

Implementasi *Naive Bayes* Untuk Analisis Sentimen Pada data Twitter Tentang Isu Politik di Indonesia

Fathurrahman Putra Syah^a, Tasrif Hasanuddin^b, Nia Kurniati^c

Universitas Muslim Indonesia, Makassar, Indonesia

^a13020210266@umi.ac.id; ^btasrif.hasanuddin@umi.ac.id; ^cnia.kurniati@umi.ac.id;

Received: 21-02-2025 | Revised: 02-08-2025 | Accepted: 12-09-2025 | Published: 29-09-2025

Abstrak

Perkembangan media sosial, khususnya Twitter, telah menjadi sarana utama bagi masyarakat Indonesia dalam menyampaikan opini dan pandangan terhadap berbagai isu, termasuk politik. Analisis sentimen menjadi metode yang efektif untuk memahami kecenderungan opini publik melalui media sosial. Penelitian ini bertujuan untuk menganalisis sentimen tweet netizen terkait isu politik di Indonesia dengan menerapkan algoritma *Naïve Bayes Classifier*. Data dikumpulkan dari tweet yang membahas isu politik dan diproses melalui tahap preprocessing, seperti pembersihan data, normalisasi kata, serta penghapusan kata yang tidak relevan. Hasil penelitian menunjukkan bahwa model *Naïve Bayes* memiliki akurasi sebesar 72% setelah dilakukan optimasi, dengan kinerja yang baik dalam mendeteksi sentimen positif dan netral, namun masih menghadapi tantangan dalam mengidentifikasi sentimen negatif akibat ketidakseimbangan data. Sentimen negatif umumnya berisi kritik terhadap kebijakan pemerintah, sementara sentimen positif mencerminkan harapan dan dukungan terhadap proses politik. Kelemahan dalam penelitian ini disebabkan oleh asumsi independensi fitur dalam *Naïve Bayes* serta kompleksitas bahasa dalam tweet, seperti sarkasme dan konteks yang ambigu. Oleh karena itu, penelitian lebih lanjut disarankan untuk mengembangkan pendekatan yang lebih canggih, seperti *deep learning* atau model hibrida, guna meningkatkan akurasi klasifikasi sentimen politik di media sosial.

Kata kunci: Analisa Sentimen, Twitter, Naive Bayes

Pendahuluan

Pada saat ini perkembangan jejaring sosial sebagai alat komunikasi yang sangat populer, masyarakat menggunakan jejaring sosial sebagai salah satu media untuk berkomunikasi. Salah satunya jejaring sosial twitter, twitter digunakan sebagai sarana promosi produk, iklan, kampanye politik maupun sebagai sarana menyampaikan pendapat terkait kritik, saran, isu-isu dan opini-opini publik [1]. seringkali muncul isu-isu terkait kecurangan, polarisasi politik, dan pendapat publik yang beragam terhadap isu politik di Indonesia [2]. Twitter di anggap lebih diminati oleh para masyarakat indonesia karna dirasa lebih mudah dan simpel dalam merepresentasikan opininya [3]. Oleh karna itu analisis sentimen melalui media sosial menjadi metode yang efektif untuk mengidentifikasi pandangan dan sentimen masyarakat terkait isu politik di Indonesia [4]. Analisis sentimen media sosial adalah metode untuk memahami pendapat dan emosi individu melalui platform media sosial. Dengan metode ini, kita dapat mengidentifikasi dan menganalisis opini, tanggapan, dan ekspresi emosional terhadap topik tertentu yang dibagikan di media sosial [5].

Analisis sentimen adalah proses pengolahan data tekstual untuk memahami dan mengekstraksi informasi dari rumor kontroversial. Melalui analisis ini, kita dapat mengidentifikasi pendapat atau sentimen yang terkait dengan rumor tersebut, baik positif, negatif, atau netral. Analisis sentimen membantu memahami data tekstual dan memperoleh informasi relevan untuk mengatasi rumor kontroversial [6]. Selain itu, pengklasifikasi berbasis *Naive Bayes* juga dapat digunakan untuk menyelesaikan masalah klasifikasi dan menghasilkan solusi yang dapat diimplementasikan [7]. *Naive Bayes* adalah salah satu metode untuk mengklasifikasi data dengan probabilitas sederhana yang mengaplikasikan teorema bayes dengan karakter independen yang tinggi. Metode ini sesuai untuk banyak dataset dengan performa yang cepat dalam mengklasifikasi data dan memiliki akurasi tinggi [8]. Dalam proses analisis, Term Frequency-Inverse Document Frequency (TF-IDF) sering digunakan untuk mengubah teks menjadi representasi numerik dengan menilai seberapa penting suatu kata dalam sebuah dokumen dibandingkan dengan keseluruhan koleksi dokumen. Metode ini membantu meningkatkan akurasi model dengan menekan bobot kata-kata umum dan memberi penekanan lebih pada kata-kata yang lebih bermakna dalam konteks tertentu [6].

Beberapa penelitian membahas tentang perbandingan beberapa algoritma pada tahap *preprocessing* (*url elimination*, *word normalization*, and *stop words elimination*), pemilihan atribut (*CfsSubsetEval* dan *ClassifierSubsetEval*), klasifikasi (*Naïve bayes Classifier* dan *Support Vector Machine*), dengan menggunakan 949 tweet di kota Bandung. Hasil klasifikasi terbaik adalah 89,04%, dicapai dengan kombinasi penggunaan *word normalization* pada tahap *Preprocessing*, *CfsSubsetEval* pada tahap pemilihan atribut dan algoritma *Naïve bayes* pada tahap klasifikasi [4]. Adapun penelitian sebelumnya Analisis sentimen juga telah dipelajari di bidang ekonomi. menganalisis sentimen konsumen pada toko online JD.com. Opini konsumen ini diklasifikasikan menjadi sentimen positif, negatif, dan netral menggunakan algoritma pengklasifikasi *Naïve Bayes*. Hasil analisis menunjukkan akurasi 96,44% [9]. Penelitian selanjutnya dalam Penelitian ini membandingkan akurasi algoritma *Naïve Bayes*, *K-Nearest Neighbor*, dan *Decision Tree* dalam analisis sentimen tweet tentang kuliah daring. Hasilnya menunjukkan bahwa *Naïve Bayes* memiliki akurasi tertinggi sebesar 81,57%, sedangkan *K-Nearest Neighbor* dan *Decision Tree* memiliki akurasi yang lebih rendah, yaitu 62,10% dan 51,60%. Oleh karena itu, *Naïve Bayes* merupakan metode yang paling efektif dalam analisis sentimen ini [10].

Berdasarkan permasalahan yang telah dibahas dan temuan sebelumnya, penelitian ini berfokus pada penerapan metode *Naïve Bayes* untuk analisis sentimen terhadap data Twitter yang berkaitan dengan isu politik di Indonesia. Data penelitian akan diperoleh dari tweet yang membahas berbagai topik politik di Indonesia, kemudian dikategorikan ke dalam sentimen positif, negatif, atau netral. Metode *Naïve Bayes* dipilih karena kemampuannya dalam memproses data teks dengan efisien serta menghasilkan klasifikasi yang akurat, sehingga dapat memberikan wawasan lebih dalam mengenai opini publik terhadap isu-isu politik yang sedang berkembang.

