

Analisis Sentimen Ulasan Pengguna terhadap Ekspedisi Layanan Logistik dengan Metode Naïve Bayes

1st Wahyu Budi Prayogo

Sistem Informasi Universitas Telkom
Purwokerto

Universitas Telkom Purwokerto
Purwokerto, Indonesia

wahyubp@student.telkomuniversity.ac.id

2nd Sena Wijayanto

Sistem Informasi Universitas Telkom
Purwokerto

Universitas Telkom Purwokerto
Purwokerto, Indonesia

senawijayanto@telkomuniversity.ac.id

3rd Mahazam Afrad

Sistem Informasi Universitas Telkom
Purwokerto

Universitas Telkom Purwokerto
Purwokerto, Indonesia

mahazama@telkomuniversity.ac.id

Abstrak — Perkembangan teknologi informasi memungkinkan pengguna menyampaikan ulasan melalui platform seperti *google maps* dengan memberikan informasi penting seperti penilaian kualitas dan pengalaman pengguna layanan. Penelitian ini berfokus pada pengelolaan ulasan pelanggan di JNE Purwokerto yang belum secara efektif dimanfaatkan untuk keperluan dalam evaluasi layanan. Dengan menggunakan pendekatan *Machine Learning*, penelitian ini bertujuan untuk mengklasifikasikan sentimen ulasan negatif dan positif pengguna sebagai bahan evaluasi bagi penyedia layanan. Tahapan penelitian dimulai dengan pengambilan data ulasan yang terdiri dari 968 ulasan dengan rincian 111 (positif) dan 857 (negatif) yang diperoleh melalui kode *Python* dan *API*. Lalu *preprocessing data*; *case folding*, *cleaning*, tokenisasi, *stopword removal*, dan *stemming*. Ekstraksi fitur dengan metode *Bag of Words* untuk mengubah data menjadi nilai numerik. *Splitting data* 80:20 menghasilkan 774 data latih dan 194 data uji. Penggunaan teknik *Synthetic Minority Over-sampling Technique* terhadap ketidakseimbangan data latih 89 (positif) dan 685 (negatif) menjadi data yang seimbang. Algoritma *Multinomial Naive Bayes* untuk pengujian. *Confusion matrix* untuk mengevaluasi performa metode algoritma *Naive Bayes*. Hasil dari proses evaluasi memperoleh nilai akurasi sebesar 89%, nilai presisi 93%, *recall* 94%, dan *f1-score* 93% (negatif). Sedangkan, nilai presisi 58%, *recall* 54%, dan *f1-score* 56% (positif). Hasil analisis menunjukkan ketidakpuasan layanan dengan tingginya sentimen negatif pada ulasan pelanggan.

Kata kunci — analisis sentimen, naive bayes, pembelajaran mesin, ulasan

I. PENDAHULUAN

Di era sekarang perkembangan teknologi adalah faktor yang sangat berpengaruh dalam kehidupan sehari-hari[1]. Perkembangan teknologi informasi membawa dampak besar dalam berbagai aspek kehidupan termasuk dalam sektor sosial[2]. Salah satu yang memainkan peran penting dalam memfasilitasi interaksi antara pengunjung dan penyedia layanan adalah media internet, terutama *google maps*, yang

menjadi tempat bagi pengunjung untuk memberikan ulasan atau *review* mengenai berbagai layanan suatu tempat[3].

Google maps telah menjadi sebuah platform yang sangat penting dalam membantu pengguna dalam menemukan dan menilai berbagai layanan, termasuk layanan pengiriman barang yang disediakan oleh Jalur Nugraha Ekakurir[4]. Melalui ulasan dan *review* yang ditinggalkan oleh pengguna, *google maps* memberikan kesempatan bagi pengguna untuk berbagi pengalaman ke pengguna lain dengan layanan yang diberikan. Hal ini menciptakan sebuah sumber daya informasi yang berharga bagi penyedia layanan seperti JNE Sub Agen Purwokerto, karena dapat memberikan wawasan langsung dari perspektif pelanggan mengenai kualitas layanan dan pengalaman pengguna[5]. Memahami dan mengelola data ulasan di *google maps* secara efektif dapat membuka peluang bagi JNE Sub Agen Purwokerto untuk mengidentifikasi kualitas layanan pengiriman barang. Dengan menganalisis ulasan dan *review* secara menyeluruh diharapkan JNE Sub Agen Purwokerto dapat mengambil tindakan yang tepat untuk meningkatkan kepuasan pengguna, merespons keluhan atau masukan yang diberikan, serta memperbaiki reputasi layanan. Selain itu, dengan menggunakan pendekatan yang terstruktur dan teknologi yang tepat, data ulasan di *google maps* juga dapat digunakan untuk memprediksi tren dan menggambarkan pengalaman pengguna terhadap layanan yang diberikan[6]. Namun data ulasan tersebut belum dimanfaatkan secara baik oleh pengelola manajemen layanan pengiriman JNE Sub Agen Purwokerto. Kurangnya perhatian terhadap data ulasan ini menyebabkan potensi informasi berharga terabaikan. Data ulasan yang belum dimanfaatkan dapat memberikan wawasan mendalam terkait kelemahan layanan pengiriman barang oleh JNE Sub Agen Purwokerto.

Penggunaan metode *machine learning* dalam analisis sentimen berdasarkan opini yang dihasilkan oleh pengguna menjadi relevan dan penting dalam penelitian ini. Metode *machine learning* yang dipakai yaitu metode *naive bayes classifier* (NBC). Metode *naive bayes classifier* model *multinomial naive bayes* (MNB) merupakan algoritma

klasifikasi yang berbasis *teorema bayes* dan mudah untuk digunakan dalam pengklasifikasian dokumen teks ke dalam kategori sentimen positif dan negatif terutama dengan dataset yang banyak[7]. Hasil dari pengumpulan data ulasan pengguna tersebut akan menjadi bahan evaluasi untuk mengidentifikasi kebutuhan pengguna, proses perencanaan dan peningkatan kualitas layanan sehingga dapat meningkatkan kepuasan pengguna terhadap JNE Sub Purwokerto. Berdasarkan penelitian terdahulu oleh Fildza Sakinah Alnaz dan Warih Maharani berjudul Analisis Sentimen Pengguna MY JNE Menggunakan Algoritma *Naïve Bayes* dengan menggunakan dataset yang bersumber dari cuitan pengguna berbahasa Indonesia berjumlah 4401 data *tweet*. Metode klasifikasi dengan *naïve bayes* model multinominal *naïve bayes* dan perbandingan ekstraksi fitur *n-gram* dan pembobotan kata dengan *tf-idf*. Hasil penelitian didapatkan nilai akurasi dengan jumlah 65%, nilai rata-rata *precision* yaitu 69%, *recall* yaitu 64% dan *f1-score* nya yaitu 66%[8]. Dari permasalahan sebelumnya, penelitian berjudul "Analisis sentimen ulasan pengguna terhadap ekspedisi layanan logistik dengan metode *naïve bayes*" memiliki urgensi untuk memahami perasaan dan tanggapan emosional dari ulasan pelanggan terkait layanan dalam pengiriman barang oleh pihak JNE Sub Agen Purwokerto.

