadap Kesehatan Mental Generasi Z,” *J. Appl. Comput. Sci. Technol.*, vol. 5, no. 1, pp. 16–25, 2024, doi: 10.52158/jacost.v5i1.715.
- [11] I. Markoulidakis, I. Rallis, I. Georgoulas, G. Kopsiaftis, A. Doulamis, and N. Doulamis, “A Machine Learning Based Classification Method for Customer Experience Survey Analysis,” *Technologies*, vol. 8, no. 4, 2020, doi: 10.3390/technologies8040076.
- [12] Rina Noviana and Isram Rasal, “Penerapan Algoritma Naive Bayes Dan Svm Untuk Analisis Sentimen Boy Band Bts Pada Media Sosial Twitter,” *J. Tek. dan Sci.*, vol. 2, no. 2, pp. 51–60, 2023, doi: 10.56127/jts.v2i2.791.