Metode

Pendekatan *Naïve Bayes Classifier* pertama kali diperkenalkan oleh Thomas Bayes dan kemudian disempurnakan oleh Pierre-Simon Laplace pada tahun 1763. Metode ini digunakan untuk memperkirakan kemungkinan suatu peristiwa terjadi di masa depan berdasarkan data dari kejadian sebelumnya. *Naïve Bayes Classifier* didasarkan pada prinsip probabilitas dan berfungsi dengan mengklasifikasikan serta menyusun data ke dalam kategori tertentu.

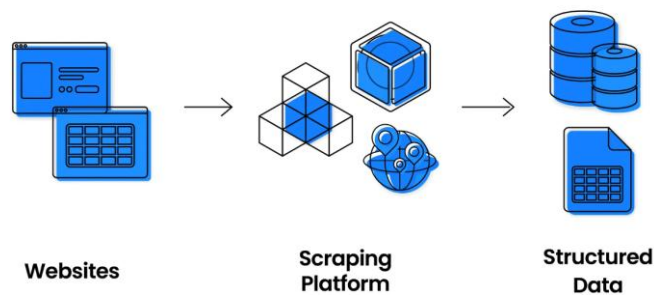
Naïve Bayes Classifier merupakan jenis pengklasifikasi berbasis probabilistik yang memanfaatkan kombinasi nilai frekuensi dari suatu dataset [11]. Dalam proses klasifikasi, *Naïve Bayes Classifier* menggunakan model statistik untuk menghitung probabilitas suatu kelas berdasarkan kelompok atribut yang tersedia, lalu menentukan kelas yang paling sesuai. Dalam metode ini, setiap atribut memiliki peran dalam pengambilan keputusan, dengan anggapan bahwa semua atribut memiliki bobot yang sama dan bersifat independen satu sama lain [12]. Berikut ini adalah bentuk umum dari teorema *Naïve Bayes Classifier* :

$$P(H|X) = \frac{(H|X).P(H)}{P(X)} \quad (1)$$

Keterangan :

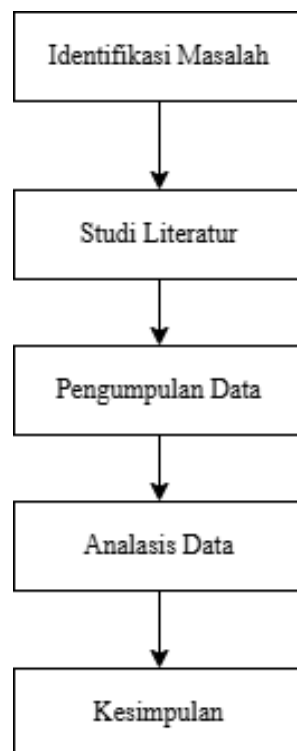
- X : Data dengan kelas yang belum diketahui
- H : Hipotesis bahwa data X termasuk dalam suatu kelas tertentu
- P(H|X) : Probabilitas hipotesis H berdasarkan kondisi X (*posteriori probability*)
- P(H) : Probabilitas hipotesis H (*prior probability*)
- P(X|H) : Probabilitas X berdasarkan kondisi pada hipotesis H
- P(X) : Probabilitas X

Web scraping adalah proses otomatisasi pengambilan data dari situs web menggunakan skrip atau program tertentu. Proses ini biasanya melibatkan pengiriman permintaan ke halaman web, mengambil kode HTML, lalu mengekstrak informasi yang relevan, seperti teks, gambar, atau tautan, dengan menggunakan teknik parsing seperti BeautifulSoup atau Scrapy dalam Python. Data yang diperoleh kemudian dapat disimpan dalam format terstruktur, seperti CSV, JSON, atau database, untuk dianalisis lebih lanjut. Web scraping sering digunakan dalam berbagai aplikasi, seperti pengumpulan data untuk analisis sentimen, pemantauan harga produk, serta penelitian tren di media sosial [13].

Gambar 1. Tahap *Web Scraping*

Long Short-Term Memory (LSTM) adalah jenis jaringan saraf tiruan yang termasuk dalam *Recurrent Neural Network* (RNN) dan dirancang untuk mengatasi masalah vanishing gradient, sehingga dapat menyimpan informasi dalam jangka waktu yang lebih panjang pada data berurutan. LSTM memiliki tiga gate utama input gate, forget gate, dan output gate yang berperan dalam mengatur aliran informasi, mempertahankan data yang penting, serta menghapus informasi yang kurang relevan. Dengan struktur ini, LSTM mampu menangani berbagai tugas pemrosesan bahasa alami (*Natural Language Processing*), seperti analisis sentimen, penerjemahan otomatis, dan prediksi teks, karena lebih efektif dalam memahami keterkaitan antar kata dalam sebuah urutan dibandingkan dengan RNN standar.

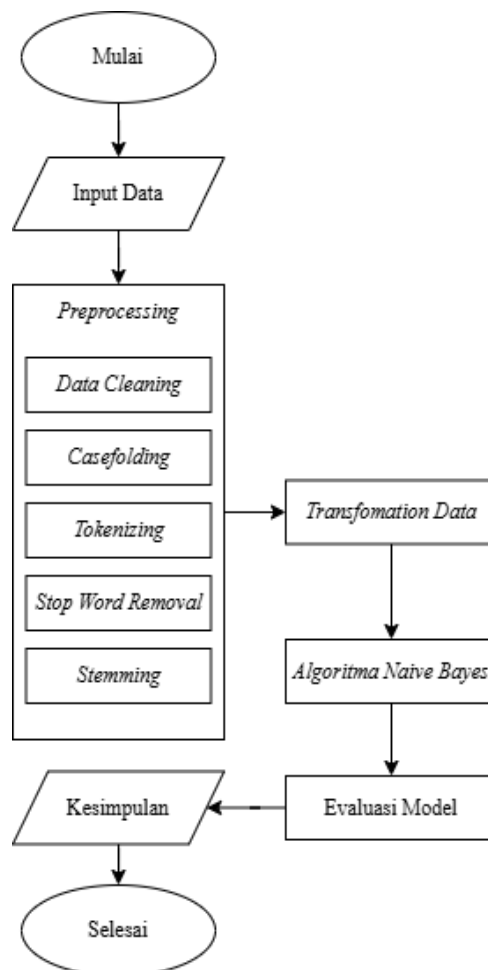
Dalam *Naïve Bayes Classifier*, X merupakan data yang kelasnya belum diketahui, sedangkan H adalah hipotesis bahwa X termasuk dalam suatu kelas tertentu. $P(H|X)$ mewakili probabilitas hipotesis H setelah mempertimbangkan data X , yang disebut *posteriori probability*. $P(H)$ adalah probabilitas awal hipotesis H sebelum mempertimbangkan data baru (*prior probability*). $P(X|H)$ menunjukkan peluang munculnya X jika hipotesis H benar, sementara $P(X)$ adalah probabilitas keseluruhan dari X dalam dataset [13]. Penulis memberikan gambaran terhadap alur penelitian yang dilakukan. Tahapan penelitian yang dibangun ditujukan pada Gambar 1.



