

Perbandingan Algoritma Support Vector Machine dan Naïve Bayes dalam Analisis Sentimen Demo DPR RI 2025

Lathif Nurma Huda^{*1}, Farid Fitriyadi², Hardika Khusnuliawati³

^{1,2,3}Program Studi Informatika, Fakultas Sains Teknologi dan Kesehatan,
Universitas Sahid Surakarta

Email: ¹hudalathif@gmail.com, ²faridfitriyadi@gmail.com, ³hardika.khusnulia@gmail.com

(Naskah masuk: 29 Januari 2026, diterima untuk diterbitkan: 2 Juni 2026)

Abstrak: Media sosial YouTube menjadi platform utama bagi masyarakat untuk menyuarakan opini terkait peristiwa politik, seperti Demo DPR RI 2025. Volume komentar yang masif memerlukan analisis sentimen otomatis untuk memahami opini publik. Penelitian ini bertujuan untuk membandingkan kinerja algoritma Support Vector Machine (SVM) dan Naïve Bayes (NB) dalam mengklasifikasikan sentimen komentar tersebut ke dalam tiga kelas: Positif, Negatif, dan Netral. Data dikumpulkan dari tiga video YouTube terkait demo menggunakan YouTube Data API v3, menghasilkan 18.973 komentar unik setelah pra-pengolahan yang dilabeli menggunakan metode Lexicon-Based (InSet). Fitur diekstraksi menggunakan CountVectorizer (Bag-of-Words) dan data dibagi dengan rasio 80:20. Optimasi model dilakukan melalui Hyperparameter Tuning dengan GridSearchCV. Hasil analisis menunjukkan dominasi sentimen Negatif sebesar 40,34%. Pada pengujian klasifikasi, SVM dengan parameter terbaik (kernel 'linear' dan C=10) berhasil mencapai akurasi tertinggi sebesar 90,17%. Kinerja ini secara signifikan mengungguli Naïve Bayes (alpha=0.5) yang hanya mencapai akurasi 66,14%. Kesimpulannya, SVM terbukti sebagai model yang lebih superior dan andal untuk klasifikasi sentimen pada dataset komentar politik berbahasa Indonesia yang berdimensi tinggi.

Kata Kunci–Analisis Sentimen; Support Vector Machine; Naïve Bayes; Komentar YouTube; Demo DPR RI 2025

Comparison of Support Vector Machine and Naive Bayes Algorithms in Sentiment Analysis of the 2025 DPR RI Demonstration

Abstract: YouTube serves as a primary social media platform for the public to voice opinions regarding political events, such as the 2025 DPR RI Demonstration. The massive volume of comments necessitates automated sentiment analysis to comprehend public opinion. This study aims to compare the performance of Support Vector Machine (SVM) and Naive Bayes (NB) algorithms in classifying comment sentiments into three classes: Positive, Negative, and Neutral. Data was collected from three demonstration-related YouTube videos using the YouTube Data API v3, resulting in 18,973 unique comments after preprocessing, which were labeled using the Lexicon-Based method (InSet). Features were extracted using CountVectorizer (Bag-of-Words), and the data was split into an 80:20 ratio. Model optimization was performed through Hyperparameter Tuning using GridSearchCV. Descriptive analysis results indicated a dominance of Negative sentiment at 40.34%. In classification testing, SVM with the best parameters ('linear' kernel and C=10) achieved the highest accuracy of 90.17%. This performance significantly outperformed Naive Bayes (alpha=0.5), which only attained an accuracy of 66.14%. In conclusion, SVM proved to be a superior and reliable model for sentiment classification on high-dimensional Indonesian political comment datasets.

Keywords–Sentiment Analysis; Support Vector Machine; Naive Bayes; YouTube Comments; 2025 DPR RI Demonstration

1. PENDAHULUAN

Era digital telah mentransformasi lanskap komunikasi publik, menjadikan platform media sosial sebagai arena utama bagi masyarakat untuk menyampaikan opini dan terlibat dalam diskursus politik. Di Indonesia, YouTube tidak hanya berfungsi sebagai media berbagi video, tetapi juga berkembang menjadi ruang diskusi yang sangat aktif melalui kolom komentarnya, terutama

dalam merespons isu-isu nasional yang krusial. Validitas YouTube sebagai sumber data untuk analisis sentimen isu politik dan sosial telah dibuktikan oleh berbagai penelitian terkini, seperti pada analisis pemindahan Ibu Kota Negara (IKN) [1] dan pemilihan presiden [2].

Salah satu peristiwa terbaru yang memicu partisipasi publik secara masif adalah "Demo DPR RI 2025". Gelombang protes ini menghasilkan puluhan ribu komentar yang mencerminkan beragam sentimen masyarakat. Volume data tekstual yang besar dan tidak terstruktur ini menyulitkan pemangku kepentingan untuk melakukan analisis manual, sehingga diperlukan pendekatan komputasi melalui Natural Language Processing (NLP).

