



Deteksi Makna Kebijakan Tunjangan DPR Republik Indonesia Menggunakan Algoritma Naïve Bayes Classifier dan Lexicon

Detecting the Meaning of the Allowance Policy of the House of Representatives of the Republic of Indonesia Using Naïve Bayes Classifier and Lexicon Algorithms

Eka Fauziah^{1, a)} Erna Daniati^{2, b)} Dwi Harini^{3, c)}

¹⁾Jurusan sstem Informas, ²⁾ Jurusan Sistem Informasi, ³⁾ Jurusan Sistem Informasi
Universitas Nusantara PGRI Kediri Jl. Ahmad Dahlan No.76, Mojoroto, Kediri, Indonesia

* Corresponding Author: Erna Daniati, ernadaniati@unpkediri.ac.id

Received: 2 Juli 2026; Accepted: 28 Juli 2026

KEY WORDS

Kata kunci:

Naive Bayes Classifier
Lexicon Algorithm
DPR
Kebijakan
Sentiment Analysis

Keywords:

Naive Bayes Classifier
Lexicon Algorithm
DPR
Allowance
Sentiment Analysis

ABSTRAK

Abstrak. Media sosial berperan sebagai sumber informasi yang dapat dimanfaatkan untuk mengetahui pandangan masyarakat tentang kebijakan yang diambil pemerintah. Salah satu topik yang sering diperbincangkan oleh publik adalah kebijakan yang berkaitan dengan tunjangan anggota DPR RI. Tujuan dari penelitian ini adalah untuk mengkaji bagaimana sentimen masyarakat mengenai kebijakan tersebut dengan menggunakan metode *Naïve Bayes Classifier* dan *Lexicon Sentiment*. Pendekatan penelitian yang diterapkan adalah CRISP-DM (*Cross-Industry Standard Process for Data Mining*) yang meliputi tahap pemahaman bisnis, pemahaman data, persiapan data, pemodelan, evaluasi, serta penerapan. Data dikumpulkan dari platform media sosial X (Twitter) melalui *scraping* dan diproses dengan langkah-langkah *preprocessing* serta pemberian bobot TF-IDF. Temuan dari penelitian ini mengindikasikan bahwa metode *Naïve Bayes Classifier* mencapai akurasi sebesar 74%, sementara metode *Lexicon Sentiment* membantu dalam memahami nuansa emosional yang ada dalam teks. Kombinasi dari kedua metode ini menghasilkan analisis sentimen yang lebih komprehensif dan relevan dibandingkan dengan hanya menggunakan salah satu metode saja. Penelitian ini membuktikan bahwa gabungan pendekatan statistik dan berbasis kamus sangat berguna dalam analisis sentimen terhadap opini masyarakat berbahasa Indonesia.

Abstract. Social media serves as a source of information that can be used to gauge public opinion regarding government policies one topic frequently discussed by the public is the policy regarding allowances for members of the Indonesian House of Representatives. The objective of this study is to examine public sentiment regarding these policies using the *Naïve Bayes Classifier* and *Lexicon Sentiment* methods. The research approach applied is CRISP-DM (*Cross Industry Standard Process for Data Mining*), which encompasses the stages of business understanding, data understanding, data preparation, modeling, evaluation, and implementation. Data was collected from the social media platform X (Twitter) via *scraping* and processed through *preprocessing* steps and TF-IDF weighting. The findings of this study indicate that the *Naïve Bayes Classifier* method achieved an accuracy of 74%, while the *Lexicon Sentiment* method helped in understanding the emotional nuances present in the text. The combination of these two methods produces a more comprehensive and relevant sentiment analysis compared to using only one method alone. This study demonstrates that the combination of statistical and lexicon-based approaches is highly useful in analyzing sentiment regarding the opinions of the Indonesian-speaking public.