Gambar 2. Tahapan Penelitian

Tahapan penelitian ini diawali dengan identifikasi masalah terkait analisis sentimen pada data Twitter mengenai isu politik di Indonesia, khususnya dalam memahami opini publik. Setelah itu, dilakukan studi literatur untuk meninjau penelitian terdahulu serta dasar teori yang mendukung penerapan metode *Naïve Bayes* dalam analisis sentimen. Selanjutnya, data dikumpulkan dari Twitter menggunakan teknik *web scraping*, mencakup tweet yang mengandung kata kunci terkait isu politik tertentu. Kemudian, data yang diperoleh dianalisis dengan metode *Naïve Bayes*, yang mencakup tahap pemrosesan teks, pelatihan model, serta evaluasi akurasi. Akhirnya, penelitian ini menghasilkan kesimpulan mengenai pola sentimen masyarakat terhadap isu politik serta efektivitas metode yang digunakan dalam klasifikasi sentimen.

Perancangan



Gambar 3. Desain Penelitian

A. Input Data

Tahap awal dalam penelitian ini dilakukan pengimputan data, di mana menggunakan teknik *crawling* data untuk mengumpulkan informasi yang relevan dari media sosial terkait dengan pertanyaan penelitian, terutama terkait isu politik di Indonesia. penulis mengumpulkan tweet yang berkaitan dengan topik tersebut dengan melakukan pencarian menggunakan *Twitter API*

B. Preprocessing

Labeling data adalah proses mengkategorikan data menjadi sentimen positif, netral, atau negatif. Tahap ini penting untuk melatih model *Long Short-Term Memory* (LSTM) dalam mengklasifikasikan sentimen secara akurat. LSTM adalah jenis jaringan saraf tiruan berbasis *Recurrent Neural Network* (RNN) yang mampu mengingat informasi jangka panjang. Model ini sangat efektif dalam menangani data berurutan, seperti teks dalam analisis sentimen.

C. Preprocessing Data

Tahap *preprocessing* Data bertujuan untuk menyiapkan data sebelum proses klasifikasi [14]. Langkah-langkah *preprocessing* sebagai berikut :

1. *Cleaning*

Tahapan *cleaning* dilakukan untuk menghilangkan beberapa kata atau karakter yang tidak dibutuhkan seperti mention (@), hastag (#), link dan angka. Selain itu baris baru pada setiap data akan dijadikan spasi, menghapus karakter spasi pada kiri dan kanan teks dan menghilangkan semua tanda baca [15].

2. *Case Folding*

Proses pengubahan semua huruf menjadi huruf kecil, hal ini diperlukan untuk menyeragamkan teks [16].

3. *StopWords*

Stopword sebuah teknik untuk mengurangi jumlah kata yang terdiri dari kata hubung tetapi tidak memiliki makna signifikan atau memberikan kontribusi yang kurang penting [17].

4. *Stemming*

Proses *stemming* pada tahap praprosesing data. Dimana proses ini bertujuan untuk memperkecil jumlah indeks yang berbeda dari satu data sehingga sebuah kata yang memiliki suffix maupun prefix akan kembali ke bentuk dasarnya [15].

5. *Transformation Data*

Setelah dokumen telah diproses dan dimodifikasi, tahap selanjutnya adalah mengubah teks menjadi data numerik yang dapat secara akurat mewakili dokumen tersebut. Salah satu metode pembobotan yang digunakan adalah *Term Frequency-Inverse Document Frequency (TF-IDF)*, yang memberikan nilai bagi setiap kata yang mengindikasikan seberapa sering kata kunci atau istilah lain muncul dalam dokumen tersebut [18].

6. Algoritma *Naive Bayes*

Langkah berikutnya adalah data mining, di mana data mentah diubah menjadi informasi yang berharga. Pada tahap ini, tujuannya adalah untuk menemukan dan menganalisis pola tersembunyi dalam jumlah besar data guna mendapatkan wawasan yang berguna dan bermanfaat. Dalam langkah ini, digunakan pendekatan klasifikasi untuk mengkategorikan sekumpulan sentimen ke dalam kelompok positif, negatif, dan netral. Salah satu algoritma klasifikasi yang digunakan adalah *Naive Bayes* berdasarkan teorema Bayes. Data yang telah disiapkan akan dibagi menjadi data pelatihan dan pengujian sebagai bagian dari proses klasifikasi [19].

7. Evaluasi Model

Confusion Matrix adalah sebuah tabel yang digunakan dalam machine learning untuk mengevaluasi kinerja dari suatu model klasifikasi, yang dimana memungkinkan kita untuk mengevaluasi akurasi, presisi, recall, dan F1-Score dari algoritma tersebut [20]. Akurasi dihitung sebagai rasio antara jumlah prediksi yang benar (baik positif maupun negatif) dengan keseluruhan dataset. Perhitungan akurasi dapat dilakukan menggunakan persamaan 7.

$$Accuracy = \frac{TP+TN}{TP+FP+FN+TN} \quad (2)$$

Precision merupakan rasio antara true positive (TP) dengan total prediksi positif, yang mencakup true positive (TP) dan false positive (FP). Perhitungan precision dapat dilakukan menggunakan persamaan 8.

$$Precision = \frac{TP}{TP+FP} \quad (3)$$

Recall adalah rasio antara jumlah prediksi benar positif dengan total data yang sebenarnya positif. Perhitungan recall dapat dilakukan menggunakan persamaan 9.

$$Recall = \frac{TP}{TP+FN} \quad (4)$$

F1-Score adalah nilai rerata tertimbang dari precision dan recall. Rumus dalam menghitung F1-Score bias dilakukan dalam persamaan 10.

$$F1\ Score = 2 \times \frac{precision \times recall}{precision + recall} \quad (5)$$

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Gambar 4. Confusion Matrix

Berdasarkan Gambar 4, TP (*True Positive*) adalah prediksi yang benar sebagai kasus positif, TN (*True Negative*) adalah prediksi yang benar sebagai kasus negatif, FP (*False Positive*) adalah prediksi yang dianggap positif tetapi sebenarnya negatif, dan FN (*False Negative*) adalah prediksi yang dianggap negatif tetapi sebenarnya positif merupakan prediksi kasus positif ternyata negatif, dan FN (*False Negatif*) merupakan prediksi kasus negatif ternyata positif.

Pemodelan

A. Data Scraping

Pada tahap data scraping, data dikumpulkan dari aplikasi X menggunakan API X dengan bahasa pemrograman Python. Proses ini dilakukan dengan menggunakan Visual Studio Code (VS Code) sebagai editor utama untuk mengambil data terkait ulasan pengguna mengenai hasil Pemilu Presiden 2024. Data yang diperoleh mencakup informasi nama pengguna (user), rating yang diberikan, tanggal ulasan diposting, serta isi ulasan (review). Rentang waktu pengambilan data dilakukan dari 1 Januari 2024 hingga 31 Juli 2024, dengan jumlah total 1200 baris data.

