

ANALISIS SENTIMEN PUBLIK PADA KOMENTAR INSTAGRAM POSTINGAN AKUN TEMPODOTCO TERHADAP PENERAPAN KURIKULUM AI DI INDONESIA MENGGUNAKAN METODE NAÏVE BAYES

Dwi Aditya Budi Listiawan, Wisnu Adi Wicaksono, Wiyanto

Teknik Informatika, Universitas Pelita Bangsa

Jalan Inspeksi Kalimalang, Tegal Danas, No. 009, Bekasi, Jawa Barat, Indonesia

listiawan.312110003@mhs.pelitabangsa.ac.id

ABSTRAK

Penerapan kurikulum kecerdasan buatan (AI) di Indonesia, khususnya melalui mata pelajaran pilihan pada tingkat SD hingga SMA mulai tahun ajaran 2025/2026, menimbulkan beragam respons publik di media sosial. Permasalahan muncul dari banyaknya komentar negatif yang mencerminkan kekhawatiran masyarakat terhadap kesiapan pelaksanaan kebijakan tersebut. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap kurikulum AI melalui komentar pada platform Instagram, dengan menggunakan metode klasifikasi teks berbasis algoritma Naïve Bayes. Data berupa 1.238 komentar dikumpulkan dari akun Instagram @tempodotco, kemudian melalui tahap praproses teks (*cleaning, tokenisasi, stopword removal, stemming*) dan ekstraksi fitur menggunakan TF-IDF. Selanjutnya, tiga varian algoritma *Naïve Bayes* (*MultinomialNB, BernoulliNB, dan ComplementNB*) diujikan dan dievaluasi menggunakan metrik akurasi, presisi, *recall*, dan *F1-score*. Hasil menunjukkan bahwa algoritma *Multinomial Naïve Bayes* memberikan performa terbaik dengan akurasi 79,84% dan *F1-score* tertinggi pada kelas negatif. Mayoritas komentar (sekitar 81%) bernada negatif, menunjukkan perlunya strategi komunikasi dan kesiapan infrastruktur dalam penerapan kurikulum AI agar lebih diterima publik.

Kata kunci : analisis sentimen, klasifikasi, *Naive Bayes*, Instagram, kurikulum AI, opini publik

1. PENDAHULUAN

Integrasi teknologi kecerdasan buatan dalam pendidikan menjadi fokus strategis di Indonesia. Pemerintah melalui Kementerian Pendidikan, Kebudayaan, Riset, dan Teknologi (Kemendikbudristek) aktif mendorong pembaruan kurikulum yang mengakomodasi kemajuan teknologi digital. Mulai tahun ajaran 2025/2026, mata pelajaran coding dan kecerdasan buatan (AI) akan masuk ke dalam kurikulum sebagai mata pelajaran pilihan bagi sekolah yang siap. Langkah ini merupakan bagian dari upaya nasional meningkatkan digitalisasi pembelajaran dan mempersiapkan generasi muda yang lebih adaptif terhadap perkembangan teknologi [1].

Pandemi COVID-19 mempercepat digitalisasi pendidikan dan menegaskan kebutuhan mendesak untuk memasukkan teknologi informasi dan komunikasi secara sistemik dalam proses pembelajaran. Dalam konteks ini, pemahaman terhadap AI menjadi krusial sejak usia dini. Sejumlah negara merespons dengan menyisipkan literasi AI dalam kurikulum dari jenjang dasar hingga menengah sebagai bagian dari strategi nasional pengembangan sumber daya manusia unggul [2]. Hal ini mencerminkan kesadaran global bahwa penguasaan AI perlu ditanamkan sedini mungkin untuk menghasilkan generasi yang siap bersaing di era digital.

Pengajaran AI di sekolah tidak hanya mencakup aspek teknis seperti pemrograman, tetapi juga dimensi etika, sosial, dan filosofis yang relevan untuk penggunaan teknologi cerdas secara bertanggung jawab. literasi AI melibatkan pemahaman interaksi dengan AI, berpikir komputasional, literasi data, serta

etika AI [2]. Fakta ini menggambarkan adanya konsensus internasional bahwa AI merupakan literasi baru yang setara pentingnya dengan literasi baca-tulis dan numerasi, yang harus dimiliki oleh generasi mendatang.

