

ANALISIS SENTIMEN TERHADAP RESPON PERUBAHAN NAMA TWITTER MENJADI 'X' MENGGUNAKAN METODE NAÏVE BAYES CLASSIFIER

Fazrin Meila Azzahra Sofyan, Nina Sulistiyowati, Apriade Voutama

Sistem Informasi, Universitas Singaperbangsa Karawang

Jl. HS.Ronggo Waluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat, Indonesia

2010631250044@student.unsika.ac.id

ABSTRAK

Media sosial kini menjadi bagian esensial dalam kehidupan manusia untuk mendapatkan informasi, sebagaimana terlihat dari laporan We Are Social (*HootSuite*) yang mencatat ada 167 juta pengguna media sosial di Indonesia pada Januari 2023. Twitter menjadi platform yang populer saat itu, terutama setelah Elon Musk mengubah nama dan logo Twitter menjadi 'X' pada 24 Juli 2023. Perubahan ini memicu berbagai reaksi pengguna di Twitter, dengan tweet yang bervariasi dari positif, negatif, hingga netral. Untuk menganalisis sentimen tweet tersebut, digunakan algoritma *Naïve Bayes Classifier* dalam proses yang mengikuti metodologi KDD (*Knowledge Discovery in Database*), meliputi tahap *data selection*, *data preprocessing*, *data transformation*, *data mining*, dan *evaluation*. *InSet Lexicon* digunakan untuk melabeli sentimen berdasarkan kamus kata positif dan negatif, yang dibandingkan dengan pelabelan manual oleh ahli bahasa. Penelitian ini menemukan bahwa pelabelan dengan *InSet Lexicon* dan manual menghasilkan klasifikasi sentimen dan jumlah data yang berbeda. *InSet Lexicon* menghasilkan 1131 tweet negatif, 687 positif, dan 583 netral, sedangkan pelabelan manual menghasilkan 1198 tweet negatif, 94 positif, 1035 netral, dan 74 multitafsir. Evaluasi menunjukkan bahwa pelabelan dengan *InSet Lexicon* lebih baik daripada pelabelan manual, dengan *accuracy* 83%, *precision* 84%, dan *recall* 83%, sedangkan pelabelan manual menunjukkan *accuracy* 77%, *precision* 76%, dan *recall* 76%.

Kata kunci : Analisis Sentimen, *InSet Lexicon*, KDD (*Knowledge Discovery in Database*), *Naïve Bayes Classifier*, Twitter

1. PENDAHULUAN

Twitter adalah platform media sosial yang memberikan sarana bagi penggunanya untuk menulis, membagikan aktivitas serta opini pengguna melalui pesan singkat yang dikenal sebagai kicauan (*tweet*) [1].

Tingginya pengguna Twitter memberikan kebebasan untuk memberikan komentar maupun opini mengenai isu-isu atau berita yang sedang menjadi perbincangan hangat [2].

Salah satu topik yang tengah menjadi perhatian utama kalangan pengguna Twitter saat ini yaitu mengenai perubahan (*rebranding*) terhadap Twitter menjadi 'X' oleh Elon Musk, perubahan ini resmi ditetapkan pada tanggal 24 Juli 2023 dengan mengganti logo dan nama Twitter menjadi 'X' [3].

Keputusan yang diambil oleh Elon Musk memicu berbagai reaksi yang beragam, mulai dari dukungan pengguna yang menyambut positif atas perubahan nama dan logo Twitter menjadi 'X' hingga kritik yang tajam atas perubahan tersebut yang semuanya dibagikan dalam suatu unggahan tweet. Untuk memahami bagaimana perubahan ini terjadi, diperlukan analisis yang mendalam terhadap sentimen pengguna Twitter.

Analisis sentimen merupakan suatu proses yang melibatkan ekstraksi, pengolahan, dan pemahaman terhadap data berupa teks yang tidak terstruktur untuk memperoleh wawasan mengenai sentimen dari teks yang terkandung dalam suatu kalimat pendapat atau opini [4].

Algoritma yang umum digunakan dalam melakukan analisis sentimen yaitu *Naïve Bayes Classifier* karena kemampuannya dalam bekerja dengan cepat dan efisien serta mampu mencapai tingkat akurasi yang tinggi terutama ketika dihadapkan pada volume data yang besar [5].

Dalam analisis sentimen, pendekatan *lexicon* contohnya *InSet Lexicon (Indonesia Sentiment Lexicon)* dikenal sebagai metode pendekatan kamus dengan pemberian nilai pada setiap kata dalam suatu opini menurut kamus bobot penilaian positif dan negatif [6][7].

Untuk menangani data yang tidak seimbang, metode SMOTE (*Synthetic Minority Over-sampling Technique*) melakukan penyeimbangan dengan cara meningkatkan jumlah sampel dalam kelas minoritas melalui proses *resampling* [8].

Penelitian sebelumnya dilakukan oleh [9] membahas mengenai komparasi metode K-NN, *Naïve Bayes*, dan *Decision Tree* terhadap opini PT PAL Indonesia menghasilkan algoritma *Naïve Bayes* sebagai metode dengan tingkat akurasi yang tinggi sebesar 84,08%, metode K-NN memperoleh 83,38% dan *Decision Tree* sebesar 81,09%.