Dataset yang dikumpulkan ini akan menjadi dasar analisis sentimen dalam penelitian ini, yang bertujuan untuk memahami bagaimana persepsi publik terhadap hasil Pemilu Presiden 2024. Metode yang digunakan dalam analisis ini adalah *Naive Bayes Classifier* (NBC) untuk mengkategorikan sentimen menjadi positif, negatif, atau netral. Proses scraping dilakukan dengan akurasi tinggi, dengan menyesuaikan parameter pencarian seperti kata kunci dan tanggal, serta memastikan data yang diperoleh relevan dengan topik penelitian.

Tabel 1 adalah contoh 10 sampel dataset yang telah dikumpulkan dalam format tabel:

Tabel 1. Contoh Sampel Dataset.

User	Rating	Date	Review
ajibaums	1	2024-09-28 13:43:10	@Yadiesupri1 @ranikunty9966 @nanath27 Lu bilang ini reformasi?
ozhanqu	3	2024-09-29 16:29:13	@FDonghun Depresi? bkanya para calon itu yg depresi mikir strategi.
ShadowTeach	2	2024-09-27 14:10:45	Kritik suara dibungkam
politikasik	4	2024-09-26 11:50:32	Pemilu kali ini cukup transparan dibanding sebelumnya.
rakyatbicara	1	2024-09-25 19:20:17	Percuma nyoblos kalau akhirnya ada kecurangan begini!
pemilihbjak	5	2024-09-24 08:10:05	Senang akhirnya pemimpin baru terpilih dengan jujur.
suarapublik	2	2024-09-23 23:35:42	Debat capres kemarin kurang menarik, tidak ada gagasan baru.
demokrasiKita	4	2024-09-22 15:45:30	Partisipasi masyarakat meningkat, itu hal positif.

beritaPilpres	3	2024-09-21 10:15:55	Banyak berita hoaks menyebar, hati-hati menerima informasi.
suaraHati	1	2024-09-20 21:30:29	Kecewa, pemilu kali ini tidak adil!

B. *Preprocessing Data*

Pada tahap Text Transformation dalam proses Data Preprocessing, data ulasan yang diperoleh dari media sosial diproses melalui beberapa langkah transformasi teks untuk mempersiapkannya dalam analisis sentimen. Proses pertama adalah Case Folding, di mana seluruh teks diubah menjadi huruf kecil untuk memastikan konsistensi dalam pemrosesan data. Selanjutnya, dilakukan Text De-Slang, yaitu mengganti kata-kata tidak formal ke bentuk baku menggunakan kamus slang agar makna lebih akurat. Setelah itu, tahap Tokenizing diterapkan untuk memecah teks menjadi kata-kata individual, sehingga dapat dianalisis lebih lanjut.

Kemudian, pada proses Filtering, kata-kata tidak relevan seperti stopwords dihapus agar fokus hanya pada kata-kata yang memiliki makna signifikan dalam analisis. Setelah itu, dilakukan Text Stemming, yaitu mengembalikan kata-kata ke bentuk dasarnya untuk memastikan representasi data yang bersih dan konsisten.

Tabel 2 adalah contoh hasil Text Transformation dari proses Data Preprocessing:

Tabel 2. Contoh Hasil *Text Transformation*.

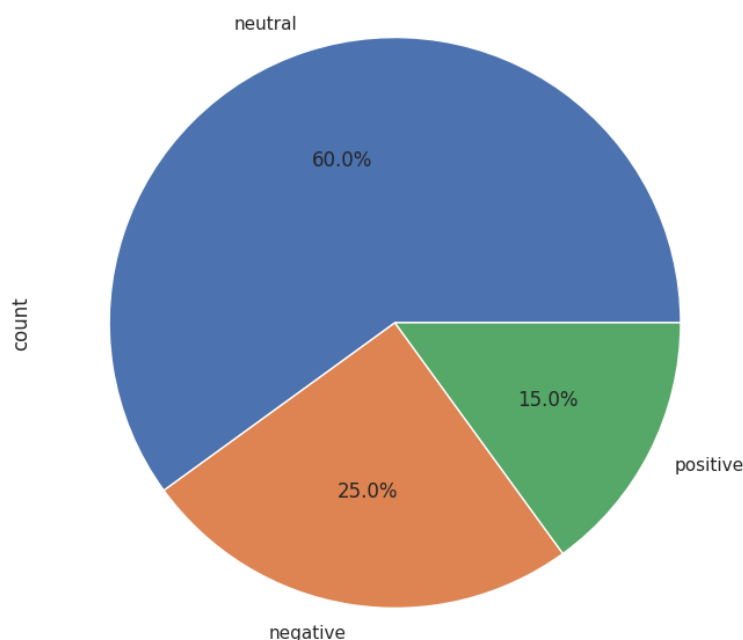
Original	Case Folding	Text De-Slang	Tokenizing	Filtering	Stemming
Reformasi birokrasi di pemerintahan masih perlu perbaikan	reformasi birokrasi di pemerintahan masih perlu perbaikan	reformasi birokrasi pemerintahan perbaikan	[reformasi, birokrasi, pemerintahan, perbaikan]	[reformasi, birokrasi, pemerintahan, perbaikan]	[reformasi, birokrasi]
Isu korupsi dalam politik semakin menjadi perhatian masyarakat	isu korupsi dalam politik semakin menjadi perhatian masyarakat	isu korupsi politik perhatian masyarakat	[isu, korupsi, politik, perhatian, masyarakat]	[isu, korupsi, politik, perhatian, masyarakat]	[isu, korupsi, politik]
Kandidat A terlihat lebih unggul dalam debat, kandidat B basis massa kuat	kandidat a terlihat lebih unggul dalam debat kandidat b basis massa kuat	kandidat terlihat unggul debat kandidat basis massa	[kandidat, a, unggul, debat, kandidat, b, basis, massa]	[kandidat, unggul, debat, kandidat, basis, massa]	[kandidat, unggul, debat]
Kandidat muda mulai mendapat perhatian, perubahan politik	kandidat muda mulai mendapat perhatian perubahan politik	kandidat muda mendapat perhatian perubahan politik	[kandidat, muda, mendapat, perhatian, perubahan, politik]	[kandidat, muda, mendapat, perhatian, perubahan, politik]	[kandidat, muda, perhatian, politik]