II. KAJIAN TEORI

Kajian teori yang berkaitan dengan pendukung dan literatur pustaka berisi tentang penelitian terdahulu yang memiliki keterkaitan permasalahan memberi landasan teoritis dan pemahaman bagaimana penyelesaian masalah. Seperti pada penelitian tentang Analisis Sentimen Pengguna MY JNE dengan Algoritma *Naïve Bayes* membahas mengenai analisis sentimen pengguna layanan aplikasi MY JNE pada *google playstore* memiliki ulasan pengguna terkait pengalaman terhadap layanan JNE. Dalam penelitian ini, sebanyak 996 ulasan pengguna yang dilabeli sentimen berdasarkan *rating* dengan rincian 1, 2, dan 3 (negatif), 4 dan 5 (positif). *Splitting* data dengan jumlah data latih sebesar 80% dan data uji 20%. Algoritma klasifikasi dengan model *naïve bayes classifier* menghasilkan model dengan nilai akurasi sebesar 96%. Hasil sentimen dari aplikasi My JNE ini mendapatkan hasil *negative* karena ulasan yang terdapat pada data terbesar 70% pada nilai *precision*, dan tingkat keberhasilan (*recall*) 45% Hasil analisis sentimen menunjukkan bahwa mayoritas ulasan (96,8%) bersifat negatif, yang mengindikasikan tingkat kepuasan pengguna terhadap aplikasi My JNE cenderung rendah[10].

A. Analisis Sentimen

Analisis Sentimen adalah cara untuk mengevaluasi dan menilai opini atau ekspresi yang ada dalam teks. Tujuannya adalah untuk mengidentifikasi apakah sentimen dalam konten tersebut positif, negatif, atau netral[10]. Dengan menggunakan teknik-teknik pengolahan bahasa alami dan pembelajaran mesin, analisis sentimen membantu kita memahami pandangan dan perasaan individu terkait topik atau produk tertentu[11]. Proses ini dapat melibatkan penggunaan teknik-teknik analisis teks, pembelajaran mesin dan pemrosesan bahasa alami *Natural Language Processing* (NLP) untuk mengidentifikasi kata-kata kunci, pola, dan konteks yang mengidentifikasi sentimen tertentu[12].

B. Confusion Matrix

Confusion matrix adalah metode untuk menilai hasil performa model klasifikasi. Evaluasi yang dihitung meliputi nilai *accuracy*, *precision*, *recall*, dan *f1- Score*. Untuk menghitung nilai tersebut dapat digunakan rumus persamaan berikut[16].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

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

$$F1 - Score = 2 \times \left(\frac{Precision \times Recall}{Precision + Recall} \right) \quad (4)$$

C. Data Mining

Data Mining adalah proses penemuan pola-pola baru dari kumpulan data yang sangat besar. Metode-metode yang digunakan dalam *Data Mining* melibatkan *artificial intelligence*, *machine learning*[13]. Proses data mining dilakukan dengan pengumpulan data yang akan dibersihkan dan dipreproses untuk menghilangkan *noise* atau ketidakpastian yang mungkin mempengaruhi hasil analisis[20].

D. Ekstraksi fitur

Ekstraksi fitur dilakukan untuk memberikan bobot nilai pada setiap kata dalam masing-masing ulasan. Metode yang digunakan adalah metode *bag of words* (BoW) model *countvectorizer*. Proses ini juga biasa dikenal sebagai pembobotan kata dengan mengubah data teks menjadi representasi numerik[14].

E. Google Maps API

Antarmuka Pemrograman aplikasi (API) milik *google* untuk mengakses berbagai layanan dan data dari *google maps*[15]. API ini memungkinkan pengembang untuk membuat aplikasi yang dapat menggunakan peta, geolokasi, dan data terkait informasi tempat, rute, dan berbagai fitur geospasial lainnya[20][25].

F. Google Maps Reviews

Google maps reviews adalah platform *google maps* dengan layanan berupa platform *feedback*. Dengan mengumpulkan dan mengevaluasi *feedback* pengguna terhadap lokasi atau fasilitas tertentu, informasi berharga dapat diperoleh[16]. Melalui analisis ini, penelitian dapat memanfaatkan opini dan pemikiran publik untuk menggambarkan gambaran yang komprehensif tentang sejauh mana suatu tempat dihargai atau dianggap penting oleh komunitas lokal atau pengunjung[17].

G. Jalur Nugraha Ekakurir

Jalur Nugraha Ekakurir (JNE) adalah sebuah perusahaan jasa pengiriman barang yang berbasis di Indonesia[4]. JNE didirikan pada tahun 1990 dan menjadi penyedia layanan ekspedisi terkemuka di Indonesia[5]. Tujuan utama dari JNE adalah untuk menyediakan layanan pengiriman yang andal dan efisien kepada pelanggan, baik itu dalam skala individu maupun bisnis[4].

H. Naïve Bayes Clasiffier

Naïve bayes classifier adalah metode klasifikasi yang berdasarkan *teorema bayes*. Metode ini melakukan perhitungan probabilitas untuk membuat prediksi dengan cepat dan akurat. Metode ini mengasumsikan independensi antara fitur, dengan memprediksi label data dengan membandingkan probabilitas munculnya fitur-fitur tertentu dalam setiap kategori, cocok untuk analisis sentimen, kategorisasi dokumen, dan deteksi spam[8]. Metode *Naïve bayes* memiliki rumus utama yang dijadikan perhitungan dalam penyelesaian yang berdasarkan *teorema bayes* yaitu[10].

$$P(H|X) = \frac{P(X|H)P(H)}{P(X)} \quad (5)$$

Keterangan:

X = Data dengan *class* yang belum diketahui
 H = Hipotesis data X adalah suatu kelas
 $P(H|X)$ = Probabilitas hipotesis H berdasarkan kondisi x
 $P(H)$ = Probabilitas hipotesis H
 $P(X|H)$ = Probabilitas X berdasarkan kondisi tersebut
 $P(X)$ = Probabilitas dari X

I. SMOTE

Teknik *Synthetic Minority Over-sampling Technique* (SMOTE) merupakan teknik untuk menangani ketidakseimbangan kelas dalam dataset[18]. Proses dimulai dengan menghitung jarak antar data pada data minoritas, selanjutnya menentukan k terdekat, dan terakhir menghasilkan ringkasan data. Untuk persamaan sebagai berikut[19].