Penelitian selanjutnya yang dilakukan oleh [10] membahas mengenai pembatalan Indonesia menjadi tuan rumah Piala Dunia U-20 menggunakan *Naïve Bayes Classifier* dan *Lexicon Based* memperoleh akurasi sebesar 85%. Hasil penelitian oleh [11] mengenai tanggapan pengguna Twitter terhadap kinerja walikota Medan menggunakan algoritma

Naïve Bayes dan penerapan SMOTE menghasilkan akurasi sebesar 78% pada rasio data latih dan uji 80:20.

Berdasarkan latar belakang di atas, penelitian ini bertujuan untuk menentukan opini pengguna Twitter terhadap respon perubahan nama Twitter menjadi 'X' dengan melakukan analisis sentimen dan mengelompokkan opini tersebut berdasarkan kelasnya yaitu positif, negatif dan netral menggunakan algoritma *Naïve Bayes Classifier* dan penerapan *InSet Lexicon* serta SMOTE untuk meningkatkan akurasi dan keseimbangan data.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis sentimen diterapkan untuk mengelompokkan atau mengklasifikasikan dokumen, pendapat atau ulasan yang dikemukakan terkait dengan suatu topik [6].

Analisis sentimen bisa diklasifikasikan menjadi sentimen positif, negatif, atau netral [12].

- Sentimen Positif**
Sentimen positif adalah sikap yang memberikan nilai lebih pada seseorang atau sesuatu.
- Sentimen Negatif**
Sentimen negatif adalah respon atau sikap yang bisa mengurangi nilai seseorang atau sesuatu, sehingga membuat tren turun. Contoh sentimen negatif ditandai oleh penggunaan kata-kata negatif untuk mengubah arah pernyataan.
- Sentimen Netral**
Sentimen netral merujuk pada reaksi yang tidak bersifat mendukung atau memihak. Kalimat sentimen netral ini tidak menunjukkan ekspresi positif atau negatif secara spesifik.

2.2. Naïve Bayes Classifier

Naïve Bayes Classifier ialah metode untuk mengklasifikasikan data menerapkan probabilitas sederhana dengan menggunakan Teorema Bayes dan mengasumsikan tingkat independensi yang tinggi. Metode ini sangat sesuai dalam menangani data dengan kinerja yang cepat untuk klasifikasi data dengan akurasi yang tinggi. Dasar dari penggunaan *Naïve Bayes* adalah prinsip dasar teorema Bayes, di mana peluang terjadinya suatu kejadian A dengan kondisi B dihitung dari peluang kejadian B saat kondisi A terjadi, peluang kejadian A, dan peluang kejadian B [13] yang ditunjukkan pada persamaan berikut.

$$P(A|B) = \frac{P(B|A) \times P(A)}{P(B)} \quad (1)$$

Pada pengaplikasiannya rumus tersebut berubah menjadi:

$$P(C_i|D) = \frac{P(D|C_i) \times P(C_i)}{P(D)} \quad (2)$$

2.3. InSet Lexicon (Indonesia Sentiment Lexicon)

InSet Lexicon adalah sebuah kamus sentimen berbahasa Indonesia yang melibatkan 3.609 kata dengan makna positif dan 6.609 kata dengan makna negatif, kamus ini dapat diakses melalui GitHub. Setiap kata dalam kamus ini memiliki nilai atau polarity score yang berkisar antara -5 sampai +5 [7].

2.4. KDD (Knowledge Discovery in Database)

Knowledge Discovery in Database (KDD) diartikan sebagai kegiatan pengambilan informasi yang berpotensi, tersirat, dan belum diketahui dari suatu kumpulan data. Tahap *knowledge discovery* melibatkan output dari proses *Data Mining* yang selanjutnya diubah dengan akurat menjadi informasi yang dapat dipahami secara mudah [14]. Proses KDD secara umum memiliki 5 tahap yaitu *data selection*, *data preprocessing*, *data transformation*, *data mining*, dan *evaluation* [15].

2.5. SMOTE (Synthetic Minority Over-sampling Technique)

SMOTE merupakan sebuah teknik untuk menangani *imbalance* data pada kelas sentimen. Dalam SMOTE, pendekatan ini bekerja dengan melakukan *oversampling* pada kelas minoritas dengan cara mensintesis data baru guna mengimbangi data [16].

2.6. TF-IDF (Term Frequency-Inverse Document Frequency)

TF-IDF adalah metode untuk menentukan sejauh mana suatu *term* mencerminkan konten dalam suatu dokumen dengan memberi skor pada masing-masing kata yang ada, nilai TF-IDF diperoleh melalui perkalian nilai TF (*Term Frequency*) dengan nilai IDF (*Inverse Document Frequency*) [17] yang ditunjukkan pada persamaan berikut.

2.7. Perhitungan Nilai TF

$$TF_{t,d} = f_{(t,d)} \quad (3)$$

Keterangan:

TF = kata muncul dalam suatu dokumen teks
f = Jumlah kata dalam dokumen tersebut

2.8. Perhitungan Nilai IDF

$$IDF = \log \frac{N}{DF_t} \quad (4)$$

Keterangan:

N = Banyaknya dokumen
DF_t = Nilai TF yang diperoleh

2.9. Perhitungan Nilai TF-IDF

$$Wt = TF \times IDF \quad (5)$$

Keterangan:

W_t = TF-IDF
TF = Nilai TF yang diperoleh
IDF = Nilai IDF yang diperoleh

2.10. Crawling Data

Crawling adalah metode untuk mengambil data pada situs web, proses *crawling* dilakukan secara otomatis di mana informasi yang diambil sesuai dengan kata kunci yang telah ditentukan oleh pengguna. Hasil dari proses *crawling* disimpan dalam file atau database yang telah disiapkan sebelumnya [18].