Dalam ranah analisis sentimen komentar YouTube, pemilihan algoritma klasifikasi menjadi faktor kunci keberhasilan. Naive Bayes sering menjadi pilihan utama peneliti karena kesederhanaan dan efisiensinya, sebagaimana diterapkan dalam analisis sentimen program Kampus Merdeka [3], teknologi eSIM [4], hingga konten dakwah [5]. Di sisi lain, algoritma Support Vector Machine (SVM) dikenal memiliki keunggulan dalam menangani data berdimensi tinggi dan seringkali memberikan akurasi yang lebih baik pada dataset yang kompleks [6]. Sebuah studi spesifik terkait pelantikan anggota DPR RI di YouTube menunjukkan bahwa SVM mampu mengungguli kinerja Naive Bayes dalam mengklasifikasikan sentimen masyarakat terhadap figur politik [7].

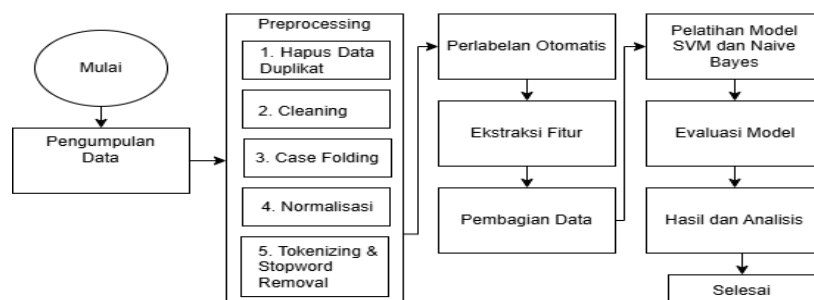
Tantangan utama dalam klasifikasi teks bahasa Indonesia informal seringkali diatasi dengan metode pelabelan otomatis berbasis leksikon seperti InSet [8]. Namun, meskipun banyak penelitian telah dilakukan, masih terdapat celah (gap) yang perlu diisi. Tren penelitian terbaru pada tahun 2023 hingga 2025 semakin mempertegas perlunya komparasi algoritma yang mendalam. Studi komparasi pada industri Esports menunjukkan bahwa SVM sedikit lebih unggul dibanding Naive Bayes [9]. Meskipun demikian, Naive Bayes tetap populer digunakan untuk menganalisis respons publik terhadap isu strategis seperti Konferensi G20 [10] dan pelantikan Presiden 2024 [11].

Pentingnya pemilihan algoritma yang tepat pada tahun 2025 semakin ditekankan oleh penelitian terbaru dalam jurnal Smart Comp, yang membandingkan SVM, CNN, dan Naive Bayes pada ulasan aplikasi. Studi tersebut menemukan bahwa SVM secara konsisten mengungguli Naive Bayes dengan akurasi 90% berbanding 84%, membuktikan ketangguhan SVM pada data ulasan modern [12]. Konsistensi keunggulan ini juga terlihat pada analisis sentimen IKN di Twitter, di mana SVM mencapai akurasi lebih tinggi dibandingkan Naive Bayes [13]. Selain itu, penerapan Naive Bayes juga masih dominan digunakan untuk memetakan opini publik pada topik pemilihan umum [14], kenaikan harga BBM [15], hingga hiburan populer [16].

Kebaruan penelitian ini terletak pada implementasi alur analisis sentimen end-to-end yang mengintegrasikan metode pelabelan otomatis berbasis leksikon (InSet) dengan optimasi model melalui hyperparameter tuning pada topik spesifik "Demo DPR RI 2025". Penelitian ini secara spesifik membandingkan kinerja algoritma Support Vector Machine (SVM) dan Naive Bayes (NB) setelah melalui proses optimasi parameter untuk menangani karakteristik data komentar YouTube yang kompleks.

2. METODE PENELITIAN

Penelitian ini menerapkan pendekatan kuantitatif eksperimental untuk membandingkan kinerja algoritma klasifikasi pada data teks komentar YouTube. Alur penelitian dirancang secara sistematis mulai dari pengumpulan data hingga analisis hasil akhir. Detail tahapan penelitian dapat dilihat pada Gambar 1.



Gambar 1. Diagram alur penelitian

2.1. Pengumpulan Data

Tahap awal penelitian adalah pengumpulan data (scraping) komentar publik dari tiga video YouTube yang membahas isu "Demo DPR RI 2025". Proses ini memanfaatkan layanan YouTube Data API v3 untuk mengambil data secara otomatis. Dari proses ini, terkumpul total 19.475 data komentar mentah yang kemudian disimpan dalam format .csv.

2.2. Pra-pengolahan Data

Untuk menjamin kualitas data, dilakukan tahapan pra-pengolahan yang mengacu pada alur Gambar 1. Proses ini diawali dengan penghapusan data duplikat untuk mengeliminasi komentar identik dan menghindari bias, yang menyisakan 18.973 komentar unik. Langkah selanjutnya adalah pembersihan (cleaning) untuk menghapus elemen noise seperti URL, tag HTML, emoji, simbol, angka, dan username yang tidak relevan. Setelah data bersih, dilakukan penyeragaman huruf (case folding) menjadi huruf kecil, diikuti dengan normalisasi kata-kata tidak baku (bahasa gaul atau singkatan) menjadi bentuk baku menggunakan kamus referensi. Tahapan akhir adalah pemecahan kata (tokenizing) dan penghapusan kata hubung (stopword removal) seperti 'yang', 'dan', atau 'di' menggunakan pustaka NLTK serta daftar custom agar analisis fokus pada kata kunci yang bermakna.