C. *Labeling Data*

Hasil pelabelan sentimen terhadap data Twitter yang membahas isu politik di Indonesia menunjukkan bahwa mayoritas opini yang dikumpulkan bersifat netral, dengan proporsi mencapai 60%. Hal ini mengindikasikan bahwa sebagian besar pengguna Twitter memberikan opini yang bersifat informatif atau deskriptif tanpa menunjukkan emosi yang kuat terhadap isu politik yang sedang berkembang. Banyak tweet hanya berisi pernyataan atau pengamatan tanpa kecenderungan mendukung atau menolak suatu pihak. Sementara itu, sebanyak 25% dari data yang dikategorikan memiliki sentimen negatif, yang

menunjukkan adanya kritik, ketidakpuasan, atau ketidakpercayaan terhadap berbagai isu politik di Indonesia. Ulasan negatif ini dapat mencerminkan berbagai faktor, seperti kebijakan pemerintah, dinamika partai politik, pemilihan umum, atau tokoh politik yang sedang menjadi sorotan. Sentimen negatif ini juga bisa muncul akibat ketidakpuasan publik terhadap situasi politik saat ini, baik dari segi ekonomi, sosial, maupun kebijakan yang diambil oleh pemimpin atau partai politik tertentu. Sebaliknya, hanya 15% dari total tweet yang memiliki sentimen positif, yang menunjukkan bahwa ada sebagian kecil pengguna yang memberikan dukungan terhadap suatu kebijakan, partai politik, atau tokoh tertentu. Tweet dengan sentimen ini umumnya berisi apresiasi terhadap kebijakan yang dianggap menguntungkan, optimisme terhadap pemimpin tertentu, atau dukungan terhadap partai politik yang dipercaya dapat membawa perubahan. Secara keseluruhan, hasil analisis ini menggambarkan bahwa persepsi publik terhadap isu politik di Indonesia lebih banyak bersifat netral, meskipun terdapat kecenderungan lebih banyak opini negatif dibandingkan positif. Hal ini menunjukkan bahwa masyarakat cenderung berhati-hati dalam menyampaikan opini politik mereka di media sosial, namun tetap ada kelompok yang secara aktif mengkritik atau mendukung suatu kebijakan atau tokoh politik tertentu.

Pie Chart of Sentiment Column



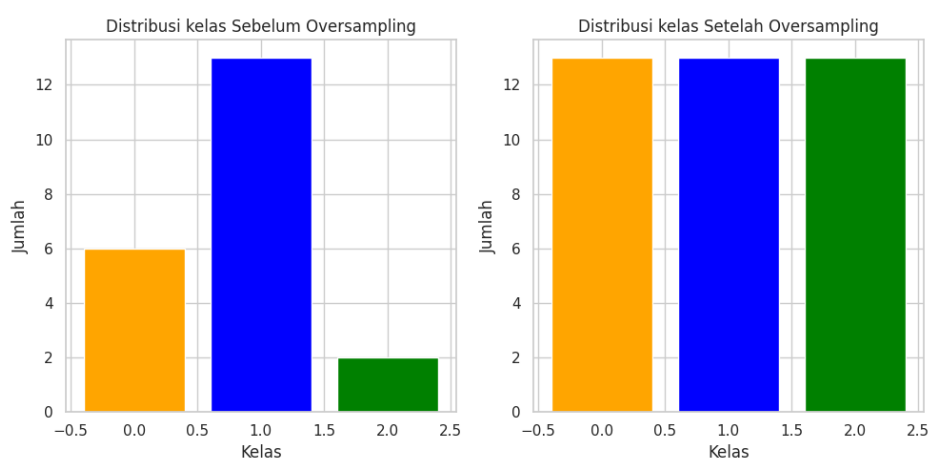
Gambar 5. Proses Labelling Data

Diagram pie chart ini menunjukkan hasil pelabelan sentimen dari tweet netizen mengenai isu politik di Indonesia yang diklasifikasikan menggunakan Naïve Bayes Classifier. Data diperoleh melalui proses scraping dari Twitter, kemudian diproses melalui tahap preprocessing sebelum dikategorikan ke dalam tiga jenis sentimen. Warna biru pada diagram mewakili sentimen netral yang mendominasi dengan persentase 60%, menunjukkan bahwa sebagian besar tweet bersifat informatif tanpa kecenderungan emosional yang kuat. Warna hijau mewakili sentimen positif sebesar 15%, yang menunjukkan adanya dukungan terhadap kebijakan, pemimpin politik, atau hasil pemilu tertentu. Sementara itu, warna oranye menunjukkan sentimen negatif sebesar 25%, yang mengindikasikan adanya kritik atau ketidakpuasan masyarakat terhadap isu politik yang dibahas. Hasil ini menunjukkan bahwa diskusi politik di Twitter lebih banyak bersifat netral, namun jumlah tweet bernada negatif lebih tinggi dibandingkan tweet positif, yang mencerminkan adanya kritik yang cukup signifikan terhadap kebijakan atau peristiwa politik yang sedang berlangsung.

D. Metode SMOTE

Hasil penerapan metode SMOTE (Synthetic Minority Over-sampling Technique) pada data sentimen

Twitter tentang isu politik di Indonesia dapat diamati dari dua grafik yang ditampilkan dalam gambar. Pada grafik Distribusi Kelas Sebelum Oversampling, terlihat bahwa jumlah sampel antar kelas tidak seimbang. Kelas tertentu memiliki jumlah data yang jauh lebih banyak dibandingkan kelas lainnya, terutama kelas dengan jumlah sampel yang paling sedikit. Ketidakseimbangan ini dapat menyebabkan bias pada model ketika melakukan analisis sentimen, di mana model akan lebih cenderung memprediksi kelas dengan jumlah data yang lebih besar dan kurang mampu mengenali pola pada kelas yang lebih kecil. Setelah menerapkan SMOTE, seperti yang terlihat pada grafik Distribusi Kelas Setelah Oversampling, jumlah sampel antar kelas menjadi lebih seimbang. SMOTE bekerja dengan cara membuat sampel sintetis berdasarkan sampel yang sudah ada dalam kelas minoritas, bukan hanya menggandakan data yang sudah ada, tetapi dengan menciptakan titik data baru yang tetap mempertahankan karakteristik kelas tersebut. Dengan demikian, model dapat belajar dengan lebih baik dan tidak terpengaruh oleh ketimpangan jumlah data antar kelas. Dari hasil ini, dapat disimpulkan bahwa penerapan SMOTE berhasil meningkatkan keseimbangan distribusi kelas, sehingga model *Naive Bayes* yang akan digunakan untuk analisis sentimen dapat melakukan klasifikasi dengan lebih akurat tanpa bias terhadap kelas mayoritas.