2.11. Word Cloud

Word cloud ialah teknik visualisasi data yang memperlihatkan frekuensi kemunculan kata-kata dalam dataset dengan menggambarkan ukuran tulisan yang berkaitan pada *word cloud* sesuai dengan frekuensi masing-masing kata. Semua kata dalam dataset dikumpulkan berdasarkan frekuensinya lalu diatur untuk membentuk tampilan yang menyerupai awan, di mana ukuran teks mencerminkan frekuensi kemunculan kata tersebut [19].

2.12. Confusion Matrix

Confusion matrix dimanfaatkan untuk mengevaluasi performa dari metode klasifikasi. Umumnya *confusion matrix* memuat informasi perbandingan hasil dari klasifikasi sistem dengan hasil klasifikasi seharusnya. Terdapat empat istilah yang mencerminkan output dari proses klasifikasi yaitu *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) [20].

		Nilai sebenarnya	
		TRUE	FALSE
Nilai prediksi	TRUE	TP (True Positive) <i>Correct result</i>	FP (False Positive) <i>Unexpected result</i>
	FALSE	FN (False Negative) <i>Missing result</i>	TN (True Negative) <i>Correct absence of result</i>

Gambar 1. Confusion Matrix

Dari Gambar 1 di atas dapat digunakan untuk menghitung nilai *accuracy*, *precision* dan *recall*.

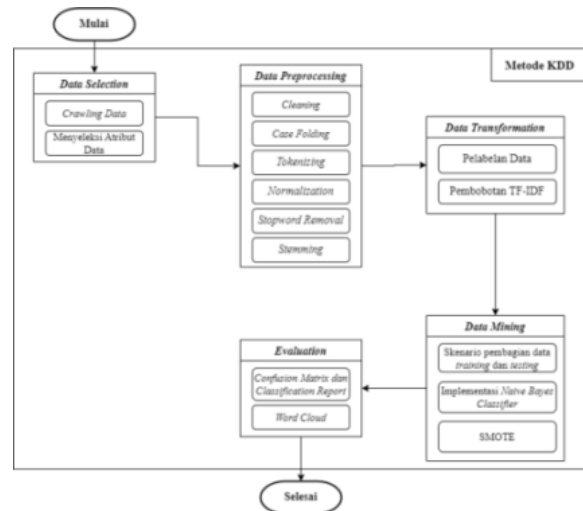
$$accuracy = \frac{TP + TN}{TP + FP + TN + FN} \times 100\% \quad (6)$$

$$precision = \frac{TP}{TP + FP} \times 100\% \quad (7)$$

$$recall = \frac{TP}{TP + FN} \times 100\% \quad (8)$$

3. METODE PENELITIAN

Metode penelitian yang digunakan ialah KDD (*Knowledge Discovery in Database*) dengan merinci tahapan-tahapan yang akan dilaksanakan sesuai pada Gambar 2.



Gambar 2. Rancangan Penelitian KDD

3.1. Data Selection

Tahap *data selection* terbagi menjadi dua yaitu *crawling* data dan menyeleksi atribut data, tahap ini dilakukan untuk pengumpulan data yang akan diproses dan di analisis pada penelitian ini.

3.2. Data Preprocessing

Tahap ini melibatkan *cleaning*, *case folding*, *tokenizing*, *normalization*, *stopword removal*, dan *stemming* untuk menghilangkan atribut yang tidak dibutuhkan pada data tweet dan memudahkan proses analisis lebih lanjut.

a. Cleaning

Cleaning menghapus elemen-elemen pada data seperti URL (<http://>), tagar (#), *username* (@), angka, spasi berlebih serta tanda baca yang bertujuan untuk memudahkan pada proses analisis.

b. Case Folding

Case folding mengubah teks pada dataset menjadi huruf kecil (*lowercase*) secara keseluruhan.

c. Tokenizing

Tokenizing melibatkan pemisahan kalimat menjadi bentuk kata per-kata untuk analisis yang lebih rinci.

d. Normalization

Tahap ini dilakukan untuk memperbaiki kata yang tidak sesuai dengan Kamus Besar Bahasa Indonesia (KBBI)

e. Stopword Removal

Tahap *stopword removal* melibatkan penghapusan kata-kata dari daftar *stopword* NLTK, seperti kata hubung 'dan', 'atau', 'yang', dan sebagainya. Selain itu, daftar *stopword* tambahan dibuat secara manual untuk mencakup kata-kata termasuk nama-nama yang tidak sesuai dengan penelitian ini dan tidak ada dalam daftar *stopword* NLTK.

f. Stemming

Tahap *stemming* bertujuan untuk mengubah kata yang memiliki imbuhan menjadi bentuk kata dasar menggunakan *library* Sastrawi.

3.3. Data Transformation

Tahap ini dilakukan pelabelan pada data menggunakan *InSet Lexicon* dan pelabelan manual serta pembobotan TF-IDF (*Term Frequency-Inverse Document Frequency*).