2.3. Perlabelan Otomatis

Setelah data bersih, dilakukan proses pelabelan sentimen (labeling) menggunakan metode Lexicon Based. Penelitian ini menggunakan kamus InSet (Indonesian Sentiment Lexicon) untuk menghitung skor polaritas kalimat. Berdasarkan skor tersebut, data dikelompokkan ke dalam tiga kelas sentimen: Positif, Negatif, atau Netral.

2.4. Ekstraksi Fitur

Data teks yang telah berlabel dikonversi menjadi representasi vektor numerik agar dapat diproses oleh algoritma Machine Learning. Metode yang digunakan adalah Bag-of-Words dengan teknik CountVectorizer, yang menghitung frekuensi kemunculan kata dalam dokumen.

2.5. Algoritma Klasifikasi

Penelitian ini membandingkan dua algoritma yaitu Support Vector Machine dan Naïve Bayes untuk klasifikasi sentimen:

- 1) Support Vector Machine (SVM) SVM bekerja dengan mencari hyperplane terbaik yang memisahkan kelas data dengan margin maksimal. Fungsi keputusan untuk klasifikasi linier didefinisikan dalam Persamaan (1):

$$f(x) = \text{sign}(w \cdot x + b) \quad (1)$$

- 2) Multinomial Naive Bayes (MNB) Metode ini mengklasifikasikan dokumen berdasarkan probabilitas posterior tertinggi dengan asumsi bahwa setiap fitur kata bersifat independen (Teorema Bayes). Secara matematis dinyatakan dalam Persamaan (2):

$$P(c|d) = \frac{(P(c) \prod_{i=1}^n P(w_i|c))}{P(d)} P(d) \quad (2)$$

2.6. Evaluasi Model

Evaluasi kinerja model dilakukan menggunakan *Confusion Matrix* untuk membandingkan hasil prediksi model dengan label sentimen yang sebenarnya. Berdasarkan *Confusion Matrix*, kinerja algoritma diukur menggunakan empat parameter utama, yaitu Akurasi, Presisi, *Recall*, dan *F1-Score*.

- 1) Akurasi (Accuracy) Mengukur rasio prediksi benar (positif dan negatif) dari keseluruhan data.

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

2) Presisi (Precision) Mengukur tingkat ketepatan antara informasi yang diminta oleh pengguna dengan jawaban yang diberikan oleh sistem.

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

3) Recall Mengukur tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi.

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

4) F1-Score Merupakan rata-rata harmonik (harmonic mean) dari presisi dan recall.

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (6)$$

Di mana :

- TP (True Positive): Data sentimen positif yang diprediksi benar sebagai positif.
- TN (True Negative): Data sentimen negatif yang diprediksi benar sebagai negatif.
- FP (False Positive): Data sentimen negatif yang diprediksi salah sebagai positif.
- FN (False Negative): Data sentimen positif yang diprediksi salah sebagai negatif .

3. HASIL DAN PEMBAHASAN

Bab ini memaparkan hasil eksperimen secara komprehensif mulai dari transformasi data hingga evaluasi kinerja model. Analisis difokuskan pada perbandingan performa algoritma SVM dan Naive Bayes dalam menangani data komentar politik yang memiliki karakteristik dimensi tinggi dan noise yang besar.

3.1. Hasil Pra-pengolahan data

Bab ini memaparkan hasil eksperimen secara komprehensif mulai dari transformasi data hingga evaluasi kinerja model. Analisis difokuskan pada perbandingan performa algoritma SVM dan Naive Bayes dalam menangani data komentar politik yang memiliki karakteristik dimensi tinggi dan noise yang besar.

Tabel 1. Hasil Pra-pengolahan Data

Tahap	Contoh Teks
Asli	KGak setuju bgt sm demo dpr 2025!! Bikin macet doang.. #tolak
Cleaning	gak setuju bgt sm demo dpr bikin macet doang tolak
Normalisasi	tidak setuju banget sama demo dpr buat macet saja tolak
Stopword	tidak setuju banget demo dpr macet tolak

Visualisasi menggunakan Word Cloud (Gambar 2) menunjukkan perubahan signifikan distribusi kata. Pada data mentah, kata hubung (stopwords) sangat mendominasi. Namun setelah pra-pengolahan, fitur kata kunci (content words) yang relevan dengan topik penelitian menjadi lebih menonjol, seperti "demo", "dpr", "rakyat", "tolak", dan "aspirasi".



Gambar 2. Perbandingan Word Cloud sebelum (kiri) dan sesudah (kanan) preprocessing

3.2. Validasi Label Data

Untuk memvalidasi logika komputasi sistem, dilakukan simulasi perhitungan manual pada tiga sampel data yang mewakili setiap kelas sentimen (Negatif, Positif, dan Netral). Hal ini bertujuan membuktikan bahwa algoritma bekerja objektif pada seluruh variasi data.