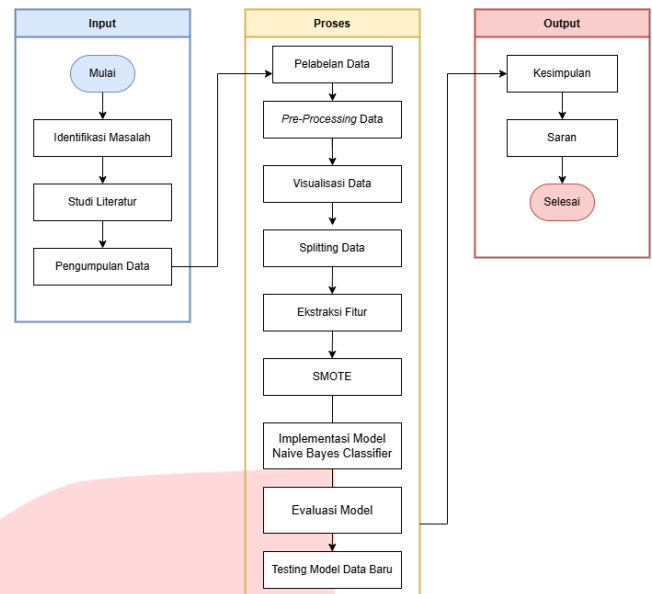
$$x_{syn} = x_i + (x_{knn} - x_i) \times \partial \quad (6)$$

Keterangan:

x_{syn} = Data sintetis yang diciptakan
 x_i = Data yang akan direplikasi (ditiru)
 x_{knn} = Data yang memiliki jarak terdekat dari data yang akan direplikasi (ditiru)
 ∂ = Nilai acak antara 0 sampai 1

III. METODE

Dalam melakukan penelitian ini, dibutuhkan prosedur pengolahan data serta tahapan dalam pengolahan data kemudian proses implemtasi model dan evaluasi model yang dibuat. Berikut adalah diagram alir dari penelitian ini.



Gambar 3.1 Diagram Alir penelitian

Pada Gambar 3.1 merupakan tahapan proses yang akan dilakukan pada penelitian ini. Berikut tahapan penelitian yaitu sebagai berikut.

A. Identifikasi Masalah

Pada tahapan pertama, peneliti membuat rumusan masalah dengan mengamati trend atau permasalahan yang sedang ramai atau populer saat ini mengenai penggunaan jasa layanan pengiriman logistik.

B. Studi Literatur

Pada tahapan kedua, peneliti mencari literatur – literatur pada penelitian sebelumnya seperti jurnal, artikel, atau dokumen relevan lainnya yang membahas terkait analisis ulasan pada fenomena tertentu.

C. Pengumpulan Data

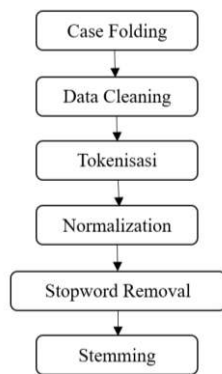
Teknik dalam pengumpulan data dilakukan lewat *website google maps* dengan menggunakan metode *crawling* data untuk pengambilan datanya. Total jumlah data yang digunakan dalam penelitian ini adalah sebanyak 968 data ulasan pengguna *google maps* dengan rincian 111 data berlabel sentimen positif dan data 867 berlabel negatif.

D. Pelabelan Data

Pelabelan data ulasan dibagi menjadi dua kategori kelas yaitu positif dan negatif, dengan membagi kelas berdasarkan *rating* pengguna. Untuk *rating* 3, 4, dan 5 diklasifikasikan sebagai kelas positif. Sedangkan, *rating* 1 dan 2 diklasifikasikan sebagai kelas negatif.

E. Pre-processing Data

Langkah selanjutnya adalah *pre-processing* data ulasan untuk membersihkan data ulasan yang nantinya akan digunakan, hal ini diperlukan karena data yang didapatkan tidak seluruhnya dapat digunakan dalam proses analisis. Dalam proses ini akan menghasilkan data ulasan yang sesuai dengan kebutuhan pada proses analisis sehingga didapatkan hasil yang lebih baik.



Gambar 3.2 Tahap dalam *pre-processing* data

F. Visualisasi Data

Tahap visualisasi data dilakukan untuk memahami pola data sentimen dengan menampilkan kategori jumlah banyaknya data muncul dalam masing-masing ulasan yang sudah berlabel sentimen.

G. Ekstraksi Fitur

Tahap ekstraksi fitur dilakukan untuk memberikan bobot nilai atau mengubah data teks ulasan menjadi representasi numerik pada data atau istilah dalam masing-masing ulasan[20]. Proses ini juga dikenal sebagai pembobotan kata, dengan metode yang digunakan adalah *Bag of Words (BoW)* model *countvectorizer*.

H. Splitting Data

Pada langkah ini adalah proses membagi dataset menjadi dua bagian untuk melatih dan mengevaluasi model pembelajaran mesin secara efektif. Pembagian data dilakukan dengan skema 80:20, dimana kondisi 80% dataset untuk data latih (*training*) dan 20% dataset untuk data uji (*testing*).

I. SMOTE

Pada langkah ini adalah proses menyeimbangkan data latih yang tidak seimbang (*imbalanced dataset*) dengan menggunakan teknik SMOTE (*Synthetic Minority Oversampling Technique*). Proses ini dilakukan untuk meningkatkan representasi kelas minoritas dalam data latih, sehingga model pembelajaran mesin tidak bias terhadap dataset.

J. Implementasi Model *Naïve Bayes*

Pada langkah ini adalah implementasi model dengan metode *naïve bayes classifier*. Metode ini menggunakan algoritma yang memprediksi probabilitas[21]. Pada Metode ini model yang akan digunakan adalah model Multinomial *Naïve Bayes*, karena implemtasi mudah terutama pada dataset yang besar dan memiliki tingkat akurasi yang cukup tinggi.