3.4. Data Mining

Dilakukan pembagian data menjadi data latih dan data uji dengan rasio 90:10, 80:20, 70:30, 60:40, dan 50:50. *Naïve Bayes Classifier* diterapkan untuk klasifikasi sentimen, sementara SMOTE digunakan untuk menangani ketidakseimbangan data guna meningkatkan akurasi hasil.

3.5. Evaluation

Tahap akhir adalah *evaluation* yang akan menghasilkan *confusion matrix*, *classification report* dan *word cloud* pada algoritma *Naïve Bayes* yang telah digunakan.

4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Pada tahap ini dilakukan *crawling* data serta mengeliminasi kolom atau atribut yang tidak dibutuhkan sehingga menghasilkan data yang sesuai untuk tahap selanjutnya. *Tweet-Harvest* digunakan untuk melakukan *crawling* pada data *tweet* pengguna Twitter dengan menuliskan kata kunci “twitter ganti nama” pada rentang waktu 24 Juli 2023 – 16 November 2023, data yang berhasil dikumpulkan sebanyak 2436 data.

```
# Drop kolom yang tidak dipakai
data = data.drop(columns=['created_at',
                          'id_str',
                          'quote_count',
                          'reply_count',
                          'retweet_count',
                          'favorite_count',
                          'lang',
                          'user_id_str',
                          'conversation_id_str',
                          'username',
                          'tweet_url'])

data.head()
```

Gambar 3. Proses Seleksi Atribut

Pada Gambar 3 di atas merupakan proses untuk menghapus atribut atau kolom yang tidak akan digunakan, pada penelitian ini hanya menggunakan atribut tau kolom “full_text”. Selain proses seleksi atribut terdapat proses untuk menghapus data yang duplikat, terdapat 20 data yang duplikat sehingga dari data awal yang diperoleh sebanyak 2436 data kini menjadi 2416 data.

4.2. Data Preprocessing

Terdapat beberapa tahap pada *preprocessing* yang bertujuan guna memudahkan dataset dalam pemodelan yang akan dilakukan. Tahap yang dilakukan ialah *cleaning*, *case folding*, *tokenizing*, *normalisasi*, *stopword removal* dan *stemming* yang ditunjukkan pada Tabel 1 berikut.

Tabel 1. Contoh pada *Data Preprocessing*

full_text	Cleaning	Case folding	Tokenizing	Normalization	Stopword Removal	Stemming
Pagi ini terlihat logo dan nama twitter udah ganti dihandphone ku, good bye twitter 🍀	Pagi ini terlihat logo dan nama twitter udah ganti dihandphone ku good bye twitter	pagi ini terlihat logo dan nama twitter udah ganti dihandphone ku good bye twitter	['pagi', 'ini', 'terlihat', 'logo', 'dan', 'nama', 'twitter', 'udah', 'ganti', 'dihandphone', 'ku', 'good', 'bye', 'twitter']	['pagi', 'ini', 'terlihat', 'logo', 'dan', 'nama', 'twitter', 'sudah', 'ganti', 'di ponsel', 'ku', 'selamat', 'tinggal', 'twitter']	['pagi', 'logo', 'nama', 'twitter', 'ganti', 'di ponsel', 'selamat', 'tinggal', 'twitter']	['pagi', 'logo', 'nama', 'twitter', 'ganti', 'di ponsel', 'selamat', 'tinggal', 'twitter']
...	...	...	...	...	...	...
ini twitter nama appsnya ganti jd X" doang?"	ini twitter nama appsnya ganti jd X doang	ini twitter nama appsnya ganti jd x doang	['ini', 'twitter', 'nama', 'appsnya', 'ganti', 'jd', 'x', 'doang']	['ini', 'twitter', 'nama', 'aplikasinya', 'ganti', 'jadi', 'x', 'saja']	['twitter', 'nama', 'aplikasinya', 'ganti', 'x']	['twitter', 'nama', 'aplikasi', 'ganti', 'x']

Dari Tabel 1 di atas menjelaskan proses data *tweet* pengguna Twitter melalui beberapa proses seperti *cleaning*, *case folding*, *tokenizing*, *normalisasi*, *stopword removal* dan *stemming* untuk menghasilkan data bersih yang akan memudahkan dalam proses selanjutnya.

4.3. Data Transformation

Tahap *data transformation* dilakukan untuk memberikan label sentimen menggunakan *InSet Lexicon* dan memberikan label secara manual yang telah divalidasi oleh ahli bahasa kemudian melakukan pembobotan TF-IDF.

Pelabelan data bertujuan untuk memberikan label positif, negatif dan netral pada setiap *tweet* pengguna berdasarkan kamus sentimen *InSet Lexicon*, kamus *InSet Lexicon* dapat diakses pada link berikut <https://github.com/fajri91/InSet>.

Untuk pemberian label secara manual dilakukan dengan menggunakan dataset hasil *crawling* data sebanyak 2436 data yang kemudian divalidasi oleh Ibu Dian Hartati, SS., M.Pd. selaku Dosen Pendidikan Bahasa Indonesia di Fakultas Keguruan dan Ilmu Pendidikan (FKIP) Universitas Singaperbangsa Karawang.

Pada saat proses validasi, Ibu Dian Hartati, SS., M.Pd. menemukan beberapa *tweet* pengguna mengarah pada kalimat “multitafsir” sehingga terdapat empat label pada dataset pelabelan manual yaitu positif, negatif, netral dan multitafsir. Berikut Tabel 2 yang menunjukkan perbandingan kelas sentimen berdasarkan *Inset Lexicon* dan ahli bahasa.

