

Analisis Sentimen Terhadap Pemilihan Umum Indonesia di Media Sosial Twitter dengan Metode Naïve Bayes Classifier

Refri Elamadani, Yusmet Rizal

Departemen Matematika, Universitas Negeri Padang

Article Info

Article history:

Received November 24, 2025

Revised Desember 09, 2025

Accepted Maret 05, 2026

Keywords:

General Election

Sentiment Analysis

Twitter

Naïve Bayes Classifier

Kata Kunci:

Pemilihan Umum

Analisis Sentimen

Twitter

Naïve Bayes Classifier

ABSTRACT

This research aims to analyze public sentiment regarding the election in Indonesia using the Naïve Bayes Classifier method, chosen for its speed, accuracy, and ability to classify *text* data. The results show that the choice of sentiment lexicon, such as Indonesian Sentimen (InSet), significantly affects the accuracy of the analysis. A larger data set improves the model's performance, although imbalanced class distribution remains a challenge. *Preprocessing* techniques such as tokenization with NLTK Word Tokenizer, feature extraction with CountVectorizer, and Hold-Out validation play key roles. With an accuracy of 67%, this study highlights the importance of developing an Indonesian sentiment lexicon and strategies to handle data imbalance.

ABSTRAK

Penelitian ini bertujuan menganalisis sentimen masyarakat terhadap pemilu di Indonesia menggunakan metode Naïve Bayes Classifier, yang dipilih karena kecepatan, akurasi, dan kemampuannya dalam mengklasifikasikan data teks. Hasil menunjukkan bahwa pemilihan lexicon sentimen, seperti Indonesian Sentimen (InSet), sangat memengaruhi akurasi analisis. Jumlah data yang lebih besar meningkatkan kinerja model, meskipun distribusi kelas yang tidak seimbang tetap menjadi tantangan. Teknik *Preprocessing*, seperti tokenisasi dengan NLTK Word Tokenizer, ekstraksi fitur menggunakan CountVectorizer, dan validasi metode Hold-Out, berperan penting. Dengan akurasi 67%, penelitian ini menyoroti pentingnya pengembangan kamus sentimen berbahasa Indonesia dan strategi untuk menangani ketidakseimbangan data.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Penulis Korespondensi:

Refri Elmadani

Departemen Matematika, Universitas Negeri Padang

Email: refrielmadani257@gmail.com

1. PENDAHULUAN

Indonesia menempatkan nilai-nilai demokrasi sebagai landasan penting dalam penyelenggaraan kehidupan bernegara, sebagaimana tercermin dalam sila keempat Pancasila yang berbunyi “Kerakyatan yang dipimpin oleh hikmat kebijaksanaan dalam permusyawaratan/perwakilan”. Implementasi nilai tersebut salah satunya diwujudkan melalui penyelenggaraan pemilihan umum, yang oleh masyarakat kerap dipandang sebagai bentuk pelaksanaan pesta demokrasi. [1]. Pemilu

pertama kali dilaksanakan secara langsung oleh rakyat pada tahun 2004 [2], dan menjadi mekanisme demokratis untuk memilih wakil rakyat serta pemimpin bangsa.

Indonesia melaksanakan pemilihan umum sebagai sarana untuk menetapkan Presiden dan Wakil Presiden, anggota DPR-RI, DPD-RI, DPRD, serta memilih Kepala Daerah. Sejak tahun 2019, seluruh jenis pemilihan tersebut digabungkan dan dilaksanakan secara serentak di bawah penyelenggaraan Komisi Pemilihan Umum (KPU). Pelaksanaan pemilu wajib berpedoman pada asas langsung, umum, bebas, rahasia, jujur, dan adil. Oleh karena itu, pemilu menjadi momentum penting bagi masyarakat dalam menentukan arah pembangunan nasional dan masa depan bangsa untuk periode lima tahun berikutnya.

Menjelang hari pemungutan suara, dinamika politik semakin menguat dan memunculkan berbagai sentimen masyarakat, baik terhadap kandidat tertentu maupun terhadap pemilu secara keseluruhan. Platform media sosial seperti Facebook, Instagram, dan Twitter telah menjadi wadah utama bagi masyarakat untuk mengekspresikan pandangan, pendapat, serta aspirasi mereka, karena layanan tersebut dimanfaatkan oleh beragam kelompok pengguna. [1].