Indonesia pun mengambil langkah progresif dalam penerapan kebijakan ini. Mulai tahun ajaran 2025/2026, Kemendikbudristek menetapkan mata pelajaran coding dan AI sebagai mata pelajaran pilihan mulai kelas 5 SD [3]. Kebijakan ini adalah bagian dari program Merdeka Belajar yang menekankan fleksibilitas kurikulum agar relevan dengan kebutuhan masa depan. Dalam berbagai pernyataan publik, termasuk oleh Wakil Presiden Gibran Rakabuming Raka, penguasaan AI disebut sebagai kunci peningkatan produktivitas dan daya saing nasional, serta kunci bagi lahirnya generasi muda yang inovatif dan adaptif secara digital [4].

Namun demikian, adopsi masif kurikulum AI menimbulkan tantangan baru. Kesiapan infrastruktur menjadi kendala utama: banyak sekolah di daerah terpencil belum memiliki fasilitas teknologi dan koneksi internet memadai [1]. Selain itu, kompetensi guru juga menjadi perhatian penting; tidak semua tenaga pendidik memiliki latar belakang teknologi sehingga dibutuhkan pelatihan khusus agar mereka siap mengajarkan coding dan AI [5]. Kekhawatiran ini diperparah oleh ketidakmerataan sarana pendidikan di seluruh wilayah Indonesia, yang dapat menghambat implementasi kebijakan tersebut.

Kebijakan kurikulum AI ini mendapat respons publik yang beragam. Media sosial, sebagai ruang wacana digital, berperan penting dalam menangkap persepsi masyarakat. Instagram, dengan jutaan

pengguna aktif di Indonesia diperkirakan sekitar 103 juta pengguna, menempatkan Indonesia pada urutan ke-4 dunia menjadi platform utama penyebaran opini public [6]. Instagram memungkinkan penyebaran cepat dukungan maupun kritik terhadap kebijakan pendidikan digital. Sifatnya yang dinamis dan real-time membuat Instagram berfungsi sebagai barometer emosi kolektif, sekaligus sumber data empirik yang kaya untuk dianalisis secara ilmiah.

Untuk menanggapi fenomena tersebut, analisis sentimen berbasis pembelajaran mesin muncul sebagai pendekatan yang relevan dan efektif. Metode ini memungkinkan ekstraksi otomatis opini publik dari beragam teks komentar, sehingga diperoleh pemahaman kuantitatif yang lebih objektif terhadap persepsi Masyarakat [7]. Pendekatan ini memanfaatkan algoritma ML yang, dengan praproses teks seperti *stemming*, *tokenisasi*, dan penghapusan *stopwords*, dapat mencapai akurasi tinggi dalam klasifikasi sentimen data media sosial.

Salah satu algoritme yang banyak digunakan untuk klasifikasi teks pendek seperti komentar Instagram adalah *Naïve Bayes*. Algoritma ini mengandalkan prinsip probabilitas sederhana, namun sangat efisien dalam mengolah data teks berukuran besar. *Naïve Bayes* lebih unggul dalam kecepatan pelatihan dan performanya tetap kompetitif dibandingkan model yang lebih kompleks, terutama pada data yang tidak seimbang [7]. Penelitian lokal oleh Khairunnisa dkk. bahkan memperlihatkan bahwa *Naïve Bayes* mampu mengklasifikasikan sentimen komentar Instagram dengan akurasi yang tinggi [8], menjadikannya alat analisis yang andal dan ekonomis.

Berdasarkan latar belakang tersebut, penelitian ini diarahkan untuk mengukur dan menganalisis sentimen publik terhadap kebijakan kurikulum AI di Indonesia dengan studi kasus komentar pada akun Instagram @tempodotco. Komentar yang dihimpun akan diklasifikasikan ke dalam kategori sentimen positif dan negatif menggunakan tiga varian model *Naïve Bayes* (*MultinomialNB*, *BernoulliNB*, dan *ComplementNB*). Tujuannya adalah mengevaluasi distribusi sentimen sekaligus membandingkan kinerja ketiga model tersebut dalam mengolah data sosial yang bersifat informal dan bising. Hasil penelitian diharapkan memberikan gambaran empirik komprehensif tentang persepsi masyarakat terhadap kebijakan kurikulum AI. Pemahaman persepsi publik berbasis data ini dapat membantu pembuat kebijakan merancang strategi komunikasi dan implementasi yang lebih inklusif, transparan, dan responsif. Selain itu, temuan penelitian diharapkan memberikan kontribusi metodologis dalam analisis opini publik digital serta menjadi referensi bagi studi lanjutan yang mengintegrasikan pendekatan ilmu sosial komputasional dalam evaluasi kebijakan berbasis teknologi [7].