1. Pendahuluan

Perkembangan teknologi informasi di era modern telah mengubah lanskap interaksi, cara menyampaikan pendapat, serta pola komunikasi masyarakat di ruang publik. Platform media sosial seperti Twitter, Facebook, Instagram, dan berbagai situs web lainnya kini menjadi sarana utama bagi masyarakat Indonesia dalam mengutarakan pandangan mereka terkait berbagai dinamika sosial, ekonomi, dan politik. Media sosial berfungsi sebagai wadah komunikasi publik yang memfasilitasi masyarakat untuk saling berbagi perspektif dan merespons berbagai fenomena yang sedang berkembang[1]. Aktivitas digital ini menghasilkan data opini publik yang melimpah dan bersifat terbuka, sehingga dapat dimanfaatkan sebagai sumber informasi untuk mengukur reaksi masyarakat terhadap regulasi pemerintah. Ulasan serta pandangan yang beredar di internet dapat menjadi indikator awal untuk mengetahui tingkat kepuasan atau cara pandang masyarakat terhadap suatu isu dan kebijakan tertentu [2]. Salah satu regulasi pemerintah yang memicu perhatian luas adalah kebijakan mengenai pemberian tunjangan bagi anggota Dewan Perwakilan Rakyat Republik Indonesia. Kebijakan ini mendatangkan berbagai reaksi dari masyarakat, mulai dari dukungan karena adanya harapan terhadap peningkatan performa dalam penyusunan undang-undang, hingga kritik akibat dinilai menciptakan kesenjangan kesejahteraan antara pejabat publik dan masyarakat luas. Respons tersebut masif dilontarkan lewat media sosial dalam bentuk komentar, unggahan, serta diskusi daring yang merefleksikan sentimen dan sikap masyarakat terhadap kebijakan itu. Kumpulan tweet tersebut memuat pandangan dari berbagai pengguna media sosial, yang kemudian dapat dikelompokkan ke dalam beberapa kategori, yaitu opini positif, negatif, dan netral [3]. Aspirasi masyarakat di media sosial tidak sekadar menunjukkan bentuk persetujuan atau penolakan terhadap suatu kebijakan, melainkan juga mencakup spektrum emosi yang luas seperti amarah, kekecewaan, sindiran jenaka, hingga dukungan normatif. Oleh sebab itu, diperlukan sebuah pendekatan analisis yang sistematis dan berbasis data agar dapat memahami pandangan masyarakat secara objektif dan menyeluruh. Analisis sentimen kerap digunakan untuk meninjau bagaimana opini publik berkembang berdasarkan teks yang diunggah di media sosial.

Metode evaluasi yang bersandar pada pengalaman pengguna serta sudut pandang masyarakat umum juga sering diimplementasikan dalam riset sistem digital untuk mengukur tingkat penerimaan publik terhadap suatu layanan atau aplikasi [4]. Analisis sentimen sendiri merupakan bagian dari teknologi pengolahan bahasa alami (Natural Language Processing/NLP) yang berfungsi untuk mengidentifikasi sekaligus mengelompokkan polaritas emosi dalam teks ke dalam kategori seperti positif, negatif, atau netral. Menurut Saron Tandiapa dan Caren Rorimpandey, masifnya perkembangan media online membuka peluang besar untuk melakukan analisis sentimen karena opini publik kini dapat dipantau secara mudah lewat platform digital [5]. Dalam studi analisis sentimen, algoritma Naïve Bayes Classifier sering digunakan karena keunggulannya dalam memproses data berskala besar secara cepat dan efektif. Algoritma Naïve Bayes dianggap cukup efektif dalam mengklasifikasikan teks karena memiliki efisiensi komputasi yang tinggi serta mampu bekerja dengan volume data yang masif [2]. Dimas Sanjaya dan timnya memaparkan bahwa metode ini bekerja berdasarkan teori probabilitas dan efisien dalam memproses teks dengan mengasumsikan bahwa setiap fitur dalam teks bersifat independen atau saling bebas satu sama lain[6]. Di samping itu, pendekatan berbasis kamus atau Lexicon-Based juga sering diterapkan dalam analisis sentimen karena mampu mendeteksi muatan emosional teks melalui daftar kata yang telah memiliki nilai polaritas. Menurut Saron Tandiapa dan Caren Rorimpandey, metode berbasis kamus ini menunjukkan performa yang baik saat digunakan untuk menganalisis opini di media sosial[5].

Sejumlah riset terdahulu mengindikasikan bahwa mengombinasikan metode berbasis pembelajaran mesin dengan metode berbasis kamus dapat mengoptimalkan tingkat akurasi klasifikasi sentimen. Studi yang dilakukan oleh Amrullah dan Solichin pada tahun 2024 menunjukkan bahwa model Multinomial Naïve Bayes mampu memperoleh tingkat akurasi hingga 90,67% dalam mengelompokkan sentimen[7]. Penelitian lain oleh Fitriono dan tim pada tahun 2025 mengungkapkan bahwa pendekatan kombinasi ini mampu menghasilkan akurasi tertinggi mencapai 96,49%. Kendati demikian, Fitriono dan timnya pada tahun 2025 juga menggarisbawahi bahwa pemrosesan sentimen berbasis bahasa Indonesia masih menghadapi banyak tantangan, seperti penggunaan kosakata tidak baku, keberagaman dialek daerah, serta tingginya unsur sarkasme pada berbagai konten media sosial [8]. Selain itu, ketidakseimbangan dalam distribusi data dapat memengaruhi performa model klasifikasi karena kecenderungan model untuk lebih memprioritaskan kelas yang jumlah sampelnya mayoritas[9]. Meskipun sudah banyak penelitian yang dilakukan, mayoritas riset tersebut masih terbatas pada pengelompokan sentimen standar ke dalam tiga kategori, yaitu positif, negatif, dan netral. Studi yang mengintegrasikan klasifikasi perasaan dengan analisis makna emosional dalam konteks opini publik, khususnya mengenai isu kebijakan politik, masih

sangat jarang dijumpai. Keterbatasan pada penelitian-penelitian sebelumnya inilah yang menjadi landasan utama dilakukannya studi ini.