Salah satu media sosial yang akrab digunakan di Indonesia adalah media sosial Twitter. Pada Twitter dikenal dengan istilah tweet. Tweet atau cuitan merupakan cara seseorang menyampaikan sesuatu baik untuk hal personal maupun publik. Jangkauan interaksi pada twitter sangat luas. Fitur twitter yang bisa memperluas jangkauan adalah hashtag. Hashtag dilambangkan dengan tanda tagar (#). Dengan tagar ini para pengguna twitter dapat mengunggah informasi yang relevan. Indonesia menjadi negara keempat terbesar pengguna twitter dengan jumlah pengguna saat ini mencapai 25,23 juta pengguna [3].

Animo yang sangat besar dari masyarakat terhadap media sosial dapat dikaitkan dengan gejolak pemilu di Indonesia. Hubungan ini dapat dimanfaatkan untuk melihat sentimen masyarakat terhadap pemilu di Indonesia. Semakin banyak pengguna maka akan semakin beragam sentimen yang muncul. Sentimen ini perlu dianalisis untuk memperoleh gambaran tentang demokrasi di Indonesia.

Analisis sentimen merupakan suatu kerangka pikiran matematis dan komputasi dengan cara mengolah informasi yang tekstual pada sebuah platform dan memperoleh suatu data kesimpulan yang menyatakan sentimen pada kalimat opini, sikap, dan emosi seseorang terhadap objek. Analisis sentiment dilakukan untuk memperoleh suatu data berharga dari dataset acak yang tidak terstruktur. Analisis ini biasanya dilakukan untuk suatu topik sensitif dengan responden yang besar. Aplikasi untuk mengukur sentimen ini juga berkembang pesat dan dimanfaatkan pada banyak keperluan dan industri [4].

Analisis sentimen atau *opinion mining* merupakan bidang dalam matematika dan informatika yang mempelajari sentimen, opini, dan emosi yang diekspresikan melalui teks. Data teks ini diolah menggunakan perangkat lunak untuk menghasilkan kelompok data melalui proses klasifikasi (Rozi dkk., 2020) [5]. Salah satu teknik yang mendukung analisis sentimen adalah *text mining*, yaitu proses semi otomatis untuk menggali informasi dan mengekstraksi pola dari data teks tidak terstruktur sehingga menghasilkan pengetahuan baru yang dapat digunakan dalam pengambilan keputusan (Indriani, 2019) [6].

Pemanfaatan *text mining* sangat luas, mulai dari data putusan pengadilan, artikel penelitian, data keuangan, hingga *review* produk di media sosial. Sebelum analisis dilakukan, diperlukan proses *crawling* untuk memperoleh data. Pada platform seperti Twitter, *crawling* dilakukan menggunakan Application Programming Interface (API) untuk mengumpulkan “cuitan” yang mengandung kata kunci tertentu sebagai data analisis (Wandani, 2021) [7].

Untuk mengklasifikasikan informasi yang diperoleh diperlukan suatu metode statistika. Salah satu metode yang paling sederhana dan umum digunakan dalam analisis sentimen adalah Naïve Bayes. Metode ini merepresentasikan probabilitas suatu kelas berdasarkan distribusi kata dalam data teks. Naïve Bayes dikenal sebagai metode yang cepat, efisien, dan mampu menghasilkan tingkat akurasi yang tinggi, sehingga banyak dimanfaatkan dalam berbagai penelitian. Meskipun demikian, pendekatan ini memiliki keterbatasan karena tidak selalu dapat memenuhi asumsi independensi antar atribut yang menjadi dasar perhitungannya. [8].

Berdasarkan paparan masalah yang telah disampaikan maka peneliti tertarik untuk mengaplikasikan ilmu matematika dalam sebuah penelitian yang berjudul “Analisis Sentimen Terhadap Pemilihan Umum Indonesia di Media Sosial Twitter Dengan Metode Naïve Bayes Classifier”. Penelitian ini akan diselesaikan dengan metode statistika dan kajian literatur yang relevan dengan judul.