Tabel 2. Perbandingan Sentimen Pelabelan InSet Lexicon dan Manual

Kalimat	Pelabelan	
	InSet Lexicon	Ahli Bahasa
Pagi ini terlihat logo dan nama twitter udah ganti dihandphone ku, good bye twitter 📱	Positif	Multitafsir
@marklde Ganti logo doang ya? Tapi nama apk masi twitter?	Netral	Netral
@Arlnaaaa Icon doank yang ganti, nama tetap twitter	Netral	Negatif
ini twitter nama appsnya ganti jd X" doang?"	Negatif	Netral

Setelah tahap pelabelan selesai, langkah berikutnya ialah menghitung bobot setiap kata menggunakan TF-IDF dengan memanfaatkan *TfidfVectorizer* pada proses ini untuk mengubah data *tweet* menjadi bentuk numerik seperti pada Gambar 4 berikut.

Menampilkan TFIDF sample ke-0

pagi logo nama twitter ganti ponsel selamat tinggal twitter

Term	TF	IDF	TF-IDF
ganti	0.111	1.066	0.118
logo	0.111	2.760	0.307
nama	0.111	1.067	0.119
pagi	0.111	5.526	0.614
ponsel	0.111	4.341	0.482
selamat	0.111	5.257	0.584
tinggal	0.111	5.173	0.575
twitter	0.222	1.061	0.236

Gambar 4. TF-IDF InSet Lexicon dan Label Manual

Gambar 4 merupakan hasil TF-IDF pada dokumen dengan indeks ke-0. Hasil dari nilai TF-IDF tersebut dapat berupa sebuah matriks dengan ukuran 1000×2401 yang ditunjukkan pada Gambar 5.

	D1	D2	D3	D4	D5	D6	D7	D8	D9	D10	...	D2392	D2393	D2394	D2395	D2396	D2397	D2398	D2399	D2400	D2401
abad	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
abadi	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
abai	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
abjad	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
acak	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
yudisiam	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
zafira	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
zaman	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
zoom	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
zuckerberg	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

1000 rows × 2401 columns

Gambar 5. Matriks TF-IDF

Pada Gambar 5 menunjukkan hasil dari proses pembobotan TF-IDF menghasilkan sebuah matriks dua dimensi, hasil matriks TF-IDF dari data *tweet* pengguna terhadap respon perubahan nama Twitter menjadi “X” memiliki ukuran 1000×2401 yang diartikan bahwa terdapat sebanyak 1000 token (fitur) dan 2401 dokumen data *tweet* pengguna Twitter yang dihasilkan pada pembobotan kata TF-IDF.

4.4. Data Mining

Tahap ini dilakukan skenario pembagian data, diikuti dengan implementasi *Naïve Bayes Classifier* untuk klasifikasi sentimen serta penerapan SMOTE guna mengatasi ketidakseimbangan data. Proses pelabelan data menggunakan metode Inset Lexicon dan label manual akan diuji dengan variasi pembagian data (90:10, 80:20, 70:30, 60:40, dan 50:50) serta penerapan metode TF-IDF (dengan dan tanpa TF-IDF) sehingga diperoleh skenario pembagian data seperti pada Tabel 3.

Tabel 3. Hasil Akurasi Dengan dan Tanpa TF-IDF

Label	Skenario	Rasio	TF-IDF	Pembagian Data		Akurasi
				Data Training	Data Testing	
InSet Lexicon	Skenario 1	90:10	Tidak Pakai	2160	241	77%
	Skenario 2	80:20		1920	481	74%
	Skenario 3	70:30		1680	721	74%
	Skenario 4	60:40		1440	961	66%
	Skenario 5	50:50		1200	1201	66%
	Skenario 1	90:10	Pakai	2160	241	78%
	Skenario 2	80:20		1920	481	76%
	Skenario 3	70:30		1680	721	76%
	Skenario 4	60:40		1440	961	75%
	Skenario 5	50:50		1200	1201	74%

Label	Skenario	Rasio	TF-IDF	Pembagian Data		Akurasi
				Data Training	Data Testing	
Manual	Skenario 1	90:10	Tidak Pakai	2160	241	65%
	Skenario 2	80:20		1920	481	68%
	Skenario 3	70:30		1680	721	68%
	Skenario 4	60:40		1440	961	67%
	Skenario 5	50:50		1200	1201	69%
	Skenario 1	90:10	Pakai	2160	241	69%
	Skenario 2	80:20		1920	481	70%
	Skenario 3	70:30		1680	721	70%
	Skenario 4	60:40		1440	961	70%
	Skenario 5	50:50		1200	1201	71%

Tabel 3 di atas merupakan skenario pembagian data pada label *InSet Lexicon* dan manual tanpa TF-IDF maupun menggunakan TF-IDF. *InSet Lexicon* memperoleh hasil *accuracy* tertinggi pada *scenario* 1 (90:10) yaitu 77% (Tanpa TF-IDF) dan 78% (menggunakan TF-IDF), sedangkan untuk label manual memperoleh hasil *accuracy* tertinggi pada *scenario* 5 (50:50) yaitu 69% (Tanpa TF-IDF) dan 71% (menggunakan TF-IDF). Untuk meningkatkan hasil akurasi dapat menggunakan SMOTE (*Synthetic*

Minority Oversampling Technique), SMOTE bertujuan untuk mengatasi ketidakseimbangan data pada kelas sentimen yang sudah diberikan sebelumnya dari hasil *InSet Lexicon* dan label manual agar hasil akurasi pemodelan menggunakan algoritma *Naïve Bayes Classifier* jauh lebih baik dibandingkan sebelum menggunakan SMOTE. Berikut perbandingan akurasi sebelum dan sesudah menggunakan SMOTE yang ditunjukkan pada Tabel 4 di bawah.