Berangkat dari permasalahan tersebut, tujuan penelitian ini adalah untuk menganalisis sikap masyarakat terhadap kebijakan tunjangan anggota DPR RI menggunakan metode Naïve Bayes Classifier untuk mengklasifikasikan sentimen dan metode Lexicon Sentiment untuk membedah makna emosionalnya. Batasan masalah dalam riset ini difokuskan pada data yang diekstraksi dari media sosial X (Twitter), dengan spesifikasi pada teks yang ditulis menggunakan bahasa Indonesia. Sentimen dalam teks tersebut dipisahkan ke dalam tiga kategori, yaitu positif, negatif, dan netral. Pengujian performa model dilakukan dengan mengukur beberapa metrik seperti akurasi, presisi, recall, dan skor F1. Pendekatan yang digunakan dalam penelitian ini mengadopsi kerangka kerja CRISP-DM, yaitu standar proses lintas industri untuk penambangan data, yang meliputi tahap memahami permasalahan, memahami data, mempersiapkan data, membangun model, mengevaluasi hasil, hingga menerapkan hasilnya. Tahapan ini diterapkan agar proses analisis data dapat berjalan secara teratur dan terorganisir. Melalui pendekatan ini, penelitian diharapkan dapat menyajikan gambaran yang lebih mendalam mengenai persepsi masyarakat terhadap kebijakan pemerintah, khususnya terkait tunjangan DPR RI. Selain itu, riset ini juga diproyeksikan mampu menjadi referensi dalam pengembangan metode analisis sentimen berbahasa Indonesia sekaligus menjadi fondasi bagi riset selanjutnya dalam merancang model analisis opini publik yang lebih komprehensif.

2. Tinjauan Pustaka

Studi terdahulu memegang peranan yang sangat krusial dalam perkembangan riset ini, lantaran membantu memberikan pemahaman mengenai teknik yang telah diterapkan serta memetakan ranah yang masih berpotensi untuk diteliti lebih mendalam. Abdillah dan Hasan (2023) menerapkan analisis sentimen terhadap kandidat presiden dengan memanfaatkan algoritma Naïve Bayes. Temuan mereka mengindikasikan bahwa metode tersebut efektif dalam memilah opini publik di media sosial, sehingga dinilai representatif untuk diaplikasikan pada studi yang berkaitan dengan isu-isu politik [10]. Riset lainnya yang dikerjakan oleh Cinta Azzaria dan timnya pada tahun 2023 mengimplementasikan metode Naïve Bayes Classifier untuk meninjau data penjualan produk, di mana hasil pengujiannya mampu menyentuh tingkat akurasi yang sangat optimal. Hal ini membuktikan bahwa metode Naïve Bayes mempunyai fleksibilitas yang tinggi dan dapat diimplementasikan pada beragam karakteristik data [11]. Alviyanti dan tim peneliti (2024) juga melangsungkan analisis terhadap ulasan pengguna aplikasi menggunakan algoritma Naïve Bayes Classifier.

Studi tersebut membuktikan bahwa metode ini mampu memberikan performa yang optimal dalam mengelompokkan data berbasis teks [12]. Di samping itu, Harwenda dan tim peneliti (2025) menggarisbawahi bahwa metode Lexicon Sentiment sangat membantu dalam membedah sudut pandang masyarakat terkait regulasi pemerintah melalui pemanfaatan pendekatan yang bersandar pada kamus kata [13]. Pada penelitian yang berbeda, Asri dan timnya pada tahun 2022 menggabungkan metode VADER Lexicon dengan Naïve Bayes untuk mengkaji ulasan pengguna pada aplikasi PLN Mobile. Output dari riset tersebut mengindikasikan bahwa integrasi antara pendekatan berbasis kamus dan teknik pembelajaran mesin mampu memberikan hasil yang lebih efektif dalam pemrosesan analisis sentimen. Temuan ini memperkuat argumen bahwa penerapan metode campuran berpotensi melahirkan performa yang jauh lebih baik [14]. Mengacu pada peta jalan berbagai literatur yang telah dipaparkan, dapat ditarik kesimpulan bahwa mayoritas studi terdahulu masih berfokus pada kategorisasi sentimen standar ke dalam kelompok positif, negatif, dan netral saja. Oleh sebab itu, penelitian ini berupaya menyajikan analisis yang lebih komprehensif dengan mengintegrasikan metode Naïve Bayes Classifier dan Lexicon Sentiment guna membedah opini masyarakat terkait kebijakan alokasi tunjangan anggota DPR RI.

2.1 Analisis Sentimen

Analisis sentimen adalah bagian dari *Natural Language Processing* (NLP) yang berfungsi untuk mengidentifikasi, mengelompokkan, dan menganalisis opini atau emosi yang terdapat dalam teks. Menurut Abdillah dan Hasan (2023), analisis sentimen dapat menawarkan gambaran tentang kecenderungan opini publik berdasarkan data teks yang bersumber dari media sosial. Pendekatan ini sering digunakan untuk memahami respons masyarakat terhadap isu sosial, ekonomi, dan politik [10]. Dalam lingkup kebijakan publik, analisis sentimen sangat penting karena mampu memberikan informasi berbasis data tentang pandangan masyarakat dengan cepat dan objektif.