2. METODE

Metode penelitian dalam studi ini disusun melalui serangkaian tahapan sistematis untuk memperoleh hasil yang terukur dan dapat dipertanggungjawabkan. Proses penelitian diawali dengan pengumpulan data, yaitu pengambilan tweet dari platform Twitter melalui Google Colab Python dengan menggunakan kata kunci “pemilu,” “pemilihan umum,” dan tagar “#pemilu2024.” Data yang diperoleh mencakup unggahan yang relevan dengan topik dan diambil dalam rentang waktu 14 Januari hingga 14 Maret. Selanjutnya, data tersebut dipilah menjadi dua bagian, yakni data latih yang digunakan untuk membangun model Naïve Bayes dan data uji yang diklasifikasikan secara manual sebagai dasar evaluasi.

Tahap berikutnya adalah studi literatur yang bertujuan memperoleh pemahaman mendalam mengenai teori serta penelitian terdahulu terkait analisis sentimen di Twitter menggunakan algoritma Naïve Bayes. Hasil kajian ini menjadi landasan dalam merancang sistem, yang meliputi struktur data, algoritma, serta alur program yang akan diimplementasikan.

Setelah rancangan sistem tersusun, tahap implementasi dilakukan dengan mengaplikasikan desain tersebut ke dalam bahasa pemrograman yang telah ditetapkan. Proses ini dilanjutkan dengan tahap pengujian dan evaluasi untuk memastikan bahwa sistem berfungsi sesuai tujuan dan untuk mengidentifikasi kemungkinan kesalahan. Preprocessing data juga dilakukan dengan menghapus stopwords dan membersihkan teks sehingga data siap untuk dianalisis.

Metode Naïve Bayes kemudian digunakan untuk menghitung bobot sentimen dan mengklasifikasikan setiap tweet ke dalam kategori positif atau negatif. Evaluasi dilakukan dengan membandingkan hasil klasifikasi otomatis dengan hasil klasifikasi manual pada data uji. Apabila tingkat akurasi model mencapai 80%, proses penelitian dinyatakan berhasil; jika belum, pengambilan dan pemrosesan data perlu diulang. Dengan pendekatan ini, penelitian diharapkan mampu menghasilkan model analisis sentimen yang memiliki akurasi optimal.

3. HASIL DAN PEMBAHASAN

3.1. Deskripsi Data

Data penelitian ini diperoleh melalui crawling di media sosial Twitter (X) dengan kata kunci “pemilu,” “pemilihan umum,” dan “#pemilu2024,” mencakup periode 14 Januari hingga 14 Maret 2024. Rentang waktu ini dianggap krusial karena mencakup masa kampanye, debat, pengumuman program politik, hingga pemungutan suara, di mana opini publik terhadap calon dan partai politik cenderung terpolarisasi. Aktivitas politik yang intens pada periode ini memungkinkan pengambilan sampel sentimen masyarakat yang relevan dan terkini. Analisis data dalam rentang ini memberikan gambaran sikap publik sebelum dan sesudah momen penting pemilu, sehingga dapat menangkap perubahan sentimen dan suasana hati masyarakat secara lebih akurat.

3.2. Crawling Data

Data tweet diperoleh melalui proses crawling menggunakan Python di Google Colab. Dengan kata kunci seperti “pemilu,” “pemilihan umum,” dan “#pemilu2024,” data diambil dalam rentang waktu 14 Januari hingga 14 Maret 2024, dengan batasan 500 tweet per kata kunci. Proses crawling ini memanfaatkan pustaka Python dan alat seperti tweet-harvest, yang memungkinkan pengumpulan data secara efisien. Data yang dihasilkan berupa kumpulan tweet dalam bahasa Indonesia, yang nantinya digunakan sebagai sampel analisis sentimen terkait isu pemilu. Berdasarkan hasil *crawling data* diperoleh 497 data tweet.

3.3. Preprocessing Data

Proses preprocessing data bertujuan untuk mempersiapkan data sebelum tahap pelabelan dan klasifikasi model. Hanya bagian “full_text” dari data yang digunakan, sedangkan elemen lain diabaikan.