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Metode *Naïve Bayes* merupakan salah satu algoritma klasifikasi yang banyak digunakan dalam analisis sentimen media sosial di Indonesia. Wahyudi et al. melakukan penelitian terhadap tweet yang membahas Program Indonesia Pintar Kuliah (KIP-K) dengan menerapkan metode *Naïve Bayes*. Hasil penelitian tersebut menunjukkan tingkat akurasi mencapai 84,99%, menegaskan bahwa pendekatan ini cukup efektif dalam memetakan opini publik terkait program pendidikan pemerintah [9].

Dalam konteks lain, Arum et al. melakukan analisis sentimen publik Indonesia terhadap ajang olahraga internasional SEA Games 2023 dengan data dari *Twitter*. Dengan pembagian data latih dan data uji 40:60, penelitian ini menghasilkan akurasi tertinggi sebesar 92,70%. Temuan ini menunjukkan bahwa *Naïve Bayes* mampu mengklasifikasikan sentimen dalam skenario opini massal terkait *event* besar [10].

Selanjutnya, Prihandana dan Lestari meneliti komentar pengguna di *Instagram* terkait *brand* kosmetik lokal *Rose All Day*. Dengan menerapkan *Naïve Bayes*, mereka menemukan distribusi sentimen sebesar 53,53% positif, 25,84% netral, dan 20,63% negatif. Penelitian ini memperluas bukti bahwa *Naïve Bayes* tetap relevan meskipun diterapkan pada platform berbasis visual dengan format teks singkat seperti *Instagram* [11].

Studi oleh Pandunata et al. berfokus pada opini publik mengenai Pekan Olahraga Nasional (PON) XX Papua 2021 yang juga dianalisis dari media sosial *Instagram*. Dengan konfigurasi data 90% untuk pelatihan dan 10% untuk pengujian, metode *Naïve Bayes* yang digunakan dalam penelitian ini mampu mencapai akurasi sekitar 75%, memperlihatkan kemampuan algoritma ini dalam menangani data dengan volume yang lebih kecil [12].

Terakhir, Akbar dan Sugiharto [5] menganalisis sentimen pengguna *Twitter* terhadap layanan ChatGPT di Indonesia. Penelitian ini membandingkan metode *Naïve Bayes* dan C4.5, di mana *Naïve Bayes* menunjukkan kinerja yang kompetitif dalam memproses teks dengan struktur bahasa tidak baku seperti slang dan singkatan yang umum ditemukan dalam percakapan di *Twitter* [13].

2.2. Kebijakan Kurikulum AI di Indonesia

Pemerintah Indonesia semakin mengintegrasikan literasi kecerdasan buatan (KA) ke dalam kurikulum nasional. Naskah akademik Kemendikbud Ristek (2023) menegaskan inisiatif penguatan literasi digital, coding, dan AI pada kurikulum pendidikan dasar dan menengah sebagai langkah strategis untuk meningkatkan daya saing sumber daya manusia. Kebijakan ini selaras dengan visi pembangunan SDM berkompeten di era digital, termasuk pengajaran aspek etika AI dalam Pendidikan. Sebagai contoh, pemerintah mendorong pemanfaatan AI untuk mempersonalisasi pembelajaran, sehingga

pengalaman belajar siswa dapat disesuaikan dengan kebutuhan individual. Langkah ini didukung oleh data bahwa Indonesia sedang memperbarui kurikulum untuk memasukkan literasi dan etika AI sebagai bagian dari strategi jangka panjang (misalnya menyertakan materi kode dan KA) [14]. Dengan demikian, kebijakan pendidikan Indonesia mencerminkan tren global dalam memasukkan AI sebagai kompetensi kunci.