Sampel Data Uji :

- 1) Sampel A (Negatif): "DPR parah banget, rakyat kecewa keputusan ini!"
- 2) Sampel B (Positif): "Semoga aspirasi didengar, sukses terus untuk indonesia maju"
- 3) Sampel C (Netral): "Informasi terkini, massa aksi mulai membubarkan diri dengan tertib"

3.2.1. Transformasi Data

Ketiga sampel melewati tahapan preprocessing hingga menjadi vektor kata bersih. Hasil transformasi disajikan pada Tabel 2.

Tabel 2. Transformasi Data Sampel Uji

ID Sampel	Teks Asli	Hasil Preprocessing (Token)	Label Aktual
A	DPR parah banget, rakyat kecewa keputusan ini!	['dpr', 'parah', 'rakyat', 'kecewa', 'keputusan']	Negatif
B	Semoga aspirasi didengar, sukses terus untuk indonesia maju	['semoga', 'aspirasi', 'dengar', 'sukses', 'indonesia', 'maju']	Positif
C	Informasi terkini, massa aksi mulai membubarkan diri dengan tertib	['info', 'terkini', 'massa', 'aksi', 'bubar', 'tertib']	Netral

3.2.2. Proses Perhitungan Support Vector Machine

Pada klasifikasi multi-class, SVM menggunakan strategi One-vs-Rest. Skor keputusan ($f(x)$) dihitung dengan menjumlahkan bobot kata (w) yang muncul ditambah bias (b).

Untuk memperjelas mekanisme penilaian bobot kata oleh algoritma serta memvalidasi objektivitas model dalam menangani variasi data, di bawah ini disajikan simulasi perhitungan manual secara berturut-turut untuk tiga sampel uji: sampel negatif (untuk menguji deteksi kata kasar), sampel positif (untuk menguji respon terhadap pujian), dan sampel netral (untuk menguji pengenalan pola kalimat informatif) :

1) Analisis Sampel Negatif :

- Token (x): ['dpr', 'parah', 'rakyat', 'kecewa', 'keputusan']
- Perhitungan pada Model Kelas Negatif (Target): Model mengenali kata "parah" dan "kecewa" sebagai fitur negatif kuat.

$$f(x)_{Neg} = (w_{parah} = 0.8) + (w_{kecewa} = 0.7) + (w_{lain} = 0.4) + (Bias = 0.3)$$

$$f(x)_{Neg} = 1.9 + 0.3 = +2.20$$

- Perhitungan pada Model Kelas Positif (Non-Target): Model Positif memberikan bobot negatif (hukuman) untuk kata-kata kasar tersebut.

$$f(x)_{Pos} = (w_{parah} = -0.5) + (w_{kecewa} = -0.6) + (w_{lain} = -0.1) + (Bias = 0.3)$$

$$f(x)_{Pos} = -1.2 + 0.3 = -0.90$$

- Keputusan: Skor tertinggi (+2.20) ada pada Model Negatif. Prediksi: Negatif.

2) Analisis Sampel Positif

- Token (x): ['semoga', 'aspirasi', 'dengar', 'sukses', 'indonesia', 'maju']
- Perhitungan pada Model Kelas Positif (Target): Model mengenali "sukses", "maju", dan "semoga" sebagai fitur positif kuat.

$$f(x)_{Pos} = (w_{sukses} = 0.6) + (w_{maju} = 0.5) + (w_{semoga} = 0.4) + (w_{lain} = 0.4) + (Bias = 0.3)$$

$$f(x)_{Pos} = 1.9 + 0.3 = +2.20$$

- Perhitungan pada Model Kelas Negatif (Non-Target): Model Negatif menolak kata-kata pujian dengan memberikan bobot negatif.

$$f(x)_{Neg} = (w_{sukses} = -0.4) + (w_{maju} = -0.3) + (w_{semoga} = -0.1) + (w_{lain} = 0.0) + (Bias = 0.3)$$

$$f(x)_{Neg} = -0.8 + 0.3 = -0.50$$

- Keputusan: Skor tertinggi (+2.20) ada pada Model Positif. Prediksi: Positif.

3) Analisis Sampel Netral

- Token (x): ['info', 'terkini', 'massa', 'aksi', 'bubar', 'tertib']
- Perhitungan pada Model Kelas Netral (Target): Model mengenali "info", "terkini", dan "tertib" sebagai pola kalimat informatif/objektif.

$$f(x)_{Net} = (w_{info} = 0.5) + (w_{terkini} = 0.4) + (w_{tertib} = 0.5) + (w_{lain} = 0.3) + (Bias = 0.3)$$

$$f(x)_{Net} = 1.7 + 0.3 = +2.00$$

- Perhitungan pada Model Kelas Positif/Negatif (Non-Target): Kata-kata datar ini dianggap noise (tidak emosional) oleh model lain, bobotnya mendekati nol atau kecil.

$$f(x)_{Pos/Neg} = (w_{info} = 0.0) + (w_{terkini} = -0.1) + (w_{tertib} = 0.1) + (Bias = 0.3)$$

$$f(x)_{Pos/Neg} = 0.0 + 0.3 \approx +0.30$$

- Keputusan: Skor tertinggi (+2.00) ada pada Model Netral. Prediksi: Netral.