Tabel 4. Hasil Akurasi Setelah Penerapan SMOTE

Skenario	Pelabelan	Tanpa TFIDF	Naïve Bayes + TFIDF	Naïve Bayes + TFIDF + SMOTE	Pengaruh SMOTE
90:10	Inset Lexicon	0,77	0,78	0,83	↑ 0,05
	Manual	0,65	0,69	0,74	↑ 0,05
80:20	Inset Lexicon	0,74	0,76	0,83	↑ 0,07
	Manual	0,68	0,70	0,77	↑ 0,07
70:30	Inset Lexicon	0,74	0,76	0,82	↑ 0,06
	Manual	0,68	0,70	0,76	↑ 0,06
60:40	Inset Lexicon	0,66	0,75	0,80	↑ 0,05
	Manual	0,67	0,70	0,75	↑ 0,05
50:50	Inset Lexicon	0,66	0,74	0,79	↑ 0,05
	Manual	0,69	0,71	0,73	↑ 0,02

Tabel 4 di atas merupakan asil akurasi setelah menerapkan SMOTE, hasil akurasi yang baik terdapat pada skenario 2 (80:20) dengan akurasi sebesar 83% untuk *InSet Lexicon* dan 77% untuk label manual.

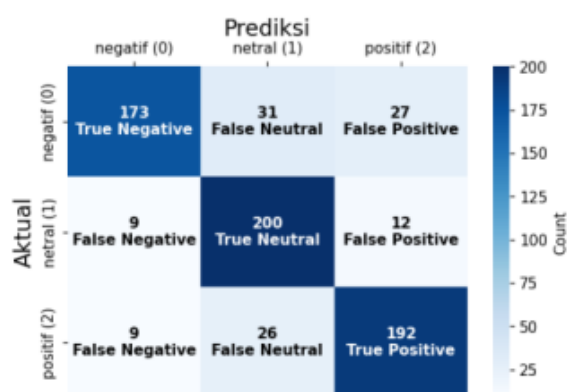
4.5. Evaluation

Tahap ini merupakan hasil dari keempat skenario *data mining* yang telah diproses sebelumnya menggunakan algoritma *Naïve Bayes Classifier* yang akan menghasilkan *confusion matrix*, *classification report* seperti nilai *accuracy*, *precision* dan *recall* serta *word cloud* setiap sentimen. Berikut *confusion matrix* dan *classification report* pada skenario 2 (80:20) pada *InSet Lexicon* dan label manual.

4.6. InSet Lexicon

Gambar 6 merupakan *confusion matrix* pada skenario 2 (80:20) pada pelabelan *InSet Lexicon*, untuk mengetahui hasil *accuracy* pada *confusion matrix* tersebut dilakukan perhitungan untuk mencari nilai *accuracy* dengan rumus sebagai berikut.

Inset Lexicon TF-IDF+SMOTE



Gambar 6. Confusion Matrix InSet Lexicon

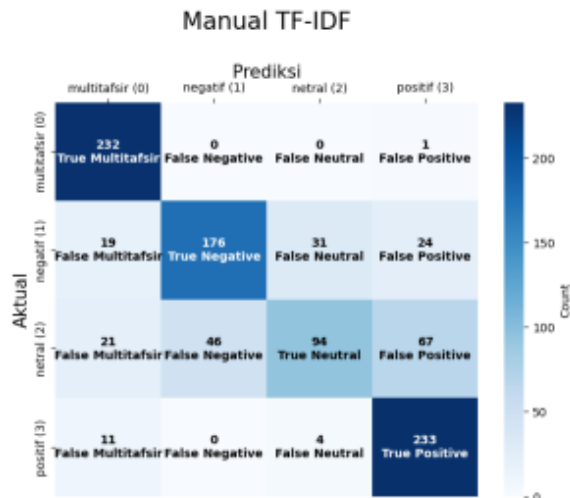
$$\begin{aligned}
 \text{Accuracy} &= \frac{TP + TN + Tnet}{TP + TN + Tnet + FP + FN + FNet} \times 100\% \\
 &= \frac{192 + 173 + 200 + (27 + 12) + (9 + 9) + (31 + 26)}{679} = 0,83 \times 100\% = 83\%
 \end{aligned}$$

Nilai *accuracy* yang diperoleh menggunakan pelabelan *InSet Lexicon* sebesar 83%, hal itu sesuai dengan *classification report* yang dihasilkan seperti pada Gambar 7 berikut.

	precision	recall	f1-score	support
negatif	0.91	0.75	0.82	231
netral	0.78	0.90	0.84	221
positif	0.83	0.85	0.84	227
accuracy			0.83	679
macro avg	0.84	0.83	0.83	679
weighted avg	0.84	0.83	0.83	679

Gambar 7. *Classification Report InSet Lexicon*

4.7. Manual



Gambar 8. *Confusion Matrix* Label Manual

Gambar 8 merupakan *confusion matrix* pada skenario 2 (80:20) pada pelabelan manual, untuk mengetahui hasil *accuracy* pada *confusion matrix* tersebut dilakukan perhitungan untuk mencari nilai *accuracy* dengan rumus sebagai berikut.