K. Evaluasi Model

Pada langkah ini model akan menggunakan *confusion matrix* untuk menganalisis performa model klasifikasi dengan berdasarkan kriteria nilai penilaian hasil akurasi, presisi, *recall*, dan *f1-score*.

L. Testing Model

Pada langkah ini model yang sudah didapatkan pada proses pelatihan dengan metode yang telah dibuat. Proses ini akan melakukan pengujian dengan data yang diinputkan ke dalam mesin pengolah data *google colab*. Hasil dari *testing* model ini peneliti akan mengetahui seberapa baik model tersebut dalam mengklasifikasikan data baru.

IV. HASIL DAN PEMBAHASAN

A. Pengumpulan Data

Proses *crawling* data dengan kode baris *python* dengan menggunakan *API key google*. Total jumlah data yang terkumpul sebanyak 968 data ulasan pengguna *google maps* dengan rincian 111 data (positif) dan 867 (negatif).

Tabel 4.1 Sampel dataset ulasan

Pengulas	Rating	Waktu	Ulasan
Nurfaiza S	1	4 bulan lalu	Buruk coba dipercepat nya padahal cuma 10 menit dari kos saya
Ayu Sirri Aini	2	1 tahun lalu	Berkali-kali belanja online klo dapet selalu lamaaaaa. Sangat kecewa!
Gentur Ageng	3	1 tahun lalu	sangat cepat waktu, sampai harus dijemput ke konter
Dimas Hilaal Permana	4	1 bulan lalu	Cepat

B. Pelabelan Data

Proses pelabelan data dengan membagi kelas sentimen menjadi 2 klasifikasi kelas, yaitu kelas positif dan kelas negatif dimana klasifikasi label berdasarkan *rating* pada tiap ulasan, untuk *rating* 3, 4, dan 5 (positif) dan untuk *rating* 1 dan 2 (negatif).

Tabel 4.2 Sampel pelabelan data

Rating	Ulasan	Label
1	Buruk coba dipercepatnya padahal cuma 10 menit	negatif
2	Berkali-kali belanja online klo dapet selalu lamaaaaa. Sangat kecewa!	negatif
3	sangat cepat waktu,sampai harus dijemput ke konter	positif
4	Cepat	positif

C. *Pre-processing* Data

Tahapan selanjutnya adalah *pre-processing* data atau pembersihan data ulasan dari *noise* seperti *punctuation*, *number*, *emoticon*[18]. Dalam tahap ini akan dilakukan pemilihan atribut dataset yang nantinya akan diproses. Untuk atribut data yang digunakan hanya data ulasan, dan label sentimen. Berikut tahapan dalam *preprocessing* data:

1. *Case Folding*

Proses mengubah semua huruf dalam teks menjadi bentuk yang seragam huruf kecil (*lowercase*)[18].

Tabel 4.3 Sampel hasil *case folding*

Ulasan	Case folding
Wehh paket gua tiba tiba di retur gimana nih ceritanya gak jelas jne nih	wehh paket gua tiba tiba di retur gimana nih ceritanya gak jelas jne nih
Najis bgt pesen di toped kirimnya pake jne.. 100% pasti bermasalah	najis bgt pesen di toped kirimnya pake jne.. 100% pasti bermasalah

2. Data Cleaning

Proses menghilangkan karakter, angka, tanda baca, emoji, spasi ganda seperti “, ”, “.”, “?”, “\$”, “1”, “0”, “%” dan sebagainya[18].

Tabel 4.4 Sampel hasil data *cleaning*

Case folding	Data clean
coba dipercepat pengiriman paketnya padahal cuma 10 menit dari kos saya	coba dipercepat pengiriman paketnya padahal cuma menit dari kos saya
najis bgt pesen di toped kirimnya pake jne.. 100% pasti bermasalah	najis bgt pesen di toped kirimnya pake jne pasti bermasalah

3. Tokenisasi

Proses membagi data ke dalam unit lebih kecil atau disebut dengan token[18].

Tabel 4.5 Sampel hasil data tokenisasi

Data Clean	Tokenisasi
coba dipercepat pengiriman paketnya padahal cuma menit dari kos saya	['coba','dipercepat','pengiriman','paketnya','padahal','cuma','menit','dari','kos','saya']
najis bgt pesen di toped kirimnya pake jne pasti bermasalah	['najis','bgt','pesen','di','toped','kirimnya','pake','jne','pasti','bermasalah']

4. Normalization

Proses mengubah kata-kata atau frasa yang bersifat informal atau populer (*slang*) menjadi bentuk yang lebih standar[18].

Tabel 4.6 Sampel hasil *normalization*

Tokenisasi	Data Normal
['coba','dipercepat','pengiriman','paketnya','padahal','cuma','menit','dari','kos','saya']	['coba','dipercepat','pengiriman','paketnya','padahal','cuma','menit','dari','kos','saya']
['najis','bgt','pesen','di','toped','kirimnya','pake','jne','pasti','bermasalah']	['najis','banget','pesan','di','tokopedia','kirimnya','pakai','jalur','nugraha','ekakurir','pasti','bermasalah']

5. Stopword Removal

Proses menghapus kata-kata umum yang tidak mengandung makna dalam teks[18].

Tabel 4.7 Sampel hasil *stopword removal*

Data Normal	Teks Stopword
['coba','dipercepat','pengiriman','paketnya','padahal','cuma','menit','dari','kos','saya']	['coba','dipercepat','pengiriman','paketnya','menit','kos']
['najis','banget','pesan','di','tokopedia','kirimnya','pakai','pasti','bermasalah']	['najis','banget','pesan','tokopedia','kirimnya','pakai','bermasalah']

6. Stemming

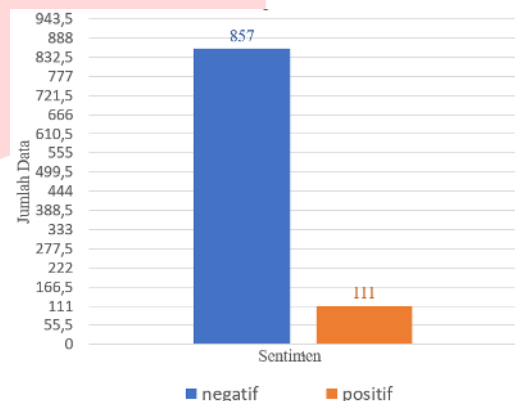
Proses untuk mengubah kata menjadi bentuk dasar dengan menghapus awalan atau akhiran tertentu[18].