3.3.1 Cleaning data

Cleaning data melibatkan penghapusan elemen tidak relevan seperti username, tautan, hashtags, emotikon, karakter khusus, angka, dan kata dengan satu huruf. Selain itu, pengecekan ejaan dilakukan untuk memastikan kualitas data.

Hasil dari proses ini adalah data yang lebih bersih, rapi, dan siap untuk analisis. Pembersihan data memastikan bahwa elemen yang tidak relevan tidak mengganggu hasil analisis atau klasifikasi model, sehingga meningkatkan akurasi dan validitas penelitian. Proses ini juga mencatat jumlah kata yang dihapus untuk mengevaluasi tingkat pembersihan data.

3.3.2 Case Folding

Langkah ini bertujuan mengonversi seluruh teks menjadi huruf kecil untuk menyederhanakan tokenisasi dan memastikan analisis teks lebih konsisten dan efisien.

	conversation_id_str	created_at	favorite_count	full_text	id_str	image_url	in_reply_to_screen_name	lang
0	1757055110279004421	Tue Feb 13 23:59:26 +0000 2024	0	https://t.co/8pTegpKvSP mengucapkan Selamat H...	1757055110279004421	https://pbs.twimg.com/media/GGQV4mGwAAAK0Bc.jpg	NaN	
1	1757055007392772440	Tue Feb 13 23:59:02 +0000 2024	12	NYUSU DULU SEBELUM DIOBLOS EH NYOBLOS I #Ratu...	1757055007392772440	https://pbs.twimg.com/ext_tw_video_thumb/17575...	NaN	
2	175755495401298559	Tue Feb 13 23:58:59 +0000 2024	12	Moa Moa Moa sudah 14 Februari nih tapi bukan V...	175755495401298559	https://pbs.twimg.com/ext_tw_video_thumb/17575...	NaN	
3	1757554936211581822	Tue Feb 13 23:58:45 +0000 2024	2	Polnet Pengarahan Pemdu 2024 Wakapolda Kampo...	1757554936211581822	https://pbs.twimg.com/media/GGQNH_bx2AEqgEL.jpg	NaN	
4	1757554845830643354	Tue Feb 13 23:58:24 +0000 2024	3	Aper Cadangan TMI Post Jaga Kondisi Rtas & K...	1757554845830643354	https://pbs.twimg.com/ext_tw_video_thumb/17575...	NaN	
499	1757401063263204748	Tue Feb 13 13:47:18 +0000 2024	0	Pemisahan surat sukur kari Pemdu 2024 dalep...	1757401063263204748	https://pbs.twimg.com/media/GGCK8TVhgAAQgHd.jpg	NaN	
499	1757401037430159952	Tue Feb 13 13:47:12 +0000 2024	0	Awan terhadap Hasi terkait Pemdu 2024 yang me...	1757401037430159952	https://pbs.twimg.com/ext_tw_video_thumb/17574...	NaN	

Gambar 1. Hasil Case Folding

3.3.3 Tokenization

Tokenisasi memecah kalimat dalam tweet menjadi kumpulan kata terpisah, seperti "Dia makan nasi" menjadi ['dia', 'makan', 'nasi'], tanpa menghapus data. Proses ini mempermudah analisis teks, seperti penghitungan frekuensi atau analisis sentimen, dan menjadi langkah awal untuk normalisasi dan stemming. Normalisasi mengubah kata tidak baku menjadi bentuk resmi, sedangkan stemming mengembalikan kata ke bentuk dasar, sehingga data lebih konsisten dan siap diolah.

3.3.4 Normalization

Normalisasi mengubah kata-kata tidak baku menjadi bentuk sesuai kaidah bahasa standar, seperti "mkn" menjadi "makan," menggunakan referensi kamus kata baku. Proses ini memastikan teks lebih terstruktur dan siap untuk analisis lanjut, seperti stemming atau analisis semantik.