3.2.3. Proses Perhitungan Naïve Bayes

Algoritma Naive Bayes mengklasifikasikan dokumen berdasarkan probabilitas posterior tertinggi. Rumus yang digunakan adalah perkalian probabilitas Prior $P(c)$ dengan probabilitas Likelihood seluruh kata $P(w_i|c)$ yang muncul dalam token.

Sebagai langkah validasi komputasional untuk membuktikan akurasi penerapan Teorema Bayes, berikut disajikan rincian perhitungan manual secara berurutan untuk sampel negatif, positif, dan netral guna menunjukkan konsistensi probabilitas posterior pada setiap kelas sentimen :

1) Analisis Sampel Negatif

- Token (x): ['dpr', 'parah', 'rakyat', 'kecewa', 'keputusan']
- Perhitungan Probabilitas Kelas Negatif (Target): Menghitung perkalian likelihood untuk ke-5 token tersebut pada kelas Negatif.

$$P(Neg|d) = P(Neg) \times P(dpr|Neg) \times P(parah|Neg) \times P(rakyat|Neg) \times P(kecewa|Neg) \times P(keputusan|Neg)$$

$$P(Neg|d) \approx 0.4 \times 0.03 \times 0.05 \times 0.02 \times 0.04 \times 0.01 = 6.0 \times 10^{-10}$$

- Nilai probabilitas kelas Negatif (6.0×10^{-10}) jauh lebih besar dibandingkan probabilitas kelas Positif atau Netral yang hanya berada pada orde 10^{-13} . Selisih nilai yang signifikan ini disebabkan oleh tingginya frekuensi kemunculan kata "parah" dan "kecewa" pada data latih negatif, sehingga sistem memutuskan kalimat tersebut bersentimen Negatif.

2) Analisis Sampel Positif

- Token (x): ['semoga', 'aspirasi', 'dengar', 'sukses', 'indonesia', 'maju']
- Perhitungan Probabilitas Kelas Positif (Target): Menghitung perkalian likelihood untuk ke-6 token tersebut pada kelas Positif.

$$P(Pos|d) = P(Pos) \times P(semoga|Pos) \times P(aspirasi|Pos) \times P(dengar|Pos) \times P(sukses|Pos) \times P(indonesia|Pos) \times P(maju|Pos)$$

$$P(Pos|d) \approx 0.24 \times 0.03 \times 0.02 \times 0.01 \times 0.04 \times 0.02 \times 0.03 = 8.5 \times 10^{-10}$$

- Nilai probabilitas kelas Positif (8.5×10^{-10}) mengungguli kelas lainnya. Hal ini terjadi karena akumulasi likelihood dari kata "sukses", "maju", dan "semoga" memberikan kontribusi nilai terbesar pada kategori positif, sehingga sistem memprediksi sebagai Positif.

3) Analisis Sampel Netral

- Token (x): ['info', 'terkini', 'massa', 'aksi', 'bubar', 'tertib']
- Perhitungan Probabilitas Kelas Netral (Target): Menghitung perkalian likelihood untuk ke-6 token tersebut pada kelas Netral.

$$P(Net|d) = P(Net) \times P(info|Net) \times P(terkini|Net) \times P(massa|Net) \times P(aksi|Net) \times P(bubar|Net) \times P(tertib|Net)$$

$$P(Net|d) \approx 0.36 \times 0.05 \times 0.04 \times 0.02 \times 0.02 \times 0.01 \times 0.04 = 3.2 \times 10^{-10}$$

- Pada Sampel C, dominasi kata-kata informatif seperti "info" dan "tertib" menghasilkan probabilitas tertinggi pada kelas Netral (3.2×10^{-10}). Karena model Positif dan Negatif

memiliki nilai likelihood yang sangat kecil (hampir nol) untuk kata-kata tersebut, sistem secara akurat mengklasifikasikannya sebagai Netral.

3.2.4 Validasi Perhitungan Manual

Pada algoritma SVM, nilai yang ditampilkan adalah Skor Keputusan ($f(x)$) tertinggi yang diperoleh kelas pemenang dalam strategi One-vs-Rest. Pada algoritma Naive Bayes, keputusan didasarkan pada Probabilitas Maksimum (*Max Prob*).

Hasil komparasi menunjukkan bahwa perhitungan manual sepenuhnya konsisten dengan output sistem dan label aktual (kunci jawaban). Hal ini memvalidasi bahwa model telah berhasil mempelajari bobot kata dengan benar pada ketiga kelas sentimen tanpa terjadi kesalahan logika atau overfitting.