Tabel 4.8 Sampel hasil *stemming*

Full Kalimat	Data Bersih
coba dipercepat pengiriman paketnya menit kos	coba cepat kirim paket menit kos
najis banget pesan tokopedia kirimnya pakai jalur nugraha ekakurir bermasalah	najis banget pesan tokopedia kirim pakai jalur nugraha ekakurir masalah

D. Visualisasi Data

Setelah data ulasan melewati proses pembersihan data atau *pre-processing* data, kemudian melakukan visualisasi data untuk melihat jumlah dataset sentimen yang berlabel positif dan negatif dalam teks ulasan.

**Gambar 4.1** Hasil visualisasi data sentimen

E. Ekstraksi Fitur

Tahap ekstraksi fitur dilakukan untuk memberikan bobot nilai atau numerik pada data dalam masing-masing ulasan. Pada proses ini metode yang digunakan adalah *Bag of Words (BoW)* model *countvectorizer*. Berikut Tabel 4.9, adalah ilustrasi bagaimana penggunaan *Bag of Words (BoW)* model *countvectorizer* dengan mengambil contoh dari data ulasan pada dataset yang telah diproses sebelumnya.

Tabel 4.9 Contoh dataset ulasan

Ulasan	Label
barang kirim ambil tidak ramah	Negatif
jelek banget paket	Negatif
kirim paket cepat	Positif
respon cepat ramah	Positif

Berdasarkan Tabel 4.9, proses ilustrasi dalam ekstraksi fitur dengan BoW model *countvectorizer* adalah sebagai berikut.

1. Memecah Kata

Disini akan dilakukan pemecahan kalimat ulasan menjadi kata-kata individual (token). Berikut contoh token/kata untuk setiap dokumen:

- Dok 1: ['barang', 'kirim', 'ambil', 'tidak', 'ramah']
- Dok 2: ['jelek', 'banget', 'paket']
- Dok 3: ['kirim', 'paket', 'cepat']
- Dok 4: ['respon', 'cepat', 'ramah']

2. Membangun kosakata

Membuat daftar semua kata unik yang muncul dalam kumpulan dokumen (*corpus*). Berikut kosakata yang dihasilkan adalah:

['ambil', 'banget', 'barang', 'cepat', 'jelek', 'kirim', 'paket', 'ramah', 'respon', 'tidak']

3. Membentuk matriks frekuensi kata

Model menghitung jumlah kemunculan setiap kata (frekuensi) dalam masing-masing dokumen dengan nilai *binary* 1 menunjukkan bahwa kata tersebut muncul dalam dokumen, sedangkan nilai *binary* 0 menunjukkan ketidakhadirannya. Berikut hasil mengubah data menjadi nilai numerik.

Tabel 4.10 Hasil matriks frekuensi kata

Dokum	ambil	banget	barang	cepat	jelek	kirim	paket	ramah	respon	tahan
D1	1	0	1	0	0	1	0	1	0	0
D2	0	1	0	0	1	0	1	0	0	1
D3	0	0	0	1	0	1	1	0	0	0
D4	0	0	0	1	0	0	0	1	1	0

F. Splitting Data

Pada langkah ini adalah proses membagi dataset menjadi dua bagian untuk melatih dan menguji model pembelajaran mesin. Proses ini untuk mengukur performa model secara obyektif pada data yang baru bukan dari data latih. Pembagian dilakukan dengan skema 80:20, dengan rincian 80% data latih (*training*) dan 20% data uji (*testing*). Pembagian menghasilkan 774 data latih dan 194 data uji.

G. SMOTE

Pada langkah ini adalah proses menyeimbangkan data latih yang tidak seimbang. Proses ini dilakukan untuk meningkatkan representasi kelas minoritas dalam data latih. Pada penggunaan teknik ini terhadap ketidakseimbangan data untuk pelatihan yang semula dengan rincian data latih 89 (positif) dan 685 (negatif) menjadi data yang seimbang dengan rincian 685 (positif) dan 685 (negatif).

$$x_{syn} = x_i + (x_{knn} - x_i) \times \partial \quad (7)$$

Keterangan:

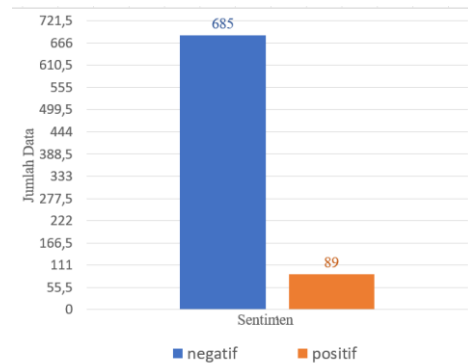
x_{syn} = Data sintetis yang diciptakan

x_i = Data yang akan direplikasi (ditiru)

x_{knn} = Data yang memiliki jarak terdekat dari data yang akan direplikasi (ditiru)

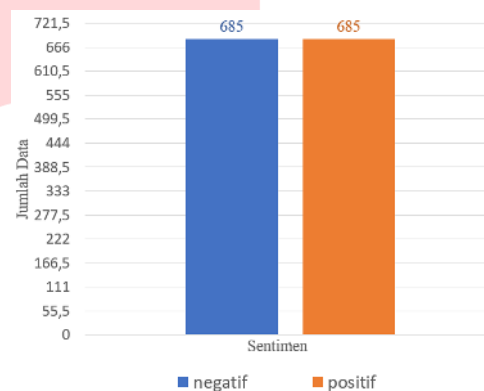
∂ = Nilai acak antara 0 sampai 1

Untuk hasil visualisasi data awal yang tidak seimbang dapat dilihat pada Gambar 4.2 berikut ini.



Gambar 4.2 Perbandingan data latih sebelum SMOTE

Pada Gambar 4.3 berikut ini menunjukkan perbandingan kelas data latih antara mayoritas dan minoritas yang sudah seimbang.