3.3.5 Stemming

Stemming mengubah kata dalam teks menjadi bentuk dasar menggunakan pustaka seperti Sastrawi. Proses ini menghilangkan imbuhan untuk mengenali kata dengan makna dasar yang sama, seperti "memakan" menjadi "makan," sehingga mempermudah analisis teks seperti klasifikasi dan pencarian kata kunci.

boyan : dayan
 memara : alara
 mutan : mutan
 good : good
 luck : luck
 wara : wara
 gileya : gila
 sanayan : sanayan
 pinad : pinad
 my : my
 tweet : tweet
 narai : narai
 disuaraka : suara
 pinak : pinak
 Inyalitaa : Inyalitaa
 Tuga : Tuga
 gva : gva
 lu : lu
 ketahuan : tahu
 manajakan : tunjok
 sob : sob
 ujlan : ujlan
 idang : idang
 kersahan : rasah
 ketasuran : basah
 kekatan : kuat
 manapi : manapi
 netijen : netijen
 sks : sks
 mengapi : mengapi
 sambil : sambil
 dengerkan : denger
 tatalitaa : tatalitaa
 perjuangan : juang
 Taxi : Taxi
 wekl : wekl
 nervous : nervous
 rasaya : rasa
 tetapakoh : tetap
 perbandingan : saudara
 arja : arja
 nanatin : nanatin
 gawengalan : gawengalan

```

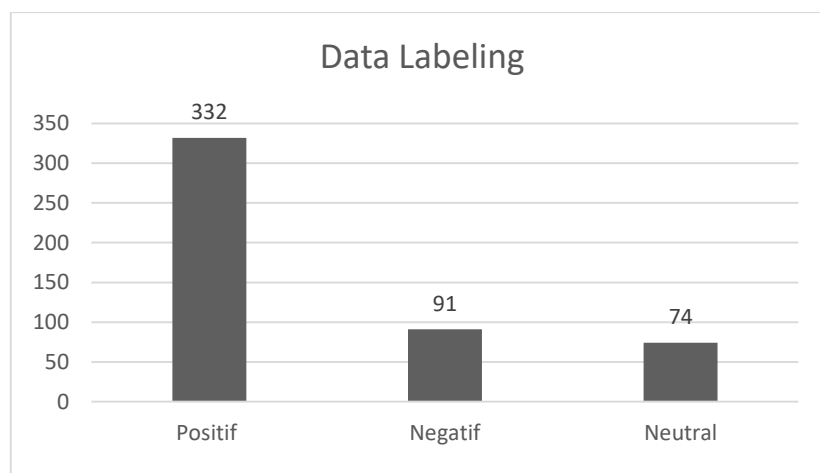
[mltk_data] downloading package stopwords from nltk_data...
[nltk_data] Package stopwords is already up-to-date
Original Text: ['meropokkan', 'palekat', 'harit', 'kayih', 'sayang', 'bagi', 'yang', 'merayabeknya', 'pikuhu', 'darit', 'ang', 'adulah', 'kasih', 'sayang', 'orang', 'yang',
Text after Stopword Removal: ['selamat', 'kashih', 'sayang', 'raya', 'gijak', 'kasih', 'orang', 'orang', 'makin', 'darit', 'maral', 'orang', 'makin', 'kashih', 'sayang']
-----
Original Text: ['nyaua', 'dulu', 'kemana', 'dikabotan', 'ah', 'nyotika']
Text after Stopword Removal: ['nyaua', 'kabotan', 'ah', 'nyotika']
-----
Original Text: ['ada', 'ada', 'kalah', 'tersebut', 'sini', 'tapi', 'bukan', 'kalimatnya', 'kalimatnya', 'Indonesia', 'kemarin', 'bagian', 'apa', 'kempi', 'ke', 'tpe', 'a',
Text after Stopword Removal: ['ada', 'ada', 'ada', 'aduh', 'aduh', 'kalimat', 'Indonesia', 'gila', 'apa', 'kempi', 'tpe', 'aduh', 'kempi', 'kempi']
-----
Original Text: ['petret', 'penggunaan', 'pamili', 'kaskapoles', 'kempi', 'pa-istatistik', 'berasa', 'kashih', 'propen', 'cek', 'peramili', 'penggunaan', 'gasing', 'logistik',
Text after Stopword Removal: ['petret', 'adan', 'sila', 'kaskapoles', 'kempi', 'pa-istatistik', 'kashih', 'propen', 'cek', 'peramili', 'ada', 'gasing', 'logistik', 'pok', 'kem']