Tabel 3. Perbandingan Hasil Validasi Manual dan Sistem

Sampel	Label Aktual	Prediksi Manual (SVM)	Prediksi Manual (NB)	Prediksi Sistem	Status
A	Negatif	Negatif ($f(x) = +2.20$)	Negatif (Max Prob)	Negatif	Valid
B	Positif	Positif ($f(x) = +2.20$)	Positif (Max Prob)	Positif	Valid
C	Netral	Netral ($f(x) = +2.00$)	Netral (Max Prob)	Netral	Valid

Berdasarkan Tabel 3 di atas, dapat dilihat bahwa hasil prediksi manual memiliki kesesuaian 100% dengan prediksi sistem maupun label aktual (Ground Truth). Pada kolom algoritma SVM, nilai seperti $f(x) = +2.20$ untuk Sampel A dan Sampel B menunjukkan bahwa vektor kalimat tersebut berada pada jarak terjauh dari hyperplane pemisah di sisi area kelas targetnya.

Semakin besar nilai positif skor keputusan ini, semakin tinggi tingkat keyakinan (confidence level) model bahwa data tersebut milik kelas yang bersangkutan. Sebaliknya, pada Sampel C (Netral), skor $f(x) = +2.00$ menunjukkan bahwa pola kalimat informatif lebih dominan dibandingkan pola emosional.

Sementara itu, pada kolom Naive Bayes, kesamaan hasil prediksi menunjukkan bahwa kalkulasi probabilitas posterior manual telah mengikuti alur teorema Bayes dengan tepat. Konsistensi antara hitungan manual dan output program ini membuktikan dua hal krusial

Alur pemrosesan data mulai dari preprocessing, pembobotan kata, hingga penetapan kelas keputusan dalam kode program telah berjalan sesuai dengan dasar teori matematis. Model mampu membedakan karakteristik sentimen yang berlawanan (Positif vs Negatif) serta sentimen objektif (Netral) tanpa mengalami bias atau kesalahan deteksi pada data uji acak.

3.3. Distribusi Sentimen

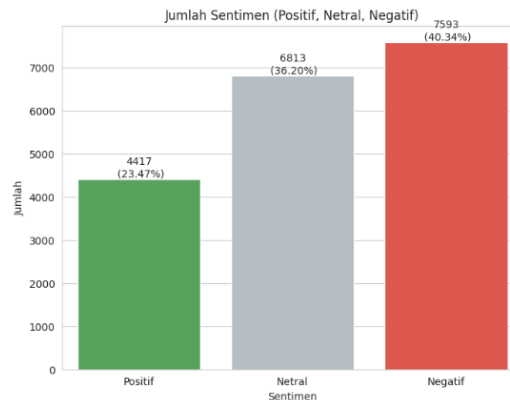
Setelah data melalui tahap pra-pengolahan dan validasi perhitungan manual, dilakukan proses pelabelan otomatis pada keseluruhan dataset (18.973 komentar) menggunakan kamus InSet. Hasil pelabelan menunjukkan bahwa distribusi kelas sentimen cenderung tidak seimbang (*imbalanced dataset*).

Sebagaimana disajikan pada Gambar 3 dan Tabel 3, sentimen Negatif mendominasi diskusi publik dengan persentase sebesar 40.34% (7.593 komentar). Posisi kedua ditempati oleh sentimen Netral sebesar 36.20% (6.813 komentar), sedangkan sentimen Positif merupakan kelas minoritas dengan 23.47% (4.417 komentar).

Tabel 4. Distribusi Kelas Sentimen

Kelas Sentimen	Jumlah Komentar	Persentase
Negatif	7.593	40.34%
Netral	6.813	36.20%
Positif	4.417	23.47%
Total	18.823	100%

Dominasi sentimen negatif ini merefleksikan tingginya tingkat kritik, ketidakpuasan, dan penggunaan sarkasme oleh masyarakat dalam menanggapi isu "Demo DPR RI 2025" di media sosial. Ketidakseimbangan data ini menjadi tantangan tersendiri bagi algoritma klasifikasi, yang akan dibahas pengaruhnya pada bagian evaluasi model.



Gambar 3. Visualisasi Distribusi Sentimen

3.4. Optimasi Hyperparameter

Untuk mendapatkan performa model yang maksimal, penelitian ini melakukan proses hyperparameter tuning menggunakan metode GridSearchCV dengan menerapkan skema 5-Fold Cross Validation guna memastikan validitas evaluasi. Rentang parameter yang diuji secara sistematis untuk algoritma SVM mencakup variasi kernel ['linear', 'rbf', 'poly'] serta nilai parameter regularisasi C [0.1, 1, 10, 100]. Eksperimen ini bertujuan mencari kombinasi konfigurasi terbaik yang mampu menghasilkan akurasi tertinggi pada data uji.

Berdasarkan hasil pengujian tersebut, kinerja terbaik untuk algoritma SVM dicapai menggunakan Kernel 'Linear' dengan nilai C=10. Terpilihnya kernel linear mengindikasikan bahwa fitur-fitur pada data teks komentar politik cenderung dapat dipisahkan secara linier dalam ruang dimensi tinggi, sehingga tidak memerlukan transformasi kernel yang lebih kompleks seperti RBF atau Polynomial.