2.3. Media Sosial (Instagram) sebagai Sumber Opini Publik

Media sosial telah menjadi saluran utama ekspresi opini publik, dan Instagram khususnya memiliki peran strategis dalam konteks Indonesia. Pada Januari 2024 terdapat sekitar 89,9 juta pengguna Instagram di Indonesia (31,6% populasi), menjadikannya salah satu negara dengan pengguna Instagram terbanyak di dunia. Indonesia menempati peringkat ke-4 global pengguna Instagram (total 1,6 miliar pengguna). Platform ini memungkinkan masyarakat mengekspresikan pendapat melalui *posting*, *komentar*, dan *tagar*, sehingga sering dimanfaatkan untuk pemantauan opini publik. Dengan basis data teks yang luas dan *real-time* dari Instagram, peneliti dan pembuat kebijakan dapat menganalisis sentimen publik terhadap isu tertentu secara lebih dinamis dan akurat [15].

2.4. Analisis Sentimen Berbasis Machine Learning

Analisis sentimen adalah proses klasifikasi teks menjadi kategori positif, negatif, atau netral berdasarkan emosi atau sikap yang terkandung [16]. Metode ini memanfaatkan algoritma machine learning untuk memproses data teks berskala besar secara otomatis. Dengan pendekatan ML, sistem komputer dapat mengenali pola bahasa dan makna emosional dalam opini publik di media sosial. Secara umum, penelitian menunjukkan bahwa kombinasi teknik NLP (natural language processing) dan ML dapat secara efektif menginterpretasi opini Masyarakat [17]. Sebagai hasilnya, analisis sentimen berbasis ML banyak digunakan untuk memantau persepsi publik terhadap kebijakan, produk, atau tren sosial secara kuantitatif dan real-time.

2.5. Kinerja Algoritma Naïve Bayes dalam Klasifikasi

Algoritma *Naive Bayes* adalah model klasifikasi berbasis probabilitas yang mengimplementasikan *Teorema Bayes* dengan asumsi independensi antar fitur (atribut). Model ini tergolong sederhana namun efektif, terutama pada data berjenis teks dan klasifikasi biner. Dalam konteks klasifikasi, *Naive Bayes* digunakan untuk memprediksi kemungkinan suatu objek termasuk dalam suatu kelas berdasarkan nilai fitur yang dimilikinya [18]. *Teorema Bayes* yang digunakan dalam algoritma ini dapat dirumuskan sebagai berikut:

$$P(C | X) = \frac{(P(X | C) \cdot P(C))}{P(X)} \quad (1)$$

dimana:

- $P(C | X)$: Probabilitas kelas C terhadap fitur X (*posterior probability*),
- $P(X | C)$: Probabilitas fitur X muncul dalam kelas C (*likelihood*),
- $P(C)$: Probabilitas awal dari kelas C (*prior probability*),
- $P(X)$: Probabilitas keseluruhan dari fitur X (*evidence*).

Evaluasi kinerja model bertujuan untuk mengukur sejauh mana model klasifikasi mampu menghasilkan prediksi yang benar. Evaluasi dilakukan menggunakan metrik-metrik statistik berdasarkan hasil prediksi dan data aktual, di antaranya:

Akurasi (Accuracy) adalah persentase jumlah prediksi yang benar dari total data:

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

Presisi (Precision) adalah kemampuan model dalam menghindari false positive:

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

Recall (Sensitivitas) adalah kemampuan model dalam menangkap semua data positif:

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

F1-Score adalah harmonik dari Precision dan Recall:

$$F1 = 2 \cdot \frac{Precision \cdot Recall}{Precision + Recall} \quad (5)$$

Area Under Curve (AUC) Digunakan untuk mengukur kinerja model klasifikasi biner berdasarkan *ROC Curve*.

Confusion Matrix Matriks yang menggambarkan hasil klasifikasi berdasarkan kelas sebenarnya dan prediksi [19].