Gambar 6. Metode SMOTE

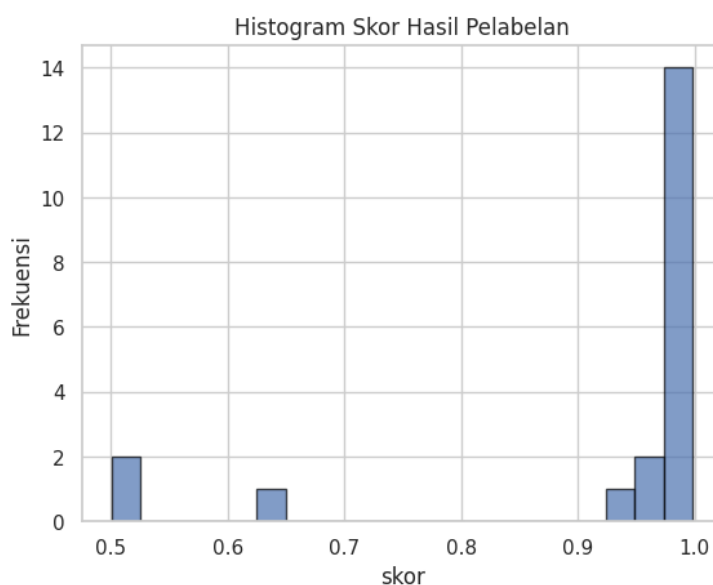
Gambar 6 menunjukkan hasil penerapan SMOTE (Synthetic Minority Over-sampling Technique) dalam menyeimbangkan jumlah data pada masing-masing kelas sentimen dalam analisis tweet mengenai isu politik di Indonesia. Diagram di sebelah kiri menggambarkan distribusi kelas sebelum dilakukan oversampling, di mana terlihat bahwa salah satu kelas memiliki jumlah data yang jauh lebih sedikit dibandingkan kelas lainnya. Kelas netral yang ditandai dengan warna biru memiliki jumlah data yang paling dominan, sementara kelas negatif yang ditandai dengan warna oranye berada di posisi kedua, dan kelas positif yang ditandai dengan warna hijau memiliki jumlah data yang paling sedikit.

Diagram di sebelah kanan menunjukkan hasil setelah dilakukan SMOTE, di mana jumlah data untuk ketiga kelas menjadi seimbang. Teknik SMOTE bekerja dengan menambahkan sampel sintetis ke kelas yang memiliki jumlah data lebih sedikit sehingga distribusi data menjadi lebih merata. Dengan adanya oversampling ini, model *Naive Bayes* yang digunakan dalam analisis sentimen dapat melakukan klasifikasi dengan lebih akurat tanpa bias terhadap kelas yang memiliki jumlah data lebih dominan. Hal ini membantu model dalam mengenali pola sentimen dari masing-masing kelas secara lebih seimbang, sehingga meningkatkan performa klasifikasi dalam analisis sentimen pada data Twitter tentang isu politik di Indonesia.

E. Skor Hasil Pelabelan

Histogram skor hasil pelabelan pada gambar menunjukkan distribusi tingkat keyakinan model dalam menentukan kategori sentimen pada data Twitter yang membahas isu politik di Indonesia. Dari grafik, terlihat bahwa mayoritas skor prediksi berada di sekitar 1.0, yang menunjukkan bahwa model memiliki tingkat kepercayaan yang tinggi dalam menentukan sentimen pada sebagian besar data. Namun, terdapat beberapa sampel dengan skor lebih rendah, sekitar 0.5 hingga 0.9, yang menandakan bahwa model mengalami sedikit ketidakpastian dalam menentukan sentimen untuk beberapa data tertentu. Distribusi ini

menunjukkan bahwa model *Naive Bayes* yang digunakan dalam analisis sentimen telah bekerja dengan cukup baik dalam memberikan prediksi dengan tingkat keyakinan yang tinggi untuk sebagian besar data. Namun, adanya beberapa skor yang lebih rendah mengindikasikan bahwa beberapa tweet memiliki karakteristik yang mungkin lebih sulit diklasifikasikan, misalnya karena konteks yang ambigu, kata-kata yang dapat memiliki makna ganda, atau perbedaan gaya bahasa pengguna Twitter. Secara keseluruhan, hasil ini mengindikasikan bahwa model dapat melakukan klasifikasi dengan cukup akurat, tetapi tetap ada beberapa data yang memerlukan perhatian lebih lanjut, seperti dengan melakukan pembersihan data tambahan atau pelabelan manual untuk meningkatkan akurasi model dalam menangani teks yang lebih kompleks.



Gambar 7. Histogram Skor Hasil Pelabelan

F. Word Cloud

Word cloud merupakan teknik visualisasi teks yang digunakan untuk menampilkan kata-kata yang paling sering muncul dalam kumpulan data teks, dengan ukuran kata yang lebih besar menunjukkan frekuensi kemunculan yang lebih tinggi. Dalam penelitian ini, word cloud digunakan untuk menganalisis data Twitter yang membahas isu politik di Indonesia, guna mengidentifikasi kata-kata yang paling dominan dalam percakapan publik di media sosial.

Dalam konteks analisis sentimen menggunakan *Naive Bayes*, word cloud dibuat berdasarkan tiga kategori utama, yaitu sentimen positif, negatif, dan netral. Word cloud dari masing-masing kategori sentimen ini membantu mengungkapkan bagaimana masyarakat mengekspresikan opini mereka terhadap berbagai isu politik, seperti pemilu, kebijakan pemerintah, partai politik, hingga tokoh politik yang menjadi sorotan publik. Dengan cara ini, perbedaan dalam pola penggunaan kata dapat diamati secara lebih jelas.

1) Word Cloud sentimen positif

Dalam word cloud ini, beberapa kata yang dominan di antaranya adalah "pimpin", "rakyat", "harap", "publik", "benar", dan "retorika". Kata "pimpin" dan "rakyat" yang muncul dengan ukuran besar mengindikasikan bahwa banyak tweet yang membahas kepemimpinan dan keterlibatan rakyat dalam politik dengan konteks positif. Hal ini dapat menunjukkan adanya dukungan publik terhadap pemimpin atau kebijakan tertentu. Kata "harap" juga sering muncul, yang bisa menggambarkan harapan masyarakat terhadap pemerintahan dan kebijakan politik yang lebih baik di Indonesia. Selain itu, kehadiran kata "publik" dan "benar" menegaskan bahwa dalam diskusi politik yang bersentimen positif, masyarakat menyoroti transparansi dan kejujuran dalam pemerintahan. Kata "retorika" yang muncul di antara kata-kata ini juga bisa mencerminkan bagaimana cara komunikasi para pemimpin politik mendapat apresiasi atau dianggap berhasil dalam membangun kepercayaan publik.

Secara keseluruhan, word cloud ini menunjukkan bahwa dalam sentimen positif, netizen cenderung membahas kepemimpinan, harapan terhadap perbaikan politik, serta bagaimana publik merespons

kebijakan yang dianggap positif atau sesuai dengan ekspektasi mereka. Visualisasi ini membantu dalam memahami pola komunikasi dalam isu politik yang memperoleh reaksi positif dari masyarakat.