Gambar 4.3 Perbandingan data latih sesudah SMOTE

H. Implementasi Model Naïve Bayes

Pada langkah ini adalah implementasi model dengan metode *naïve bayes classifier*. Metode *naïve bayes* memanfaatkan perhitungan probabilitas untuk membuat prediksi dengan cepat dan akurat[21].

$$P(H|X) = \frac{(P(X|H)P(H))}{P(X)} \quad (8)$$

Keterangan:

X = Data dengan class yang belum diketahui

H = Hipotesis data X adalah suatu kelas

$P(H|X)$ = Probabilitas hipotesis H berdasarkan kondisi x

$P(H)$ = Probabilitas hipotesis H

$P(X|H)$ = Probabilitas X berdasarkan kondisi tersebut

$P(X)$ = Probabilitas dari X

Pada implementasi *Naïve Bayes* (NB), khususnya model *Multinomial Naïve Bayes* (MNB), Metode ini bekerja berdasarkan prinsip *teorema bayes* dan mengasumsikan bahwa setiap fitur dalam dataset bersifat independen (*naïve assumption*)[21]. Berikut adalah contoh perhitungan menggunakan MNB untuk analisis sentimen dengan mengambil beberapa data ulasan. Pada Tabel 4.11 di bawah ini adalah contoh dataset yang akan digunakan sebagai contoh dalam penggunaan metode NB model MNB.

Tabel 4.11 Sampel dataset ulasan

Ulasan	Label
barang kirim ambil layanan tidak ramah	Negatif
jelek banget paket tahan	Negatif
kirim paket cepat tuju beli	Positif
respon cepat ramah	Positif

Penjelasan ilustrasi implementasi metode *naive bayes* model *multinomial naive bayes* sebagai berikut:

1. Melakukan pemecahan teks menjadi kata-kata unik (*vocabulary*):

Vocabulary: {barang, kirim, ambil, layanan, tidak, ramah, jelek, banget, paket, tahan, cepat, tuju, beli, respon}

Dimana ukuran kosakata/jumlah kata di atas (V) = 14.

2. Menghitung probabilitas prior ($P(c)$)
Terdapat jumlah total dokumen = 4, kemudian jumlah dokumen Positif = 2

$$P(\text{positif}) = \frac{(\text{Jumlah Dok di Kelas Positif})}{(\text{Total jumlah Dok})} = \frac{2}{4} = 0.5$$

Lalu jumlah dokumen Negatif = 2

$$P(\text{negatif}) = \frac{(\text{Jumlah Dok di Kelas Negatif})}{(\text{Total jumlah Dok})} = \frac{2}{4} = 0.5$$

3. Menghitung *Likelihood* / probabilitas data (fitur atau kata-kata) ($P(w | c)$)

Dengan untuk rumus menggunakan ($\alpha=1$):

Lalu Frekuensi kata dalam kelas Positif dengan total jumlah kata Positif = 7, dimana terdiri dari kata: {kirim:1, ramah:1, paket:1, cepat:2, tuju:1, beli:1, respon:1}

$$P(w | \text{Positif}) = \frac{(\text{Count}(w, \text{Positif}) + 1)}{(\text{Total Kata Positif} + V)} \quad (9)$$

Contoh perhitungan dengan kata: “ramah”

$$P(\text{ramah} | \text{Positif}) = \frac{(1 + 1)}{(7 + 14)} = \frac{2}{21}$$

Lalu Frekuensi kata dalam kelas Negatif dengan total jumlah kata Negatif = 10, dimana terdiri dari kata: {barang:1, kirim:1, ambil:1, layanan:1, tidak:1, ramah:1, jelek:1, banget:1, paket:1, tahan:1}

$$P(w | \text{Negatif}) = \frac{(\text{Count}(w, \text{Negatif}) + 1)}{(\text{Total Kata Negatif} + V)} \quad (10)$$

Contoh perhitungan dengan kata: “jelek”

$$P(\text{jelek} | \text{Negatif}) = \frac{(1 + 1)}{(10 + 14)} = \frac{2}{24}$$

4. Klasifikasi dokumen baru

Misalkan dengan menggunakan contoh untuk mengklasifikasikan teks dokumen:

Teks dokumen: “kirim paket cepat tuju beli”

Kemudian menghitung *likelihood* / probabilitas data (fitur atau kata-kata) untuk kelas Positif dengan:

$$\{P(\text{Positif} | d) = P(\text{Positif}) \cdot P(\text{kirim} | \text{Positif}) \cdot P(\text{paket} | \text{Positif}) \cdot P(\text{cepat} | \text{Positif}) \cdot P(\text{tju} | \text{Positif}) \cdot P(\text{beli} | \text{Positif})\}$$

Lalu lakukan substitusi nilai:

$$P(\text{Positif} | d) = 0.5 \cdot \frac{2}{21} \cdot \frac{2}{21} \cdot \frac{2}{21} \cdot \frac{2}{21} \cdot \frac{2}{21}$$

$$P(\text{Positif} | d) = 0.5 \cdot \frac{32}{4.084.101} = \frac{16}{4.084.101}$$

$$= 3.9176 \times 10^{-6}$$

Kemudian menghitung *likelihood* / probabilitas data (fitur atau kata-kata) untuk kelas Negatif dengan:

$$\{P(\text{Negatif} | d) = P(\text{Negatif}) \cdot P(\text{kirim} | \text{Negatif}) \cdot P(\text{paket} | \text{Negatif}) \cdot P(\text{cepat} | \text{Negatif}) \cdot P(\text{tju} | \text{Negatif}) \cdot P(\text{beli} | \text{Negatif})\}$$

Lalu lakukan substitusi nilai:

$$P(\text{Negatif} | d) = 0.5 \cdot \frac{2}{24} \cdot \frac{2}{24} \cdot \frac{1}{24} \cdot \frac{1}{24} \cdot \frac{1}{24}$$

$$P(\text{Negatif} | d) = 0.5 \cdot \frac{4}{7.962.624} = \frac{2}{7.962.624}$$

$$= 2.512 \times 10^{-7}$$

5. Tentukan nilai

Bandingkan hasil probabilitas untuk menentukan kelas dengan probabilitas lebih tinggi:

- $P(\text{Positif} | d) = 3.9176 \times 10^{-6}$
- $P(\text{Negatif} | d) = 2.512 \times 10^{-7}$

6. Analisis hasil

- a. Perbandingan probabilitas

Probabilitas dokumen d termasuk dalam kelas *Positif* $P(\text{Positif} | d)$ lebih besar dibandingkan probabilitas dokumen termasuk dalam kelas *Negatif* $P(\text{Negatif} | d)$. Dimana secara relatif hasilnya adalah:

$$P(\text{Positif} | d) > P(\text{Negatif} | d)$$

- b. Interpretasi

Berdasarkan hasil ini, dokumen d lebih cenderung diklasifikasikan ke dalam kelas *Positif* daripada *Negatif*. Probabilitas untuk kedua kelas kecil karena nilai *likelihood* (kemungkinan kata-kata dalam dokumen cocok dengan kelas tertentu) rendah, tetapi perbandingan menunjukkan bahwa kelas *Positif* lebih mungkin.