3. METODE PENELITIAN

3.1. Jenis dan Pendekatan Penelitian

Penelitian ini bersifat kuantitatif deskriptif analitik, karena bertujuan mengukur dan menggambarkan pola sentimen pada data komentar dengan menggunakan metode klasifikasi statistik. Data yang dianalisis berupa teks dan hasilnya disajikan dalam bentuk frekuensi dan persentase kategori sentimen. Pendekatan ini sesuai untuk analisis opini publik melalui data media sosial karena bersifat menguji hipotesis kuantitatif secara sistematis.

3.2. Objek dan Lokasi Penelitian

Objek penelitian adalah komentar publik pada akun Instagram @tempodotco yang terkait dengan wacana penerapan kurikulum AI di Indonesia. Lokasi penelitian bersifat daring (online) sehingga lokasi fisik tidak relevan. Data diperoleh dari komentar yang tersedia pada postingan terkait di akun Instagram tersebut. Komentar-komentar tersebut menjadi populasi data untuk dianalisis sentimennya.

3.3. Pengumpulan Data

Sumber data berupa komentar teks publik pada media sosial Instagram (*posted comments*). Pengumpulan data dilakukan dengan teknik *web scraping* menggunakan bahasa pemrograman Python). Proses pengambilan data dilakukan pada bulan Juni 2025 untuk mendapatkan sampel yang representatif. Setiap komentar dikumpulkan beserta metadata pendukung (waktu, ID komentar). Sebelum diklasifikasikan, komentar diberi kode ID anonim agar kepatuhan privasi terjaga. Data dikumpulkan secara periodik dan difilter hanya komentar yang berhubungan dengan topik kurikulum AI.

3.4. Prapemrosesan Data

Data teks yang diperoleh belum siap untuk analisis. Oleh karena itu dilakukan tahapan prapemrosesan teks meliputi beberapa langkah umum (*case folding*, tokenisasi, penghilangan stopword, stemming, dan pelabelan manual). Prapemrosesan ini mengubah data mentah menjadi format terstruktur yang dapat dimasukkan ke model. Urutan langkahnya adalah:

- Case folding*: mengubah semua huruf menjadi huruf kecil untuk menyamakan format kata.
- Tokenisasi*: memecah kalimat komentar menjadi token (kata-kata) yang terpisah.
- Stopword removal*: menghapus kata-kata umum yang tidak memiliki makna sentimen (misalnya "dan", "yang", "di", dll.).
- Stemming*: mereduksi kata ke akar katanya, sehingga variasi kata (prefiks/sufiks) disamakan bentuk dasarnya. Untuk teks Bahasa Indonesia digunakan pustaka *Sastrawi* yang mampu mengubah kata berimbuhan menjadi bentuk dasarnya.
- Pelabelan*: setiap komentar diberi label sentiment (positif/negatif/netral) menggunakan metode VADER (Valence Aware Dictionary and sEntiment Reasoner) dengan menerjemahkan terlebih dahulu komentar berbahasa Indonesia ke dalam bahasa Inggris. Setiap komentar kemudian diberikan skor compound, di mana nilai > 0 dikategorikan sebagai sentimen positif, dan < 0 sebagai sentimen negatif.

3.5. Ekstraksi Fitur

Setelah praproses, setiap komentar diubah menjadi representasi numerik. Teknik yang digunakan adalah TF-IDF (*Term Frequency-Inverse Document*

Frequency). TF-IDF mengukur frekuensi suatu kata dalam satu dokumen dikalikan dengan bobot kebalikan frekuensi kata tersebut dalam seluruh dokumen, sehingga kata yang umum di seluruh korpus diberi bobot rendah. Metode TF-IDF banyak digunakan dalam ekstraksi fitur teks untuk pembelajaran mesin. Fitur TF-IDF bersifat vektorial dan sesuai untuk input ke model *Naïve Bayes*. Langkahnya: menghitung TF-IDF untuk setiap token pada semua komentar, menghasilkan matriks fitur dokumen.

3.6. Penerapan Algoritma Klasifikasi

Penelitian ini menguji tiga varian *Naïve Bayes*: *MultinomialNB*, *BernoulliNB*, dan *ComplementNB*. Pemilihan ketiganya didasarkan pada karakteristik data teks dan dokumentasi literatur:

- MultinomialNB* sering digunakan untuk klasifikasi teks berbasis hitungan (count/TF) dan cocok dengan fitur TF-IDF (frekuensi kata).
- BernoulliNB* menggunakan fitur biner (keberadaan/ketiadaan kata) sehingga dapat berguna bila perhatian hanya pada ada/tidaknya kata tertentu.
- ComplementNB* dikembangkan untuk mengatasi kelas data yang tidak seimbang dengan memperhitungkan distribusi komplementer kelas lain.