-----
Original Text: ['apel', 'gubang', 'tdi', 'polri', 'jaga', 'konduktifitas', 'di', 'kapanasat', 'sakarjo', 'basa', 'tanang', 'pamili', 'adulatur', 'aklamasi']
Text after Stopword Removal: ['apel', 'gasing', 'tdi', 'polri', 'jaga', 'konduktifitas', 'kapanasat', 'sakarjo', 'tanang', 'sila', 'dulu', 'aklamasi']
-----
Original Text: ['alhamdulillah', 'hai', 'jod', 'jam', 'sant', 'rangkap', 'perjalanan', 'menjila', 'tpe', 'harjo', 'dan', 'bagian', 'kemaga', 'sewah', 'berkah']
Text after Stopword Removal: ['alhamdulillah', 'saki', 'rangkap', 'jalan', 'tpe', 'tpe', 'harjo', 'tpe', 'bagian', 'kemaga', 'sewah', 'berkah']
-----
Original Text: ['renes', 'shaved', 'lir', 'tetas', 'buka', 'tanggai', 'bisa', 'di', 'pasa', 'mudi', 'jam', 'bagi', 'lir', 'dan']
Text after Stopword Removal: ['onus', 'tanggap', 'lir', 'buka', 'tanggai', 'pasa', 'jam', 'bagi', 'lir']
-----
Original Text: ['capres', 'bakal', 'yakal', 'hak', 'pilirnya', 'di', 'tpe', 'dumet', 'rumahnya', 'di', 'hambalang', 'boger', 'rubi']
Text after Stopword Removal: ['capres', 'pasa', 'hak', 'pilir', 'tpe', 'rumah', 'hambalang', 'boger', 'rubi']
-----
Original Text: ['sukaika', 'ada', 'dapat', 'serangan', 'fajer', 'saya', 'suka', 'bolang', 'paku', 'ditidink']
Text after Stopword Removal: ['serang', 'fajer', 'bolang', 'paku', 'tidink']
-----
Original Text: ['yuk', 'kubuskan', 'pamili', 'pukun', 'hak', 'pilirnya', 'ya', 'sebat', 'buteys']
Text after Stopword Removal: ['yuk', 'kubus', 'sila', 'hak', 'sila', 'sebat', 'buteys']
-----
Original Text: ['tin', 'paling', 'maya', 'hak', 'cuara', 'dar', 'tiap', 'nyobin', 'di']
Text after Stopword Removal: ['hak', 'cuara', 'nyobin']
-----
Original Text: ['an', 'kemang', 'jantang', 'jangan', 'gelant', 'ya', 'yuk', 'bakikan', 'kintana', 'peda', 'negeri']
Text after Stopword Removal: ['an', 'kemang', 'jantang', 'gelant', 'yuk', 'bakit', 'kintana', 'negeri']
-----
Original Text: ['aru', 'kupas', 'artianlase', 'pukat', 'mengikuti', 'peta', 'demokrat', 'dari', 'asa', 'ke', 'masa']
Text after Stopword Removal: ['aru', 'kupas', 'kompa', 'artianlase', 'pukat', 'demokrat']
-----
Original Text: ['agar', 'kern', 'tatar', 'sebat', 'mengawal', 'pamili']
Text after Stopword Removal: ['kern', 'sebat', 'sila']