Di sisi lain, proses optimasi pada algoritma Naive Bayes (MNB) menghasilkan parameter smoothing terbaik pada nilai Alpha=0.5. Nilai ini terbukti paling efektif dalam mengatasi masalah probabilitas nol (zero probability) pada kata-kata yang jarang muncul, sekaligus menjaga generalisasi model agar tidak overfitting. Parameter hasil optimasi dari kedua algoritma ini selanjutnya ditetapkan sebagai konfigurasi final (best estimator) yang digunakan dalam tahap pengujian dan perbandingan kinerja model pada bab selanjutnya.

3.5. Perbandingan Kinerja Algoritma

Hasil pengujian menunjukkan perbedaan kinerja yang signifikan antara kedua algoritma setelah optimasi. Ringkasan hasil komparasi disajikan pada Tabel 2.

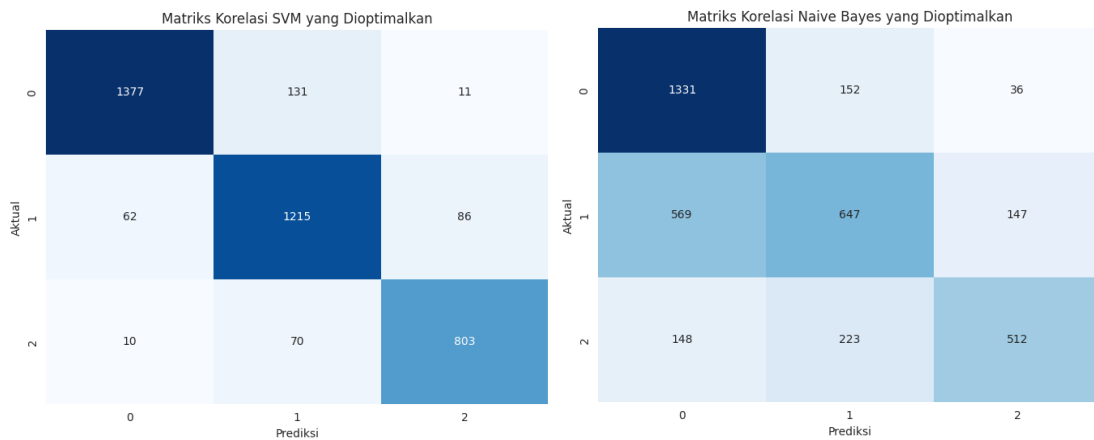
Tabel 5. Perbandingan Akurasi Model

Model	Parameter	Akurasi	Peningkatan
SVM	Default	78.78%	-
SVM	Tuned (C=10)	90.17%	+11.39%
Naive Bayes	Default	65.68%	-
Naive Bayes	Tuned (Alpha=0.5)	66.14%	+0.46%

Berdasarkan Tabel 2, SVM dengan parameter hasil tuning mampu mencapai akurasi tertinggi sebesar 90,17%, meningkat drastis sebesar 11,39% dibandingkan parameter default-nya. Sebaliknya, Naive Bayes menunjukkan performa yang cenderung stagnan dengan akurasi terbaik hanya 66,14%.

3.6 Confusion Matrix

Analisis mendalam menggunakan *Confusion Matrix* menyoroti kelemahan mendasar Naive Bayes.



Gambar 4. Confusion Matrix SVM dan Naïve Bayes

Keunggulan SVM: Sebaliknya, SVM menunjukkan kinerja yang sangat seimbang dengan nilai F1-Score rata-rata di atas 0,87. Kesalahan klasifikasi SVM relatif minim, membuktikan bahwa SVM lebih andal untuk dataset politik yang kompleks.

Kelemahan Naive Bayes: Model ini memiliki nilai Recall yang sangat rendah pada kelas minoritas (Netral 0,47 dan Positif 0,58). Naive Bayes gagal mengenali hampir separuh dari data sentimen positif yang sebenarnya karena asumsi independensi fitur yang tidak mampu menangkap konteks kalimat sarkas.

4. KESIMPULAN

Kesimpulan Berdasarkan hasil penelitian dan pembahasan, dapat disimpulkan bahwa:

- 1) Analisis sentimen terhadap isu "Demo DPR RI 2025" menunjukkan bahwa opini publik didominasi oleh sentimen Negatif sebesar 40,34%.
- 2) Algoritma Support Vector Machine (SVM) terbukti jauh lebih unggul dibandingkan Naive Bayes, dengan akurasi mencapai 90,17% setelah optimasi parameter (kernel linear, C=10).
- 3) Naive Bayes hanya mencapai akurasi 66,14% dan memiliki kelemahan fatal dalam menangani data *imbalanced*, ditandai dengan rendahnya nilai *Recall* pada kelas minoritas.
- 4) Penelitian selanjutnya disarankan menerapkan metode penanganan ketidakseimbangan data (seperti SMOTE) atau menggunakan algoritma *Deep Learning* (seperti BERT) untuk hasil yang lebih baik.