2.2 Naïve Bayes Classifier

Naïve Bayes Classifier adalah algoritma klasifikasi berbasis pada probabilitas yang menggunakan Teorema Bayes. Metode ini berfungsi dengan menghitung kemungkinan bahwa suatu dokumen termasuk dalam kategori

tertentu berdasarkan fitur kata yang ada di dalam teks. Metode ini sering digunakan dalam analisis sentimen karena kesederhanaannya, kecepatannya, dan efektivitasnya dalam menangani data teks. Alviyanti et al. (2024) menunjukkan bahwa *Naïve Bayes* dapat memberikan akurasi yang baik dalam mengklasifikasikan ulasan pengguna aplikasi [12].

2.3. *Lexicon Sentiment*

Metode *Lexicon Sentiment* adalah strategi yang menggunakan kamus kata untuk mengidentifikasi polaritas sentimen berdasarkan nilai positif dan negatif dari kata-kata. Setiap kata dalam teks akan dicocokkan dengan daftar sentimen, lalu skor yang diperoleh akan dijumlah untuk mendapatkan total nilai sentimen. Menurut Harwenda et al. (2025), teknik ini cukup efektif dalam menangkap nuansa emosional dalam teks, meskipun memiliki kelemahan dalam memahami konteks kalimat yang lebih rumit [13].

2.4. *Integrasi Naïve Bayes dan Lexicon Sentiment*

Penggabungan teknik *Naïve Bayes Classifier* dengan *Lexicon Sentiment* memiliki tujuan untuk meningkatkan akurasi dalam klasifikasi sentimen. *Naïve Bayes* berfokus pada pengenalan pola statistik data, sementara *Lexicon Sentiment* berperan dalam memahami polaritas emosional dari kata-kata. Metode ini sangat berguna untuk menganalisis pandangan masyarakat tentang kebijakan pemerintah, karena tanggapan publik sering kali berisi kritik, sindiran, dan dukungan emosional yang mungkin tidak terdeteksi oleh satu metode saja [13].

3. Metode Penelitian

Penelitian ini mengadopsi kerangka kerja CRISP-DM (*Cross-Industry Standard Process for Data Mining*), yang menjadi standar acuan utama dalam aktivitas penambangan data. Alur kerja ini terbagi menjadi enam fase krusial, meliputi pemahaman bisnis, pemahaman data, persiapan data, pemodelan, evaluasi performa, serta penyebaran model. Dalam studi ini, CRISP-DM diimplementasikan sebagai panduan menyeluruh, mulai dari mengidentifikasi rumusan masalah, mengoleksi data dari platform media sosial X (Twitter), melakukan pembersihan data, membangun model klasifikasi *Naïve Bayes* dan *Lexicon Sentiment*, menguji performa model, hingga menyajikan hasil analisis akhir. Penerapan metodologi ini berfungsi untuk memastikan bahwa seluruh rangkaian riset dapat terlaksana secara sistematis, terstruktur, dan terencana dengan matang. Pendekatan kuantitatif deskriptif kerap diaplikasikan dalam kajian sistem informasi guna mengukur aspek pengalaman pengguna serta respons mereka terhadap suatu sistem digital [4]. Jenis data yang dimanfaatkan dalam penelitian ini adalah data sekunder berupa opini publik yang diperoleh dari platform X (Twitter) melalui teknik ekstraksi data atau scraping. Proses pengumpulan data tersebut dijalankan dengan mengandalkan kata kunci spesifik yang merepresentasikan isu kebijakan tunjangan anggota Dewan Perwakilan Rakyat Republik Indonesia. Atribut data yang ditarik meliputi muatan teks dari tweet, lini masa atau tanggal pengunggahan tweet, beserta elemen informasi relevan lainnya yang berpotensi memperkuat analisis riset. Begitu tahapan pengumpulan data rampung, prosedur selanjutnya adalah preprocessing data, yaitu serangkaian proses pembersihan dan penataan data agar siap diintegrasikan ke dalam tahap pemodelan. Adapun tahapan pembersihan tersebut mencakup penghapusan URL, mention, hashtag, emoji, serta karakter simbol yang tidak diperlukan.

1. Menghapus URL, mention, hashtag, emoji, dan simbol.
2. Menghapus data duplikat dan data kosong.
3. *Case folding*, yaitu mengubah seluruh huruf menjadi huruf kecil.
4. *Tokenization*, yaitu mengubah kalimat menjadi kata.
5. *Stopword removal*, yaitu menghapus kata umum yang tidak memiliki makna penting.
6. *Stemming*, yaitu mengubah kata menjadi bentuk dasar.