Naïve Bayes dipilih karena sederhana, cepat, dan terbukti efektif pada data teks besar. Beberapa studi menunjukkan *Naïve Bayes* memberikan akurasi tinggi dalam analisis sentimen Bahasa Indonesia, serta memiliki performa yang baik dalam pengolahan data berukuran besar. Ketiga varian akan diuji untuk menentukan performa terbaik melalui eksperimen. Model dilatih menggunakan pustaka *Scikit-learn* di Python, yang menyediakan implementasi *MultinomialNB*, *BernoulliNB*, dan *ComplementNB*.

3.7. Evaluasi Model

Kinerja model klasifikasi dievaluasi dengan beberapa metrik standar. Pertama, *Confusion Matrix* digunakan untuk melihat distribusi prediksi benar/salah pada setiap kelas sentimen. Kemudian dihitung metrik:

- Akurasi*: proporsi prediksi benar dari keseluruhan data.
- Presisi (precision)*: ketepatan prediksi positif (benar positif dibagi total positif hasil prediksi).
- Recall (sensitivitas)*: cakupan prediksi positif sejati (benar positif dibagi total sebenarnya positif).
- F1-score*: rata-rata harmonik antara presisi dan recall, berguna saat kelas tidak seimbang.

Penggunaan *presisi*, *recall*, dan *F1-score* diperlukan karena akurasi saja kurang informatif pada data tidak seimbang. Sebagai contoh, dalam studi evaluasi sentimen, *precision* dan *recall* diadopsi untuk mengukur performa lebih mendalam dibanding akurasi saja.

3.8. Alat dan Perangkat Lunak

Implementasi penelitian dilakukan dengan Python sebagai bahasa pemrograman utama. Analisis data dan model dibangun menggunakan pustaka seperti *Scikit-learn* (untuk algoritma klasifikasi dan fungsi evaluasi), *Pandas* (manajemen data), *NumPy* (perhitungan numerik), dan *NLTK* (pemisahan kata/token). Untuk Bahasa Indonesia, digunakan pustaka Sastrawi khusus stemming. Lingkungan pengembangan yang digunakan adalah Jupyter Notebook atau Google Colab yang mendukung eksekusi Python interaktif. Penggunaan Python dan Colab dipilih karena dukungan ekosistem yang lengkap untuk NLP dan kemudahan kolaborasi. Praproses teks juga dapat menggunakan *library* untuk penghapusan pola (misal emoji, tanda baca), serta perpustakaan lain untuk stopwords.

4. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan bahasa pemrograman Python dengan bantuan platform Google Colaboratory untuk membangun model klasifikasi sentimen terhadap komentar Instagram terkait penerapan Kurikulum AI. Model yang diimplementasikan meliputi *Multinomial Naïve Bayes*, *Bernoulli Naïve Bayes*, dan *Complement Naïve Bayes* dengan pendekatan TF-IDF dan n-gram.

4.1. Scraping Data

Data dikumpulkan dari akun Instagram @tempodotco yang secara aktif mempublikasikan konten tentang kebijakan nasional, termasuk topik terkait Kurikulum AI. Scraping dilakukan pada komentar dari beberapa unggahan terpilih yang relevan, menggunakan bantuan pustaka *Instaloader* untuk mengekstraksi komentar dalam format teks.