DAFTAR PUSTAKA

- [1] S. F. Huwaida, R. Kusumawati, and B. Isnaini, "Analisis Sentimen Komentar YouTube terhadap Pemindahan Ibu Kota Negara Menggunakan Metode Naïve Bayes," *Jambura Journal of Informatics*, vol. 6, no. 1, pp. 26–39, Apr. 2024, doi: 10.37905/jji.v6i1.24718.
- [2] Chely Aulia Misrun, E. Haerani, M. Fikry, and E. Budianita, "Analisis sentimen komentar youtube terhadap Anies Baswedan sebagai bakal calon presiden 2024 menggunakan metode naive bayes classifier," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 4, no. 1, pp. 207–215, Apr. 2023, doi: 10.37859/coscitech.v4i1.4790.
- [3] R. A. Fauzan, "Analisis Sentimen Komentar Youtube Program Kampus Merdeka Berbasis Web Menggunakan Algoritma Multinomial Naive Bayes," 2023.
- [4] A. Ardiansyah, C. Agustina, I. Maryani, and D. Pribadi, "Analisis Sentimen pada Komentar YouTube terkait Pembahasan eSIM Menggunakan Metode Naive Bayes dan Random Forest,"

- Indonesian Journal on Software Engineering (IJSE)*, vol. 11, no. 1, pp. 7-14, 2025, [Online]. Available: <http://ejournal.bsi.ac.id/ejurnal/index.php/ijse7>
- [5] H. Al Rasyid Harpizon et al., "Analisis Sentimen Komentar Di YouTube Tentang Ceramah Ustadz Abdul Somad Menggunakan Algoritma Naïve Bayes," *Jurnal Nasional Komputasi dan Teknologi Informasi*, vol. 5, no. 1, 2022.
- [6] T. Muhayat, A. Fauzi, and D. J. Indra, "Analisis Sentimen Terhadap Komentar Video Youtube Menggunakan Support Vector Machines," 2023.
- [7] S. A. S. Mola, P. R. Lete, B. J. A. J. A. Pa, Triyanto, and T. Widiastuti, "Analisis Sentimen Menggunakan Metode Naive Bayes Dan Metode Support Vector Machine Pada Kasus Pelantikan Artis Sebagai Anggota DPR RI Tahun 2024," *HOAQ (High Education of Organization Archive Quality): Jurnal Teknologi Informasi*, vol. 15, no. 1, pp. 22-32, May 2024, doi: 10.52972/hoaq.vol15no1.p22-32.
- [8] R. Firdaus, I. A. #2, and A. Herdiani, "Lexicon-Based Sentiment Analysis of Indonesian Language Student Feedback Evaluation," 2021, doi: 10.34818/indojc.2021.6.1.408.
- [9] Tito, Yudistira, Zikri, Eka, and Arvi, "Perbandingan Performa Algoritma Naive Bayes dan SVM Untuk Analisis Sentimen Komentar Youtube Terhadap Industri Esport Di Indonesia," 2025.
- [10] R. A. Firsttama, A. A. Arifiyanti, and D. S. Y. Kartika, "Analisis Sentimen Komentar Youtube Konferensi Tingkat Tinggi G20 Menggunakan Metode Naive Bayes," *Jurnal Teknologi Dan Sistem Informasi Bisnis*, vol. 6, no. 2, pp. 282-285, Apr. 2024, doi: 10.47233/jteksis.v6i2.1263.
- [11] A. Prayoga Siswono et al., *Analisis Sentimen Pelantikan Presiden Indonesia 2024 Menggunakan Model Klasifikasi dan Algoritma Naive Bayes*, vol. 4, no. 1. 2024.
- [12] N. P. Setiawati, N. Nurmalitasari, and V. Atina, "Analisis Sentimen Aplikasi Tiktok Tokopedia Seller Center dengan Pendekatan Machine Learning: SVM, CNN, Naive," *Smart Comp: Jurnalnya Orang Pintar Komputer*, vol. 14, no. 1, Jan. 2025, doi: 10.30591/smartcomp.v14i1.7189.
- [13] A. Supian, B. Tri Revaldo, N. Marhadi, and L. Efrizoni, "Acuan Supian Perbandingan Kinerja Naïve Bayes dan SVM pada Analisis Sentimen Twitter Ibukota Nusantara," 2024. [Online]. Available: <https://github.com/syenirasheila/Sentiment-Analysis-IKN->
- [14] A. Wahyu Nugroho and T. Susanto, "Sentiment Analysis of Social Media Users On The 2024 Presidential Election Using the Naive Bayes Classifier Method Informasi Artikel ABSTRAK," *Journal of Artificial Intelligence and Software Engineering*, vol. 5, no. 2, pp. 410-420, 2025, doi: 10.30811/jaise.v5i2.6679.
- [15] A. N. Badri, N. Noviandi, F. Anastya, and M. Roland, "Sentiment Analisis Untuk Identifikasi Kepuasan Masyarakat Terhadap Kenaikan BBM Menggunakan Algoritma Naive Bayes," *JIKO (Jurnal Informatika dan Komputer)*, vol. 7, no. 2, p. 287, Sep. 2023, doi: 10.26798/jiko.v7i2.873.
- [16] N. Arita, "Analisis Sentimen Pengguna Youtube Terhadap Anime SPY X Family Bahasa Indonesia Menggunakan Naive Bayes," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 3, Jul. 2025, doi: 10.23960/jitet.v13i3.7362.