Data teks yang telah diproses kemudian diubah ke bentuk numerik menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*). Metode ini digunakan untuk memberikan bobot pada kata berdasarkan tingkat kepentingannya dalam dokumen. Tahap preprocessing dan transformasi data menjadi langkah penting sebelum proses klasifikasi dilakukan agar model mampu mengenali pola data dengan lebih baik[9].

$$TFIDF(v\ d) = \frac{TF(w,d)}{DF(w)} \quad (1)$$

Metode *Naïve Bayes Classifier* digunakan untuk mengklasifikasikan sentimen ke dalam kategori positif, negatif, dan netral. Metode ini bekerja berdasarkan probabilitas kemunculan kata dalam setiap kelas sentimen.

$$P(C|X) \propto PC) \prod_{i=1}^n P(x_i | C) \quad (2)$$

Selain menggunakan metode *Naïve Bayes Classifier*, penelitian ini juga menerapkan metode *Lexicon Sentiment* untuk mengidentifikasi makna emosional pada teks, dengan cara mencocokkan kata-kata pada tweet dengan kamus sentimen positif dan negatif, kemudian menghitung skor polaritas untuk menentukan kecenderungan sentimen dari setiap tweet. Evaluasi model dilakukan menggunakan *confusion matrix* untuk mengukur performa klasifikasi sentimen yang dihasilkan oleh metode *Naïve Bayes Classifier*. Metrik evaluasi yang digunakan meliputi *accuracy*, *precision*, *recall*, dan *F1-score*. Adapun formulasi metrik evaluasi disajikan pada formula 3-6.

$$accuracy = \frac{TP+TN}{TP+FP+TN+FN} \quad (3)$$

$$precision = \frac{TP}{TP+FP} \quad (4)$$

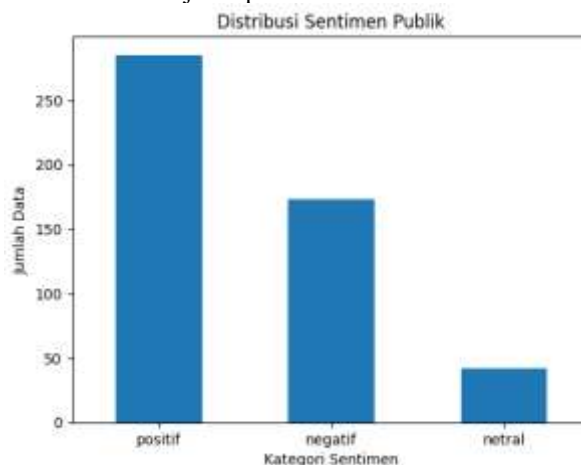
$$Recall = \frac{TP}{TP+FN} \quad (5)$$

$$f1 - score = \frac{2 \times precision \times recall}{precision + recall} \quad (6)$$

Dengan keterangan sebagai berikut TP (*True Positive*) merupakan data positif yang berhasil diprediksi benar, TN (*True Negative*) merupakan data negatif yang berhasil diprediksi benar, FP (*False Positive*) merupakan data negatif yang diprediksi positif, dan FN (*False Negative*) merupakan data positif yang diprediksi negatif. Hasil analisis kemudian disajikan dalam bentuk tabel dan visualisasi distribusi sentimen untuk menunjukkan kecenderungan opini publik terhadap kebijakan tunjangan DPR RI. Selain itu, visualisasi kata dominan dan distribusi emosi digunakan untuk membantu interpretasi hasil analisis secara lebih mendalam.

4. Hasil dan Pembahasan

Studi ini mengkaji karakteristik respons masyarakat terhadap regulasi tunjangan anggota DPR RI lewat pemanfaatan tiga skema pendekatan, yakni algoritma klasifikasi *Naïve Bayes*, Analisis Sentimen Berbasis Kamus (*Lexicon*), serta integrasi dari kedua metodologi tersebut. Data yang bersumber dari platform media sosial X (Twitter) ini selanjutnya diproses melalui serangkaian tahapan terstruktur, mulai dari penyusunan korpus data awal, pembobotan fitur menggunakan TF-IDF, proses klasifikasi sentimen, hingga pengujian performa terhadap model yang diterapkan. Sebelum memasuki tahapan klasifikasi, data opini publik wajib melalui fase pra-pemrosesan teks terlebih dahulu yang mencakup pembersihan data, pemotongan kata (*tokenisasi*), serta normalisasi teks. Prosedur pembersihan dikerjakan dengan mengeliminasi URL, simbol, tanda baca, angka, dan karakter khusus lainnya guna menghasilkan struktur teks yang lebih rapi dan seragam. Selanjutnya, tahapan pemotongan kata diterapkan untuk mengurai runtutan kalimat menjadi satuan kata tunggal agar mempermudah proses identifikasi sentimen pada tahap berikutnya. Sementara itu, proses normalisasi diarahkan untuk menyelaraskan kosakata tidak baku menjadi bentuk kata yang standar sehingga konsistensi makna di dalam teks tetap terjaga. Adapun distribusi sentimen disajikan pada Gambar 1.