Komentar-komentar tersebut kemudian dikompilasi dalam dataset akhir sebanyak 1.238 data, yang terdiri dari atribut username, timestamp, dan komentar. Dataset ini menjadi dasar dalam proses preprocessing dan klasifikasi sentimen.

	username	text
0	chachasuhatasyah	@mamanracing201m SETUJU
1	rizal_muhammad_saja	Pemerintah semakin gak punya tujuan dalam semu...
2	andiarungg	Kocak
3	umimulla_s	Mending pelajaran tentang ahklak pak, biar ber...
4	fajrinadty	@fauzan_nurh class, bener bang 🙏

Gambar 1. Hasil scraping komentar dari akun Instagram @tempodotco

4.2. Preprocessing Data

Tahapan *preprocessing* bertujuan untuk membersihkan data mentah dari elemen-elemen yang tidak relevan, serta mengubah teks menjadi bentuk yang dapat diproses lebih lanjut oleh model pembelajaran mesin. Langkah awal yang dilakukan adalah *case folding*, yaitu mengubah semua huruf dalam komentar menjadi huruf kecil (*lowercase*). Hal

ini penting untuk menjaga konsistensi dalam representasi kata.

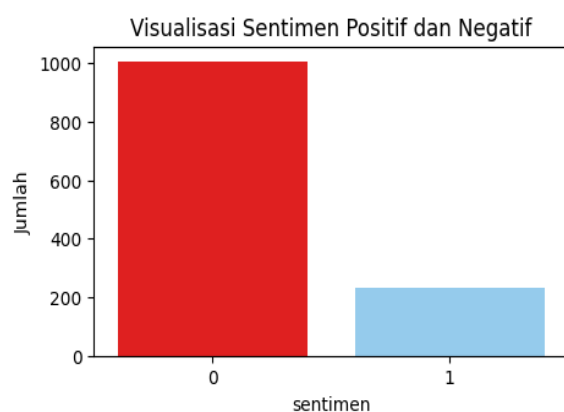
Selanjutnya dilakukan *cleansing*, yaitu menghapus simbol, tanda baca, emoji, hashtag, mention (@), URL, dan teks tidak relevan lainnya yang tidak memiliki kontribusi berarti terhadap analisis sentimen. Setelah itu, dilakukan penghapusan stopwords untuk menyaring kata-kata umum seperti "dan", "yang", "atau" yang tidak memberikan nilai informatif terhadap sentimen.

Proses dilanjutkan dengan *tokenizing*, yaitu memecah kalimat menjadi satuan kata (token). Setelah tokenisasi, dilakukan stemming menggunakan pustaka Sastrawi untuk mengubah kata menjadi bentuk dasarnya. Contohnya, kata "mengikuti" menjadi "ikut", "berjalan" menjadi "jalan". Terakhir dilakukan normalisasi kata untuk mengubah kata tidak baku atau gaul ke dalam bentuk baku. Proses ini penting untuk memastikan bahwa makna kata tidak hilang dan dapat dipahami oleh model.

Proses *preprocessing* dilakukan untuk membersihkan dan mempersiapkan data teks agar dapat diproses lebih lanjut oleh algoritma pembelajaran mesin. Langkah-langkah *preprocessing* yang diterapkan adalah sebagai berikut:

4.3. Labeling

Pelabelan sentimen dilakukan menggunakan metode VADER (*Valence Aware Dictionary and sEntiment Reasoner*) dengan menerjemahkan terlebih dahulu komentar berbahasa Indonesia ke dalam bahasa Inggris. Setiap komentar kemudian diberikan skor compound, di mana nilai > 0 dikategorikan sebagai sentimen positif, dan < 0 sebagai sentimen negatif. Hasil visualisasi distribusi label menunjukkan dominasi sentimen negatif sebesar 1006 data dan positif sebesar 232 data.



Gambar 2. Visualisasi distribusi label sentimen

4.4. Pembagian Data

Data dibagi ke dalam dua subset, yaitu data pelatihan dan data pengujian, dengan rasio 80:20. Model dilatih menggunakan data pelatihan dan diuji performanya terhadap data pengujian. Selain itu, dilakukan eksperimen menggunakan pembobotan fitur

TF-IDF serta bigram dan trigram n-gram untuk mengevaluasi dampaknya terhadap akurasi model.

4.5. Penerapan Model

Penerapan model dilakukan untuk mengevaluasi efektivitas tiga varian algoritma *Naïve Bayes*, yaitu *MultinomialNB*, *BernoulliNB*, dan *ComplementNB*, dengan pendekatan representasi data berbasis TF-IDF dan variasi n-gram (bigram dan trigram). Data komentar yang telah melewati proses preprocessing dijadikan input utama dalam pelatihan dan pengujian model. Evaluasi dilakukan menggunakan metrik akurasi, *precision*, *recall*, dan *f1-score* guna memperoleh gambaran kinerja model secara menyeluruh.