Gambar 8. Word Cloud Positif

2) Word Cloud sentimen Negatif



Gambar 9. Word Cloud Negatif

Dalam word cloud ini, beberapa kata yang dominan di antaranya adalah "politik", "masyarakat", "kampanye", "berita", "kritik", "hoaks", "jelang", dan "kontroversial". Kata "politik" dan "masyarakat" yang muncul dengan ukuran besar menunjukkan bahwa banyak tweet dengan sentimen negatif membahas bagaimana masyarakat merespons berbagai isu politik di Indonesia. Kemunculan kata "kampanye" juga cukup signifikan, yang dapat menunjukkan adanya ketidakpuasan publik terhadap dinamika politik yang terjadi selama masa kampanye atau strategi yang digunakan oleh para politisi. Selain itu, kata "berita", "hoaks", dan "jelang" mengindikasikan bahwa banyak tweet bernada negatif terkait dengan penyebaran informasi yang dianggap menyesatkan atau hoaks menjelang peristiwa politik penting seperti pemilu. Kata "kontroversial" juga muncul sebagai kata yang cukup sering digunakan, menandakan bahwa ada banyak diskusi negatif mengenai kebijakan atau tindakan politik yang dianggap tidak sesuai dengan harapan publik. Secara keseluruhan, word cloud ini

menunjukkan bahwa dalam sentimen negatif, netizen cenderung membahas ketidakpercayaan terhadap politik, kritik terhadap kebijakan atau kampanye politik, serta penyebaran berita atau informasi yang dianggap tidak akurat. Visualisasi ini memberikan wawasan lebih lanjut tentang aspek-aspek politik yang sering mendapatkan reaksi negatif dari masyarakat di media sosial, yang dapat menjadi bahan analisis lebih lanjut dalam memahami persepsi publik terhadap isu politik di Indonesia.

3) Word Cloud sentimen Netral



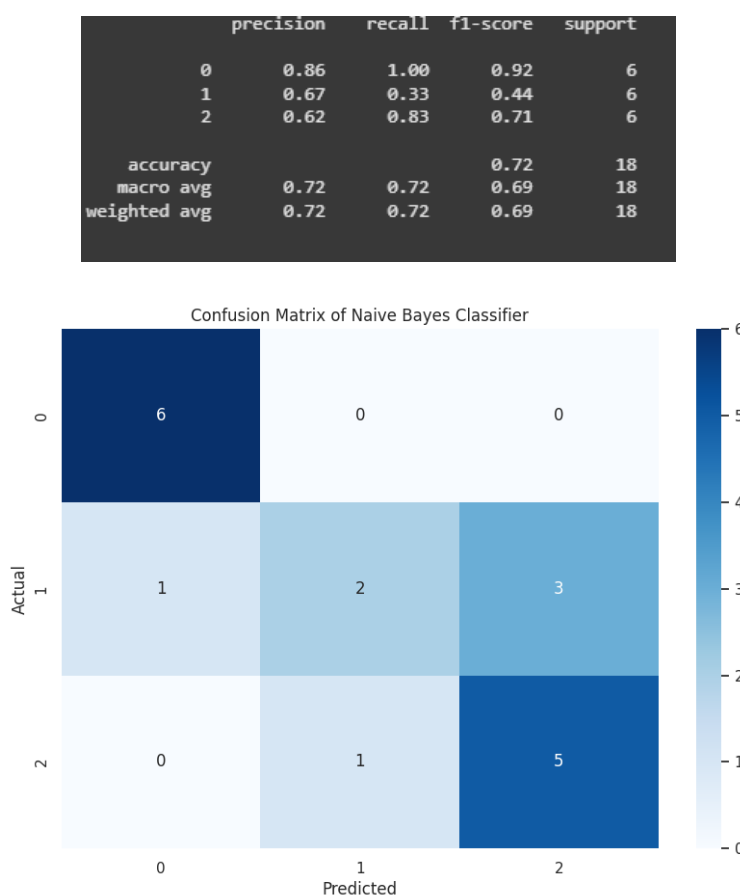
Gambar 10. Word Cloud Netral

Dalam word cloud ini, beberapa kata dominan yang muncul adalah "politik", "masyarakat", "kandidat", "pilih", "perintah", "debat", "reformasi", "birokrasi", dan "isu". Kata "politik" dan "masyarakat" menjadi yang paling sering muncul, menunjukkan bahwa diskusi netral di Twitter umumnya berfokus pada aspek umum politik tanpa ekspresi emosi yang kuat. Kata "kandidat" dan "pilih" juga muncul cukup besar, yang menandakan bahwa banyak tweet netral berisi diskusi mengenai pemilihan kandidat tanpa menunjukkan dukungan atau penolakan yang jelas. Kata "reformasi", "birokrasi", dan "perintah" juga sering muncul, menunjukkan bahwa banyak diskusi netral membahas sistem pemerintahan, kebijakan politik, dan perubahan struktural dalam politik Indonesia. Selain itu, kata "debat" sering muncul dalam konteks pemilu, yang bisa mengindikasikan pembahasan mengenai debat kandidat tanpa opini yang cenderung positif atau negatif. Secara keseluruhan, word cloud ini menunjukkan bahwa sentimen netral dalam diskusi politik di Twitter cenderung bersifat informatif atau deskriptif, dengan topik utama berkisar pada pemilihan umum, kebijakan pemerintah, dan dinamika politik tanpa adanya ekspresi emosional yang kuat. Word cloud ini memberikan wawasan mengenai bagaimana masyarakat menggunakan Twitter untuk mendiskusikan isu politik secara objektif, tanpa menunjukkan preferensi yang jelas terhadap suatu pihak.

G. Evaluasi Model

Evaluasi model menunjukkan bahwa performa klasifikasi menggunakan algoritma Naïve Bayes mengalami peningkatan setelah dilakukan optimasi. Dari confusion matrix, terlihat bahwa model mampu mengklasifikasikan sentimen netral dengan sangat baik, di mana semua sampel kelas netral berhasil diklasifikasikan dengan benar tanpa kesalahan. Untuk sentimen negatif, model masih mengalami kesulitan karena terdapat beberapa sampel yang salah diklasifikasikan sebagai sentimen positif, menunjukkan bahwa model cenderung bias terhadap sentimen yang lebih positif. Sementara itu, sentimen positif memiliki keseimbangan yang cukup baik antara precision dan recall, yang berarti model mampu mengenali dengan cukup baik namun masih mengalami beberapa kesalahan dalam klasifikasi. Hasil classification report menunjukkan bahwa akurasi model mencapai 72%, meningkat secara signifikan

dibanding sebelumnya. Precision untuk sentimen netral sangat tinggi, yang berarti prediksi sentimen ini lebih akurat dibanding kelas lainnya. Namun, recall untuk sentimen negatif masih rendah, yang menunjukkan bahwa model belum mampu menangkap semua sampel negatif dengan baik. F1-score untuk sentimen positif menunjukkan keseimbangan antara precision dan recall, sehingga model cukup stabil dalam mengenali opini positif. Secara keseluruhan, model ini cukup baik dalam menganalisis sentimen, tetapi masih perlu perbaikan terutama dalam menangani bias terhadap sentimen positif dan meningkatkan recall untuk sentimen negatif. Untuk meningkatkan akurasi lebih lanjut, model dapat diperbaiki dengan mengoptimalkan preprocessing, menyesuaikan parameter model, atau mencoba algoritma klasifikasi lain seperti SVM atau Random Forest untuk melihat apakah ada peningkatan performa yang lebih signifikan. Selain itu, penyesuaian pada representasi fitur seperti pemilihan n-gram dan stop words juga dapat membantu meningkatkan ketepatan klasifikasi, terutama dalam mendeteksi sentimen negatif yang lebih sulit diklasifikasikan dengan baik.