GAMBAR 1. Distribusi sentimen

Setelah fase pra-pemrosesan rampung, data teks tersebut akan ditransformasikan menggunakan metode TF-IDF guna mengukur frekuensi kemunculan kata sekaligus tingkat kepentingannya di dalam dokumen. Teknik TF-IDF diandalkan untuk menetapkan bobot yang proporsional pada masing-masing kata berdasarkan pada tingkat urgensinya di dalam korpus data. Apabila sebuah kata memiliki frekuensi kemunculan yang tinggi pada suatu

dokumen tertentu namun sangat jarang ditemukan pada dokumen lain, nilai bobot TF-IDF dari kata tersebut akan menjadi semakin besar. Penerapan langkah ini bertujuan agar model klasifikasi dapat lebih sensitif dalam mengenali kata-kata kunci yang krusial untuk menentukan orientasi polaritas atau sentimen. Rumus matematis untuk perhitungan TF-IDF mengacu pada persamaan nomor 1. Output nilai dari proses TF-IDF ini nantinya diintegrasikan sebagai fitur masukan dalam algoritma klasifikasi Naïve Bayes. Prosedur kategorisasi selanjutnya dijalankan dengan algoritma klasifikasi Naïve Bayes, yang prinsip kerjanya bersandar pada perhitungan peluang kemunculan suatu kata di setiap kelas sentimen. Metodologi ini mengimplementasikan Teorema Bayes untuk mengalkulasi nilai probabilitas akhir, sehingga dapat ditentukan apakah suatu teks cenderung masuk ke dalam kelompok sentimen positif, negatif, atau netral. Persamaan Naïve Bayes yang diaplikasikan dalam pengujian ini merujuk pada formula (2). Visualisasi distribusi sentimen pada Gambar 1 memaparkan bahwa kategori positif memiliki jumlah data paling dominan dibandingkan kategori negatif dan netral. Hal ini menegaskan bahwa sebagian opini publik mengandung kata-kata yang bersifat apresiatif terhadap kebijakan yang dibahas. Namun demikian, sentimen negatif juga muncul dalam jumlah cukup besar sehingga menunjukkan masih adanya kritik masyarakat terhadap kebijakan tunjangan DPR RI. Selain distribusi sentimen, digunakan visualisasi word cloud untuk melihat kata-kata yang paling sering muncul dalam data opini publik berdasarkan masing-masing kategori sentimen.



GAMBAR 2. Kategori sentimen

Berdasarkan Gambar 2, setiap kategori sentimen memiliki karakteristik kata yang berbeda. Pada sentimen positif, kata-kata yang dominan mengutarakan adanya dukungan atau apresiasi terhadap kebijakan yang dibahas. Pada sentimen negatif, kata seperti “rakyat”, “pajak”, “tunjangan”, dan “gaji” muncul lebih dominan sehingga menandakan adanya kritik masyarakat terhadap penggunaan anggaran negara dan kesejahteraan pejabat publik. Sementara itu, sentimen netral didominasi oleh kata-kata yang bersifat informatif dan tidak menunjukkan kecenderungan opini tertentu. Visualisasi word cloud juga memperlihatkan bahwa isu kesejahteraan masyarakat dan kebijakan pemerintah menjadi topik utama dalam diskusi publik di media sosial. Visualisasi data seperti word cloud dan grafik digunakan untuk membantu memahami distribusi data dan pola sentimen secara lebih mendalam [1]. Setelah proses klasifikasi dilakukan, model dievaluasi menggunakan confusion matrix dan classification report untuk mengetahui performa model. Evaluasi dilakukan menggunakan metrik accuracy, precision, recall, dan F1-score. Pengujian validitas dan reliabilitas merupakan tahapan penting untuk memastikan hasil evaluasi sistem dapat dipercaya dan memiliki konsistensi yang baik [4]. Accuracy digunakan untuk mengukur tingkat ketepatan prediksi model secara keseluruhan. Precision menunjukkan ketepatan model dalam memprediksi suatu kelas, recall mengukur kemampuan model dalam menemukan data pada kelas tertentu, sedangkan F1-score digunakan untuk melihat keseimbangan antara precision dan recall sesuai dengan formula (3)-(5) dengan hasil evaluasi seperti tabel di bawah ini.

Tabel 1.
Hasil pengujian

	Precision	Recall	F1-score	Support
Negatif	0.00	0.00	0.00	2
Negatif	0.72	1.00	0.84	67
Netral	1.00	0.25	0.40	28

Positif	0.00	0.00	0.00	3
accuracy			0.74	100
macro avg	0.43	0.31	0.31	100
weighted avg	0.76	0.74	0.67	100

Hasil analisis menunjukkan bahwa model mencapai tingkat akurasi sebesar 74%. Kategori negatif menunjukkan kinerja terbaik karena model mampu mendeteksi sebagian besar data negatif dengan baik. Di sisi lain, kinerja kategori positif masih terbilang rendah disebabkan oleh ketidakseimbangan jumlah data.