- c. Keputusan akhir

Dokumen d {“kirim paket cepat tuju beli”} dapat diklasifikasikan sebagai sentimen Positif, karena memiliki nilai $P(\text{Positif} | d)$ yang lebih besar daripada $P(\text{Negatif} | d)$.

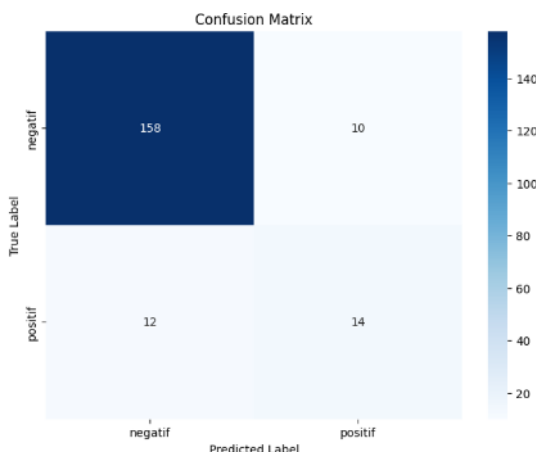
I. Evaluasi Model

Selanjutnya tahap evaluasi yaitu menentukan metrik hasil evaluasi yang terdiri dari akurasi, presisi, *recall*, *f1-score*, dan *support* untuk setiap kelas (negatif dan positif). Untuk hasil *classification report* sebagai berikut.

Tabel 4.12 Hasil evaluasi *classification report*

Hasil/ Klasifikasi	Precision	Recall	F1-Score	Support
Accuracy			0.89	194
Negatif	0.93	0.94	0.93	168
Positif	0.58	0.54	0.56	26
Macro avg	0.76	0.74	0.75	194
Wighted avg	0.88	0.89	0.88	194

Pada Tabel 4.12 di atas adalah output dari yang menghasilkan dari model yang telah diuji. Hasil klasifikasi dengan dataset ulasan 968 menghasilkan nilai *accuracy* 89%, nilai *precision* 93% (positif) dan 58% (negatif), untuk nilai *recall* sebesar 94% (positif) dan 54% (negatif), serta untuk nilai *f1-score* 93% (positif) dan 56% (negatif). Hasil evaluasi performa model klasifikasi sebagai berikut.



Gambar 4.4 Hasil performa klasifikasi *confusion matrix*

Hasil analisis sentimen menunjukkan bahwa mayoritas ulasan memiliki sentimen negatif, yang mengindikasikan adanya ketidakpuasan pengguna terhadap layanan yang dianalisis. Meskipun model memiliki akurasi yang tinggi, performa dalam mengenali ulasan negatif masih kurang optimal dibandingkan dengan ulasan positif, hasil pengujian evaluasi yang lebih rendah untuk kelas negatif. Dengan kata lain, dari seluruh data yang digunakan untuk evaluasi, sebanyak 89% sampel memiliki label prediksi yang sesuai dengan label sebenarnya. Hal ini mencerminkan tingkat keberhasilan model dalam mengklasifikasikan data secara keseluruhan dengan benar.

J. Testing Model

Pada proses *testing* model akan dilakukan untuk pengujian terhadap data baru yang diinputkan secara langsung ke dalam program. Tujuan dari *testing* model ini adalah untuk mengetahui sejauh mana keakurasian model klasifikasi yang telah dibuat.

Tabel 4.14 Sampel *testing* dengan data baru

Text ulasan	Hasil analisis
Barang sampai dengan selamat, kondisi aman dan bagus	positif
Barang datang terlambat, kondisi rusak sebagian	negatif
Barang sampai ditujuan dengan selamat, rekomended pokoknya	negatif

V. KESIMPULAN

Berdasarkan hasil analisis dan evaluasi pelatihan dengan menerapkan metode *naïve bayes classifier* terhadap data yang sudah diperoleh dengan crawling data pada layanan Logistik JNE Sub Agen Purwokerto *reviews google maps* didapatkan sejumlah 968 data ulasan pengguna, dengan rincian 111 (positif) dan 857 (negatif). Berdasarkan hasil dari pembagian data (*splitting*) 80:20 ini telah didapatkan data *training* sebesar 774 data ulasan, dan data *testing* sebesar 146 data ulasan. Pada data *training* yang memiliki jumlah ketidakseimbangan data antara sentimen negatif yang berjumlah 685 dan sentimen positif yang berjumlah 89 juga telah dilakukan *balancing* data dengan teknik *oversampling* SMOTE dan didapatkan hasil seimbang antara kelas positif dan negatif. Implementasi model pelatihan menunjukkan hasil performa dengan evaluasi *confusion matrix* memperoleh nilai akurasi 89%, diikuti untuk kelas negatif nilai presisi 93%, *recall* 94%, dan *f1-score* 93%. Sedangkan, untuk kelas positif nilai presisi 58%, nilai *recall* 54%, dan nilai *f1-score* 56%. Untuk hasil dari analisis sentimen terhadap ulasan JNE Sub Agen Purwokerto menunjukkan sentimen negatif, yang mengindikasikan adanya ketidakpuasan pengguna terhadap layanan yang dianalisis, hasil ini menunjukkan sentimen yang lebih rendah pada kelas negatif.