```

Setelah mengunggah kamus sentimen InSet, langkah selanjutnya adalah menerapkannya pada data tweet yang telah diproses. Proses ini menghasilkan Polarity Score yang menentukan sentimen tweet, apakah positif, negatif, atau netral. Karena InSet adalah kamus sentimen berbahasa Indonesia, tidak perlu translasi teks, yang menghemat waktu dan menjaga konteks asli tweet. Pendekatan ini efisien, relevan, dan memastikan analisis sentimen yang cepat serta akurat tanpa risiko distorsi makna.

```
Index(['normalized', 'filtering_stopwords'], dtype='object')
      normalized Skor
Sentimen
0  ['mengucapkan', 'selamat', 'hari', 'kasih', 's...
3
1  ['nyusu', 'dulu', 'sebelum', 'dicoblos', 'eh',...
-4
2  ['moa', 'moa', 'moa', 'sudah', 'februari', 'ni...
0
3  ['potret', 'pengamanan', 'pemilu', 'wakapolres...
0
4  ['apel', 'gabungan', 'tni', 'polri', 'jaga', '...
0
..
...
492 ['pemusnahan', 'surat', 'suara', 'lebih', 'pem...
-1
493 ['awas', 'terhadap', 'hoax', 'terkait', 'pemil...
5
494 ['ketua', 'kpu', 'hasyim', 'asyari', 'mengatak...
2
495 ['jadwal', 'rundown', 'kpps', 'hari', 'pemungu...
0
496 ['persoalan', 'logistik', 'di', 'paniai', 'ini...
0
```

Gambar 4. Sampel Hasil Analisis Data



Gambar 5. Grafik Hasil Analisis Data

Hasil analisis sentimen menunjukkan distribusi sebagai berikut: 332 tweet dengan sentimen positif, 91 tweet dengan sentimen negatif, dan 74 tweet dengan sentimen netral. Sentimen positif mendominasi, diikuti oleh sentimen netral dan negatif, memberikan gambaran umum persepsi yang berkembang dalam dataset.

3.3.8 Pembagian Data Latih dan Data Uji

Tahap pembagian dataset menjadi data latih dan data uji dilakukan dengan metode *Hold-Out Validation*, di mana 80% data dialokasikan sebagai data latih dan 20% sisanya digunakan sebagai data uji. Data latih berfungsi untuk membangun dan melatih model, sedangkan data uji digunakan untuk menilai kinerja model pada data yang tidak terlibat dalam proses pelatihan. Metode ini memastikan bahwa evaluasi yang dihasilkan mencerminkan kemampuan model secara lebih objektif. Berdasarkan proporsi tersebut, data latih terdiri atas 265 tweet berlabel positif, 72 tweet berlabel

negatif, dan 59 tweet berlabel netral. Adapun data uji mencakup 67 tweet positif, 19 tweet negatif, dan 15 tweet netral. Pada proses pembagian data, sebanyak 265 data positif untuk data latih dan 67 data positif untuk data uji dipilih secara acak, sehingga diperoleh proporsi data uji terhadap data latih sebesar 25% (67 dari 265 data).

3.3.9 Ekstraksi Fitur

Setelah pemisahan data, langkah selanjutnya adalah ekstraksi fitur menggunakan CountVectorizer untuk mengubah teks dalam data latih menjadi representasi numerik. CountVectorizer menghitung frekuensi kemunculan setiap kata dalam tweet, mengubah teks mentah menjadi matriks fitur. Setiap kolom mewakili kata unik, dan setiap baris menunjukkan frekuensi kemunculannya. Proses ini mengubah data teks yang tidak terstruktur menjadi format terstruktur yang dapat dianalisis oleh algoritma Naïve Bayes Classifier, yang nantinya digunakan untuk mengenali hubungan antara kata-kata dan sentimen dalam data uji.

3.3.10 Model Naïve-Bayes Classifier

Setelah proses penyiapan data selesai, langkah selanjutnya adalah menerapkan algoritma Multinomial Naïve Bayes untuk melakukan analisis sentimen. Algoritma ini bekerja dengan mempelajari distribusi probabilitas kemunculan kata pada masing-masing kategori sentimen, yaitu positif, negatif, dan netral. Tahapan perhitungannya mencakup penentuan nilai prior, likelihood, dan posterior yang digunakan untuk menetapkan sentimen suatu teks. Model dibangun menggunakan data latih dan kemudian diuji menggunakan data uji guna menilai performanya.