Gambar 11. Confusion Matrix

Kesimpulan

Penelitian ini berhasil mengklasifikasikan sentimen tweet netizen terkait isu politik di Indonesia menggunakan algoritma Naïve Bayes Classifier dengan tingkat akurasi yang meningkat hingga 72% setelah dilakukan optimasi. Model ini menunjukkan performa yang lebih baik dalam mendeteksi sentimen positif dan netral, tetapi masih memiliki kelemahan dalam mengidentifikasi sentimen negatif yang sering salah dikategorikan. Sentimen negatif umumnya berisi kritik terhadap kebijakan politik, sementara sentimen positif menggambarkan harapan dan dukungan terhadap pemilu serta kepemimpinan. Sentimen netral didominasi oleh diskusi informatif tanpa kecenderungan emosional. Kelemahan model terutama disebabkan oleh ketidakseimbangan data dan asumsi independensi fitur dalam Naïve Bayes. Meskipun hasilnya cukup baik, model ini masih perlu peningkatan, terutama dalam menangani bahasa yang kompleks dan sarkasme, dengan pendekatan lebih lanjut seperti deep learning atau hybrid model untuk meningkatkan akurasi klasifikasi.

Daftar Pustaka

- [1] A. V. Sudiantoro and E. Zuliarso, *Analisis Sentimen Twitter Menggunakan Text Mining Dengan Algoritma Naïve Bayes Classifier*. 2018.
- [2] F. I. Adiba, T. Islam, and M. S. Kaiser, "Effect of Corpora on Classification of Fake News using Naive Bayes Classifier," *Int J Auto AI Mach Learn*, vol. 1, no. 1, p. 80, 2020.
- [3] A. V. Sudiantoro and E. Zuliarso, *Analisis Sentimen Twitter Menggunakan Text Mining Dengan Algoritma Naïve Bayes Classifier*. 2018.
- [4] S. Afrizal, H. Nurramdhani Irmanda, N. Falih, and I. N. Isnainiyah, "Implementasi Metode Naïve Bayes untuk Analisis Sentimen Warga Jakarta Terhadap Kehadiran Mass Rapid Transit".
- [5] S. Yusuf, M. A. Fauzi, and K. C. Brata, "Sistem Temu Kembali Informasi Pasal-Pasal KUHP (Kitab Undang-Undang Hukum Pidana) Berbasis Android Menggunakan Metode Synonym Recognition dan Cosine Similarity," 2018. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [6] A. Aziz, "Analisis Sentimen Identifikasi Opini Terhadap Produk, Layanan dan Kebijakan Perusahaan Menggunakan Algoritma TF-IDF dan SentiStrength," Abdul Aziz, 2022.
- [7] G. N. Aulia and E. Patriya, "Implementasi Lexicon Based Dan Naive Bayes Pada Analisis Sentimen Pengguna Twitter Topik Pemilihan Presiden 2019," *Jurnal Ilmiah Informatika Komputer*, vol. 24, no. 2, pp. 140–153, 2019, doi: 10.35760/ik.2019.v24i2.2369.
- [8] D. Normawati and S. A. Prayogi, "Implementasi Naïve Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," 2021.
- [9] D. Aryanti, "Analisis Sentimen Ibukota Negara Baru Menggunakan Metode Naïve Bayes Classifier," *Journal of Information System Research (JOSH)*, vol. 3, no. 4, pp. 524–531, Jul. 2022, doi: 10.47065/josh.v3i4.1944.
- [10] E. Miana, A. Ernania, and A. Herliana, "Analisis Sentimen Kuliah Daring Dengan Algoritma Naïve Bayes, K-NN Dan Decision Tree," *Jurnal Responsif*, vol. 4, no. 1, pp. 70–80, 2022, [Online]. Available: <https://ejurnal.ars.ac.id/index.php/jti>
- [11] S. M. Fakultas, I. Komputer, and D. Anggreani, "Kinerja Metode Naïve Bayes dalam Prediksi Lama," *Prosiding Seminar Nasional Ilmu Komputer dan Teknologi Informasi*, vol. 3, no. 2, 2018.
- [12] I. B. Naive Bayes Untuk Menentukan Wadah Limbah and S. Karakteristik Gigih Putra Kawani, "Journal of Informatics, Information System, Software Engineering and Applications," vol. 1, no. 2, pp. 73–081, 2019, doi: 10.20895/INISTA.V1I2.
- [13] A. Z. Rizquina and C. I. Ratnasari, "Implementasi Web Scraping untuk Pengambilan Data Pada Website E-Commerce," *Jurnal Teknologi Dan Sistem Informasi Bisnis*, vol. 5, no. 4, pp. 377–383, Oct. 2023, doi: 10.47233/jteksis.v5i4.913.
- [14] H. Hartono, A. Hajjah, and Y. N. Marlim, "Penerapan Metode Naïve Bayes Classifier Untuk Klasifikasi Judul Berita Application Of The Naïve Bayes Classifier Method For News Title Classification," *Jurnal SimanteC*, vol. 12, no. 1, 2023.
- [15] F. Yuspriyadi, "Klasifikasi Sentimen Twitter Menggunakan Lstm," *Jurnal METHODIKA*, pp. 4–8, 2023, [Online]. Available: www.twitter.com
- [16] M. Romli, F. Kamarula, and N. Rochmawati, "Perbandingan CNN dan Bi-Lstm pada Analisis Sentimen dan Emosi Masyarakat Indonesia Di Media Sosial Twitter Selama Pandemi Covid-19 yang Menggunakan Metode Word2vec," *Journal of Informatics and Computer Science*, vol. 04, 2022, [Online]. Available: <http://ai.stanford.edu/~amaas/data/sentimen/>
- [17] M. Z. Rahman, Y. A. Sari, and N. Yudistira, "Analisis Sentimen Tweet COVID-19 menggunakan Word Embedding dan Metode Long Short-Term Memory (LSTM)," 2021. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [18] A. Azrul, A. I. Purnamasari, and I. Ali, "Analisis Sentimen Pengguna Twitter Terhadap Perkembangan Artificial Intelligence Dengan Penerapan Algoritma Long Short-Term Memory (Lstm)," 2024.
- [19] I. F. Rozi, R. Ardiansyah, and R. Ardiansyah, "Penerapan Normalisasi Kata Tidak Baku Menggunakan Levenshtein Distance pada Analisa Sentimen Layanan PT. KAI di Twitter," *Academia.edu*, 2019.
- [20] S. Puad and A. Susilo Yuda Irawan, "Analisis Sentimen Masyarakat Pada Twitter Terhadap Pemilihan Umum 2024 Menggunakan Algoritma Naïve Bayes," 2023.
- [21] N. Made *et al.*, "Analisis Sentimen Berbahasa Inggris Dengan Metode Lstm Studi Kasus Berita Online Pariwisata Bali English Sentiment Analysis Using The Lstm Method Case Study Of Bali Tourism Online News," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 11, pp. 1325–1334, 2024, doi: 10.25126/jtiik.2024118792.