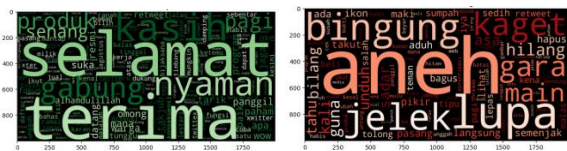
$$\begin{aligned} \text{Accuracy} &= \frac{TP + TN + T_{\text{net}} + TM}{TP + TN + T_{\text{net}} + TM + FP + FN + F_{\text{Net}} + FM} \\ &= \frac{233 + 176 + 94 + 232}{233 + 176 + 94 + 232 + (19 + 21 + 11) + (0 + 46 + 0) + (0 + 31 + 4) + (1 + 24 + 67)} \\ &= \frac{105}{137} = 0,766 \approx 0,77 \times 100\% = 77\% \end{aligned}$$

Nilai *accuracy* yang diperoleh menggunakan pelabelan manual sebesar 77%, hal itu sesuai dengan *classification report* yang dihasilkan seperti pada Gambar 9 berikut.

	precision	recall	f1-score	support
multitafsir	0.82	1.00	0.90	233
negatif	0.79	0.70	0.75	250
netral	0.73	0.41	0.53	228
positif	0.72	0.94	0.81	248
accuracy			0.77	959
macro avg	0.76	0.76	0.75	959
weighted avg	0.76	0.77	0.75	959

Gambar 9. *Classification Report* Label Manual

Dari hasil *accuracy* yang telah diperoleh menggunakan dua pelabelan yaitu *InSet Lexicon* dan Manual, *InSet Lexicon* memiliki *accuracy* yang lebih unggul dibandingkan dengan label manual. *InSet Lexicon* memperoleh *accuracy* sebesar 83%, *precision* 84% dan 83% *recall*. Sedangkan label manual memperoleh hasil *accuracy* sebesar 77%, *precision* 76% dan *recall* 76%. Setelah hasil *accuracy* diperoleh langkah selanjutnya adalah menampilkan *word cloud* pada masing-masing sentimen.



Gambar 10. *Word Cloud* Positif dan Negatif



Gambar 11. *Word Cloud* Netral dan Multitafsir

Pada Gambar 10 di atas, beberapa pengguna merasa senang dan menerima atas perubahan nama Twitter menjadi 'X'. Namun, sebagian pengguna Twitter lainnya menolak atas perubahan nama tersebut. Sedangkan pada Gambar 11 di atas, beberapa pengguna bersikap netral atas perubahan nama Twitter menjadi 'X'. Namun, ada juga pengguna bersikap ambigu seperti mendukung atas perubahan nama Twitter menjadi 'X' namun terdapat penolakan pada tweet yang ditulis.

5. KESIMPULAN DAN SARAN

Penerapan algoritma *Naive Bayes Classifier* menghasilkan klasifikasi sentimen positif, negatif dan netral pada label InSet Lexicon dan label manual dengan tambahan sentimen multitafsir. Selain itu, performa dari algoritma *Naive Bayes Classifier* menunjukkan kinerja yang sangat baik dalam melakukan pemodelan menggunakan TF-IDF (dengan dan tanpa TF-IDF) serta penerapan SMOTE. Hasil terbaik berada pada skenario 2 (80:20). Hasil evaluasi performa algoritma *Naive Bayes Classifier* pada pelabelan *InSet Lexicon* dengan menerapkan SMOTE menghasilkan nilai *accuracy* 83%, *precision* 84% dan *recall* 83%. Sedangkan pada label manual menghasilkan nilai *accuracy* 77%, *precision* 76% dan *recall* 76%.

Saran untuk penelitian selanjutnya adalah metode klasifikasi lain seperti *Support Vector Machine*, *Random Forest*, *K-Nearest Neighbor*, dan *Decision Tree*. Sehingga hasil pengujian model dapat ditingkatkan dan menghasilkan akurasi yang lebih tinggi maupun menggunakan sumber media sosial lainnya seperti YouTube, Instagram dan lain-lain.