Penelitian ini tidak hanya memanfaatkan *Naïve Bayes Classifier*, tetapi juga menerapkan metode berbasis leksikon untuk mengenali sentimen berdasarkan nilai kata yang tercantum dalam kamus leksikon bahasa Indonesia. Hasil klasifikasi menunjukkan bahwa sentimen positif lebih mendominasi dengan jumlah 285 data, diikuti oleh sentimen negatif sebanyak 173 data, dan sentimen netral sebanyak 42 data. Dominasi sentimen positif menunjukkan bahwa banyak opini publik mengandung kata-kata dengan nilai positif secara leksikal. Namun, metode *lexicon* memiliki batasan dalam memahami konteks kalimat, ironi, dan sarkasme yang sering muncul di media sosial. Oleh karena itu, hasil dari *lexicon* digabungkan dengan metode *Naïve Bayes Classifier* untuk memberikan klasifikasi sentimen yang lebih sesuai konteks dan akurat. Metode *lexicon* berfungsi dengan memberikan nilai sentimen pada tiap kata berdasarkan kamus *lexicon* Bahasa Indonesia. Penggabungan metode ini dilakukan dengan mencampurkan hasil klasifikasi dari *Naïve Bayes* dan *Lexicon Sentiment*. Pendekatan ini bertujuan untuk memadukan kemampuan probabilistik *Naïve Bayes* dalam mengenali pola data dengan kemampuan *lexicon* dalam memahami arti emosional kata. Hasil penggabungan menunjukkan distribusi klasifikasi yang lebih seimbang dibandingkan dengan menggunakan satu metode, karena dapat memadukan pendekatan probabilistik dan makna emosional dari teks. Penelitian ini menyoroti bahwa penerapan metode kombinasi antara *Naïve Bayes Classifier* dan *Lexicon Sentiment* dapat menghasilkan analisis sentimen yang lebih sesuai dengan konteks terhadap opini publik terkait kebijakan tunjangan DPR RI. Penggunaan TF-IDF membantu model dalam mengenali kata-kata kunci di dalam teks, sementara pendekatan *lexicon* berperan dalam memahami makna emosional dari pendapat masyarakat. Namun, penelitian ini masih mengalami keterbatasan dalam memahami bahasa informal, singkatan, dan sarkasme yang sering dijumpai di media sosial. Selain itu, ketidakseimbangan distribusi data juga berdampak pada kinerja model di beberapa kategori sentimen.

5. Kesimpulan

Studi ini menerapkan kerangka kerja CRISP-DM, yang menjadi standar acuan lintas industri dalam penambangan data, guna menguji karakteristik respons masyarakat terhadap regulasi alokasi tunjangan anggota DPR RI pada platform media sosial X (Twitter). Pada fase pemahaman bisnis, fokus utama penelitian ini adalah mengidentifikasi sudut pandang masyarakat mengenai kebijakan tersebut lewat analisis orientasi polaritas serta sikap publik. Pada fase pemahaman data, sekumpulan data tweet diekstraksi secara otomatis untuk kemudian dibedah karakteristiknya guna memperoleh gambaran mengenai sebaran serta kualitas informasi teks. Pada fase persiapan data, serangkaian tahapan dilakukan, meliputi pembersihan data, penyeragaman struktur huruf, pemotongan kata, eliminasi kata umum (stopword removal), pengembalian kata ke bentuk dasar (stemming), hingga perhitungan bobot TF-IDF agar korpus data siap diintegrasikan ke dalam fase pemodelan. Pada fase pemodelan, penelitian ini mengimplementasikan algoritma klasifikasi *Naïve Bayes*, pendekatan sentimen berbasis kamus (*lexicon*), serta kombinasi terintegrasi dari kedua metodologi tersebut. Berdasarkan hasil pengujian, algoritma klasifikasi *Naïve Bayes* mampu memperoleh tingkat akurasi sebesar 0,74 dengan perolehan nilai weighted average F1-score mencapai 0,67. Di sisi lain, pendekatan sentimen berbasis kamus mampu menyajikan sebaran klasifikasi yang lebih proporsional serta lebih optimal dalam menangkap muatan makna emosional teks berdasarkan pada konteks kalimatnya. Integrasi antara kedua metode ini mampu melahirkan performa klasifikasi sentimen yang lebih presisi karena mengombinasikan keunggulan pendekatan statistik dengan pendekatan berbasis kamus kata. Fase evaluasi dilaksanakan dengan mengandalkan instrumen matriks kebingungan (*confusion matrix*) serta beragam parameter performa seperti akurasi, presisi, recall, dan F1-score. Output dari pengujian ini mengonfirmasi bahwa model yang dikembangkan memiliki performa yang baik dalam mengidentifikasi tren sentimen dominan pada opini publik. Selain itu, temuan dari analisis emosi mengindikasikan bahwa mayoritas respons masyarakat terkonsentrasi pada kategori emosi netral dan marah, yang merepresentasikan adanya gelombang kritik serta ketidakpuasan emosional publik terhadap regulasi tunjangan yang ditetapkan oleh DPR RI. Pada fase penyebaran (*deployment*), hasil analisis disajikan dalam bentuk visualisasi sebaran polaritas sentimen, word cloud, serta grafik klasifikasi emosi guna mempermudah interpretasi data. Menilik keseluruhan output dari riset ini, model yang dirancang telah berhasil memenuhi target penelitian, yakni mampu mengategorikan sentimen masyarakat terkait kebijakan tunjangan DPR RI sekaligus memberikan paparan komprehensif mengenai muatan emosional di dalam aspirasi publik. Di samping itu, implementasi gabungan antara algoritma Klasifikasi

Naïve Bayes dan Sentimen Berbasis Kamus terbukti mampu menyajikan kedalaman analisis yang jauh lebih optimal ketimbang hanya bersandar pada penggunaan satu metode tunggal saja. Kesimpulan merupakan pernyataan singkat, jelas, dan tepat tentang apa yang diperoleh, memuat keunggulan dan kelemahan, dapat dibuktikan, serta terkait langsung dengan tujuan penelitian. Saran memuat berbagai usulan atau pendapat yang sebaiknya dikaitkan dengan penelitian sejenis. Saran dibuat berdasarkan kelemahan, pengalaman, kesulitan, kesalahan, temuan baru yang belum diteliti dan berbagai kemungkinan arah penelitian selanjutnya.