REFERENSI

- [1] A. Ayu, D. Lestari, and A. Merthayasa, "Peran Teknologi Dalam Perubahan Bisnis Di Era Globalisasi," vol. 7, no. 11, pp. 2548–1398, 2022.
- [2] C. A. Cholik, "Perkembangan Teknologi Informasi Komunikasi / ICT Dalam Berbagai Bidang," *Jurnal Fakultas Teknik Unisa Kuningan*, vol. 2, no.2, May. 2021.
- [3] F. Parsakh Nursyamsyi and F. Noor Hasan, "Klasifikasi Sentimen Terhadap Aplikasi Identitas Kependudukan Digital Menggunakan Algoritma Naïve Bayes dan SVM," *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 4, no. 3, pp. 1788–1798, Dec. 2023.
- [4] R. Ananda Irhasr Maha Adiprayitno and Drs. Muhammad Edwar, M.Si "Pengaruh Kualitas Layanan Dan Harga Terhadap Keputusan Penggunaan Jasa Pengiriman Barang JNE (Jalur Nugraha Ekakurir) Di Agen Putro Agung Wetan Surabaya," *Jurnal Pendidikan Tata Niaga (JPTN)*, vol.1, no.1, 2020.
- [5] Putri Wulandari, Heru Tri Sutiono, and Sri Kussujaniatun, "Pengaruh Kualitas Layanan Dan Citra Merek Terhadap Loyalitas Pelanggan Melalui Kepuasan Pelanggan Pada Pelanggan Jasa Jne Di

- Yogyakarta,” *JIMKES (Jurnal Ilmiah Manajemen Kesatuan)*, vol. 9, no. 2, pp. 293-308, June. 2021.
- [6] M. Farid Fernanda Saputra, Alex Ferdinal, and Dini Elida Putri, “Analisis Kualitas Pelayanan Terhadap Kepuasan Pelanggan Jasa Pengiriman J&T Express Koto Baru (Studi Kasus Masyarakat Jorong Seberang Piruko Timur Dusun Standart),” *INNOVATIVE: Journal Of Social Science Research*, vol. 3 no. 5, 2023.
- [7] T. S. de Souza Viana, M. de Oliveira, T. L. C. da Silva, M. S. R. Falcão, and E. J. T. Gonçalves, “A message classifier based on multinomial Naïve Bayes for online social contexts,” *Journal of Management Analytics*, Jul. 2020.
- [8] F. Sakinah Alnaz and W. Maharani, “Analisis Emosi Melalui Media Sosial Twitter Dengan Menggunakan Metode Naïve Bayes dan Perbandingan Fitur N-gram dan TF-IDF,” 2020.
- [9] D. Purwitasari, Adi Surya Suwardi Ansyah, Arya Putra Kurniawan, and Asiyah Nur Kholifah, “A Hybrid Method on Emotion Detection for Indonesian Tweets of COVID-19,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 7, no. 2, pp. 254–262, Mar. 2023.
- [10] F. Nuriy Alin, M. Hafid Totohendarto, and M. Rafi Muttaqin, “Analisis Sentimen Terhadap Kendaraan Listrik Pada Platform Twitter Menggunakan Metode Naïve Bayes,” *Informatics for Educators And Professionals: Journal of Informatics*, vol. 8, no. 2, pp. 96–107, Dec. 2023.
- [11] F. Syah, H. Fajrin, A. N. Afif, R. Saeputra, D. Mirranty, and D. D. Saputra, “Analisa Sentimen Terhadap Twitter IndihomeCare Menggunakan Perbandingan Algoritma Smote, Support Vector Machine, AdaBoost dan Particle Swarm Optimization,” *Jurnal JTIK (Jurnal Teknologi Informasi dan Komunikasi)*, vol. 7, no. 1, Jan. 2023.
- [12] S. Khomsah, R. Dias Ramadhani, and S. Wijayanto, “Big Data Analytics to Analyze Sentiment, Emotions, and Perceptions of Travelers (Case Study: Tourism Destination in Purwokerto Indonesia),” *Jurnal E-Komtek (Elektro-Komputer-Teknik)*, vol. 5, no. 2, pp. 284–297, Dec. 2021.
- [13] J. Khatib Sulaiman, M. Ikhsan, and R. R. Kurniawan, “Penerapan Text Mining pada Sistem Rekomendasi Pembimbing Skripsi Mahasiswa Menggunakan Algoritma Naïve Bayes Classifier,” *Indonesian Journal of Computer Science Attribution*, vol. 12, no. 6, pp. 2023–4196, Dec. 2023.
- [14] F. Agung, J. Ayomi, and K. E. Dewi, “Analisis Emosi Pada Media Sosial Twitter Menggunakan Metode Multinomial Naïve Bayes Dan Synthetic Minority Oversampling Technique,” *KOMPUTA: Jurnal Ilmiah Komputer dan Informatika*, vol. 12, no. 2, 2023.
- [15] S. K. Dirjen, P. Riset, D. Pengembangan, R. Dikti, E. Sutoyo, and A. Almaarif, “Terakreditasi SINTA Peringkat 2 Educational Data Mining untuk Prediksi Kelulusan Mahasiswa Menggunakan Algoritme Naïve Bayes Classifier,” *masa berlaku mulai*, vol. 1, no. 3, pp. 95–101, 2017.
- [16] Bee Shin, Sohee Ryu, Yongjun Kim, and Dongwhan Kim “Analysis on Review Data of Restaurants in Google Maps through Text Mining: Focusing on Sentiment Analysis,” vol. 9, no. 1, pp. 61–68, Mar. 2022.
- [17] W. Khofifah, D. N. Rahayu, and A. M. Yusuf, “Analisis Sentimen Menggunakan Naive Bayes Untuk Melihat Review Masyarakat Terhadap Tempat Wisata Pantai Di Kabupaten Karawang Pada Ulasan Google Maps,” *Jurnal Interkom: Jurnal Publikasi Ilmiah Bidang Teknologi Informasi dan Komunikasi*, vol. 16, no. 4, pp. 28–38, Jan. 2022.
- [18] I. N. Husada and H. Toba, “Pengaruh Metode Penyeimbangan Kelas Terhadap Tingkat Akurasi Analisis Sentimen pada Tweets Berbahasa Indonesia,” *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 6, no. 2, Aug. 2020.
- [19] A. F. Anjani, D. Anggraeni, and I. M. Tirta, “Implementasi Random Forest Menggunakan SMOTE untuk Analisis Sentimen Ulasan Aplikasi Sister for Students UNEJ,” *Jurnal Nasional Teknologi dan Sistem Informasi*, vol. 9, no. 2, pp. 163–172, Sep. 2023.
- [20] Bagus Muhammad Akbar, Ahmad Taufiq Akbar, and Rochmat Husaini “Analisis sentimen dan emosi vaksin sinovac pada twitter menggunakan naïve bayes dan valence shifter” *Jurnal Teknologi Terpadu*, vol.7, no. 2, pp. 83–92
- [21] E. A. Putra, S. Alam, and I. Kurniawan, “Analisis Sentimen Pengguna MY JNE Menggunakan Algoritma Naïve Bayes,” *Jurnal Teknologi Informatika dan Komputer*, vol. 10, no. 2, pp. 650–666, Oct. 2024.