Penilaian kinerja model dilakukan melalui beberapa metrik evaluasi, seperti accuracy, precision, recall, dan F1-score. Berdasarkan hasil pengujian, model memperoleh tingkat akurasi sebesar 67%. Selain itu, Confusion Matrix digunakan untuk memberikan gambaran lebih detail mengenai pola kesalahan klasifikasi, sehingga dapat membantu mengidentifikasi aspek yang perlu ditingkatkan

	precision	recall	f1-score	support
positif		0.00	0.00	0.00
18				
negatif		0.00	0.00	0.00
15				
neutral		0.67	1.00	0.80
67				
accuracy				0.67
100				
macro				
avg	0.22	0.33	0.27	100
weighted				
avg	0.45	0.67	0.54	100

Gambar 6. Confussion Matrix

4. KESIMPULAN

Berdasarkan hasil penelitian, analisis sentimen menggunakan pendekatan lexicon-based menunjukkan bahwa sebagian besar tweet tergolong ke dalam sentimen positif, yaitu sebanyak 332 tweet. Sementara itu, 74 tweet teridentifikasi sebagai sentimen netral dan 91 tweet sebagai sentimen negatif. Dominasi sentimen positif tersebut mengindikasikan bahwa respons masyarakat terhadap topik yang dikaji cenderung bernuansa baik, sedangkan keberadaan sentimen netral dan negatif memberikan gambaran tambahan mengenai variasi persepsi dalam dataset.

Selanjutnya, setelah proses pembagian data dilakukan, model Gaussian Naïve Bayes dilatih menggunakan data latih dan dievaluasi menggunakan data uji. Model ini memanfaatkan distribusi Gaussian dalam menghitung probabilitas setiap kelas. Hasil evaluasi menunjukkan bahwa setelah melalui tahap prapemrosesan, model Naïve Bayes Classifier menghasilkan akurasi keseluruhan

sekitar 67%, dengan kinerja paling optimal pada kelas netral dibandingkan dengan kelas positif maupun negatif.

REFERENSI

- [1] Matheos Sarimole, F., & Septian, W. (2024). *Analisis Sentimen Masyarakat Terhadap Isu Penundaan Pemilu 2024 Pada Twitter Dengan Metode Naive Bayes Dan Support Vector Machine*. Jurnal Sains dan Teknologi, 5(3), 880–889. <https://doi.org/10.55338/saintek.v5i1.2789>
- [2] Putra, T. D., Utami, E., & Kurniawan, M. P. (2019). *Analisis Sentimen Pemilu 2024 dengan Naive Bayes Berbasis Particle Swarm Optimization (PSO)*. Jurnal EXPLORE, 13(1), 1–5.
- [3] Manullang, O., & Prianto, C. (2023). *Analisis Sentimen dalam Memprediksi Hasil Pemilu Presiden dan Wakil Presiden: Systematic Literature Review*. ICOM Jurnal Informatika dan Teknologi Komputer, 04(02), 104–113. <https://ejumalunsam.id/index.php/jicom/>
- [4] Herlinawati, N., Yuliani, Y., Faizah, S., Gata, W., & Samudi. (2020). *ANALISIS SENTIMEN ZOOM CLOUD MEETINGS DI PLAY STORE MENGGUNAKAN NAÏVE BAYES DAN SUPPORT VECTOR MACHINE*. CESS (Journal of Computer Engineering System and Science), 5(2), 293–298
- [5] Rozi, I. F., Pramitarini, Y., & Puspitasari, N. (2020). *ANALISIS MENGENAI CALON PRESIDEN INDONESIA 2019 DI TWITTER MENGGUNAKAN METODE BACKPROPAGATION*. JIP (Jurnal Informatika Polinema), 6(2), 27–31.
- [6] Indriani, D. (2019). *Analisis Sentimen pada Tweet dengan Tagar #KPUJANGANCURANG Menggunakan Metode Naive Bayes*. Universitas Islam Riau.
- [7] Wandani, A. (2021). *Sentimen Analisis Pengguna Twitter pada Event Flash Sale Menggunakan Algoritma K-NN, Random Forest, dan Naive Bayes*. Jurnal Sains Komputer & Informatika (J-SAKTI), 5(2), 651–665.
- [8] Budi, S. (2017). *Text Mining Untuk Analisis Sentimen Review Film Menggunakan Algoritma K-Means*. Techno.COM, 16(1), 1–8. <http://www.cs.cornell.edu/People/pabo/movie-review-data/>.