DAFTAR PUSTAKA

- [1] A. Husnusyifa, "Pengaruh Penggunaan Media Sosial Twitter Terhadap Sikap Fanatisme Penggemar (Studi Pada Media Sosial Twitter @BTOBIndonesia Terhadap Sikap Fanatisme Penggemar)," *Idea J. Hum.*, pp. 120–133, 2019, doi: 10.29313/idea.v0i0.4935.
- [2] M. F. Fahrezi and A. A. Permana, "Sentimen Analisis Opini Masyarakat Pada Sosial Media Twitter Terhadap Organisasi Aksi Cepat Tanggap Menggunakan Naïve Bayes Classifier," *JT J. Tek.*, vol. 11, no. 02, pp. 113–121, 2022, [Online]. Available: <http://jurnal.umt.ac.id/index.php/jt/index>
- [3] B. Fallahnda, "Kenapa Twitter Ganti Logo X dan Apa Artinya Menurut Elon Musk?," 2023. <https://tirto.id/kenapa-twitter-ganti-logo-x-dan-apa-artinya-menurut-elon-musk-gNj2> (accessed Oct. 31, 2023).
- [4] P. Arsi and R. Waluyo, "Analisis Sentimen Wacana Pemindahan Ibu Kota Indonesia Menggunakan Algoritma Support Vector Machine (SVM)," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 8, no. 1, p. 147, 2021, doi: 10.25126/jtiik.0813944.
- [5] M. T. Razaq, D. Nurjanah, and H. Nurrahmi, "Analisis Sentimen Review Film Menggunakan Naive Bayes Classifier dengan Fitur TF-IDF," *e-Proceeding Eng.*, vol. 10, no. 2, pp. 1698–1712, 2023, [Online]. Available: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/19997>
- [6] M. Hamka and D. Ratna Sari, "Analisis Sentimen Dan Information Extraction Pembelajaran Daring Menggunakan Pendekatan Lexicon," *Djtechno J. Teknol. Inf.*, vol. 3, no. 1, pp. 21–32, 2022, doi: 10.46576/djtechno.v3i1.2194.
- [7] D. Musfiroh, U. Khaira, P. E. P. Utomo, and T. Suratno, "Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 1, no. 1, pp. 24–33, 2021, doi: 10.57152/malcom.v1i1.20.
- [8] N. Sulistiyowati and M. Jajuli, "Integrasi Naive Bayes Dengan Teknik Sampling Smote Untuk Menangani Data Tidak Seimbang," *Nuansa Inform.*, vol. 14, no. 1, p. 34, 2020, doi: 10.25134/nuansa.v14i1.2411.
- [9] F. S. Pattihna and H. Hendry, "Perbandingan Metode K-NN, Naive Bayes, Decision Tree untuk Analisis Sentimen Tweet Twitter Terkait Opini Terhadap PT PAL Indonesia," *JURIKOM (Jurnal Ris. Komputer)*, vol. 9, no. 2, p. 506, 2022, doi: 10.30865/jurikom.v9i2.4016.
- [10] R. Sulastiyono, A. Setiawan, and S. Nugroho, "Sentimen Analisis Pembatalan Indonesia Menjadi Tuan Rumah Piala Dunia U-20 Menggunakan Metode Naive Bayes," vol. 4, no. 4, pp. 1387–1394, 2023, doi: 10.47065/josh.v4i4.3737.
- [11] Fajar Muharram and Kana Saputra S, "Analisis Sentimen Pengguna Twitter Terhadap Kinerja Walikota Medan Menggunakan Metode Naive Bayes Classifier," *J. Sist. Inf. dan Ilmu Komput.*, vol. 1, no. 2, pp. 01–12, 2023, doi: 10.59581/jusiik-widyakarya.v1i2.17.
- [12] L. Ardiani, H. Sujaini, and T. Tursina, "Implementasi Sentiment Analysis Tanggapan Masyarakat Terhadap Pembangunan di Kota Pontianak," *J. Sist. dan Teknol. Inf.*, vol. 8, no. 2, p. 183, 2020, doi: 10.26418/justin.v8i2.36776.
- [13] D. Normawati and S. A. Prayogi, "Implementasi Naive Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," *J. Sains Komput. Inform. (J-SAKTI)*, vol. 5, no. 2, pp. 697–711, 2021, [Online]. Available: <https://ejurnal.tunasbangsa.ac.id/index.php/jsakti/article/view/369/348>
- [14] M. Jannah, "I N F O R M a T I K a Penerapan Data Mining Prediksi Nilai Ujian Nasional (Un) Siswa Smp Menggunakan Metode Naive Bayes," *J. Inform. Manaj. dan Komput.*, vol. 12, no. 2, pp. 1–6, 2020.
- [15] R. B. B. Sumantri and E. Utami, "Penentuan Status Tahapan Keluarga Sejahtera Kecamatan Sidareja Menggunakan Teknik Data Mining," *Respati*, vol. 15, no. 3, p. 71, 2020, doi: 10.35842/jtir.v15i3.375.
- [16] A. Kurniasih and K. Isyara, "Penggunaan Metode SMOTE pada Naive Bayes Gaussian untuk Klasifikasi Mahasiswa Drop Out 2 Tinjauan Pustaka," pp. 616–623, 2023.
- [17] 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.
- [18] I. F. N. Fadhilah, A. Herdiani, and W. Astuti, "Analisis Sentimen Berbasis Leksikon InSet Terhadap Partai Politik Peserta Pemilu 2019 Pada Media Sosial Twitter," *eProceedings Eng.*, vol. 6, no. 3, pp. 10397–10407, 2019, [Online]. Available: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/11216>
- [19] A. R. Randisa and A. Nurmandi, "Analisis Konten Media Sosial Twitter Sarana Pendidikan di Indonesia Study Kasus Ruang Guru," *J. Ilm. Tata Sejuta STIA Mataram*, vol. 6, no. 2, pp. 291–601, 2020, doi: 10.32666/tatasejuta.v6i2.135.
- [20] Karsito and S. Susanti, "Klasifikasi Kelayakan Peserta Pengajuan Kredit Rumah Dengan Algoritma Naive Bayes Di Perumahan Azzura Residencia," *J. Teknol. Pelita Bangsa*, vol. 9, pp. 43–48, 2019.