Daftar Pustaka

- [1] T. C. Adisti, E. Daniati, and A. Ristiyawan, "ANALISIS SENTIMEN UJARAN KEBENCIAN PADA TWEET DI TWITTER," Kediri, Apr. 2025.
- [2] A. B. Fachturroziq, J. V. Wicaksono, and E. Daniati, "Analisis Sentimen Ulasan Produk Amazon Menggunakan Algoritma Naive Bayes Untuk Prediksi Rating Produk," 2025. [Online]. Available: <https://proceeding.unpkediri.ac.id/index.php/inotek/>
- [3] E. Daniati and H. Utama, "ANALISIS SENTIMEN DENGAN PENDEKATAN ENSEMBLE LEARNING DAN WORD EMBEDDING PADA TWITTER," 2023.
- [4] M. Alfianto Alamsyah, Indriati Rini, and Harini Dwi, "Analisis Pengalaman Pengguna Terhadap Aplikasi Panjalu e-Library di Perpustakaan Mastrip Pare," Prosiding SEMNAS INOTEK (Seminar Nasional Inovasi Teknologi), 2025.
- [5] S. Saron Tandiapa and G. Caren Rorimpandey, "ANALISIS SENTIMEN ULASAN PENGGUNA PADA APLIKASI THREADS DENGAN METODE LEXICON BASED DAN NAIVE BAYES CLASSIFIER," 2024.
- [6] A. Dimas Sanjaya, T. Hendro Pudjiantoro, A. Kania Ningsih, and F. Renaldi, "Sentiment Analysis Of E-Wallets on Twitter Social Media with Naïve Bayes and Lexicon-Based Methods," Sep. 2022.
- [7] F. Amrullah and A. Solichin, "Analisis Emosi Pada Live Chat Youtube 'Mata Najwa: 3 Bacapres Bicara Gagasan' Menggunakan Pendekatan Lexicon dan Algoritma Naive Bayes," Jurnal TICOM: Technology of Information and Communication, vol. 12, no. 3, pp. 121–128, May 2024, [Online]. Available: <https://saifmohammad.com/WebPages/NRC-Emotion->
- [8] D. Fitrono, R. Indriati, and A. Ristiyawan, "Analisis Sentimen Ulasan Aplikasi Chatgpt Menggunakan Algoritma Support Vector Machine dan Lexicon-Based," JSITIK: Jurnal Sistem Informasi dan Teknologi Informasi Komputer, vol. 3, no. 2, pp. 101–111, Jun. 2025, doi: 10.53624/jsitik.v3i2.719.
- [9] A. Nugroho and D. Harini, "Teknik Random Forest untuk Meningkatkan Akurasi Data Tidak Seimbang," JSITIK, vol. 2, no. 2, 2024, doi: 10.53624/jsitik.v2i2.XX.
- [10] A. R. Abdillah and F. N. Hasan, "Analisis Sentimen Terhadap Kandidat Calon Presiden Berdasarkan Tweets Di Sosial Media Menggunakan Naive Bayes Classifier," SMATIKA JURNAL, vol. 13, no. 01, pp. 117–130, Jul. 2023, doi: 10.32664/smatika.v13i01.750.
- [11] Cinta Azzaria, Mieta Silvia Aviva, Ela Esti Susanti, Erna Daniati, and Aidina Ristiyawan, "Analisis Data Mining dengan Naive Bayes Classifier untuk Peningkatan Penjualan Produk," Aug. 2023.
- [12] S. H. Alviyanti, A. Purwandira, I. Febiyanti, E. Daniati, and A. Ristiyawan, "Klasifikasi Sentimen Pengguna Aplikasi Livin By Mandiri Pada Play Store Menggunakan Algoritma Naive Bayes," Online, Aug. 2024.
- [13] R. W. Harwenda, M. David Angelo, I. Budi, A. B. Santoso, and K. Putra, "Sentiment Analysis on Government Public Policies: A Systematic Literature Review," DIJEMSS, vol. 6, no. 5, 2025, doi: 10.38035/dijemss.v6i5.
- [14] Y. Asri, W. N. Suliyanti, D. Kuswardani, and M. Fajri, "Pelabelan Otomatis Lexicon Vader dan Klasifikasi Naive Bayes dalam menganalisis sentimen data ulasan PLN Mobile," PETIR, vol. 15, no. 2, pp. 264–275, Nov. 2022, doi: 10.33322/petir.v15i2.1733.