Secara umum, hasil implementasi menunjukkan bahwa model *MultinomialNB* dengan TF-IDF menghasilkan akurasi tertinggi yaitu 79.84%, disusul oleh *BernoulliNB* sebesar 77.02%, dan *ComplementNB* sebesar 66.13%. Sementara itu, penerapan n-gram pada *MultinomialNB* memperlihatkan bahwa akurasinya menurun menjadi 68.55% pada bigram, namun meningkat menjadi 71.37% pada trigram. Hal ini menunjukkan bahwa dalam konteks komentar singkat seperti di Instagram, penggunaan trigram mampu menangkap konteks yang lebih kaya dibandingkan bigram.

Tabel 1. Perbandingan hasil evaluasi tiga model *naïve bayes*

Metode Evaluasi	MNB	BNB	CNB
Akurasi (%)	79.84	77.02	66.13
<i>Precision</i> (Negatif)	0.82	0.78	0.83
<i>Recall</i> (Negatif)	0.94	0.97	0.68
<i>Precision</i> (Positif)	0.68	0.69	0.39
<i>Recall</i> (Positif)	0.40	0.17	0.60
F1-Score (Negatif)	0.87	0.86	0.75
F1-Score (Positif)	0.50	0.28	0.47
Macro Avg F1-Score	0.69	0.57	0.61
Weighted Avg F1-Score	0.78	0.71	0.68

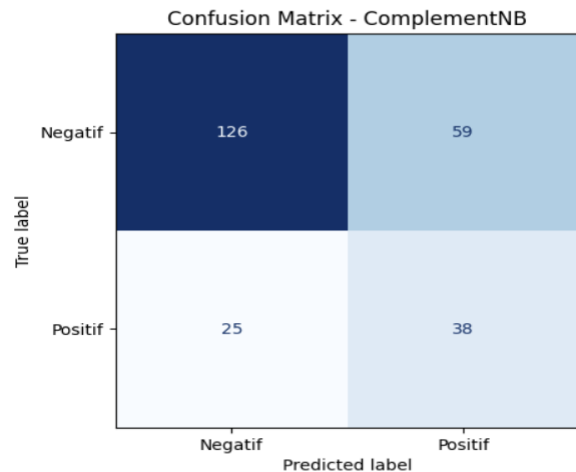
Hasil evaluasi menunjukkan bahwa:

- MultinomialNB* sangat baik dalam mendeteksi komentar negatif dengan *precision* sebesar 0.82 dan *recall* 0.94, meskipun performa pada komentar positif masih rendah.
- BernoulliNB* memiliki *recall* tertinggi untuk komentar negatif (0.97), namun performa klasifikasi pada komentar positif sangat rendah (*recall* hanya 0.17), mengindikasikan ketidakseimbangan dalam prediksi kelas.
- ComplementNB* menunjukkan performa paling seimbang dalam hal *recall* antara kelas positif dan negatif, namun memiliki *precision* yang rendah terutama pada komentar positif, sehingga menghasilkan *f1-score* yang kurang optimal.

4.6. Evaluasi Hasil

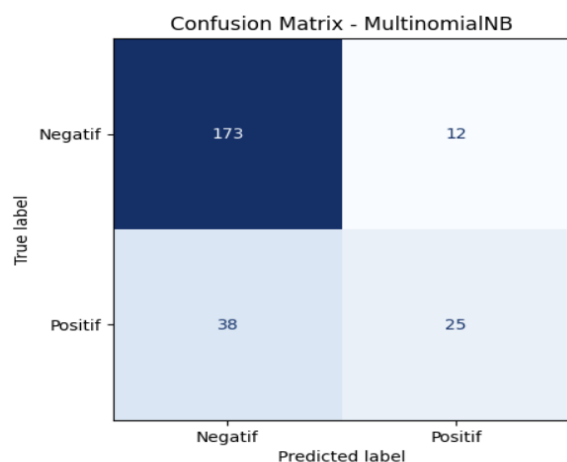
Evaluasi model dilakukan dengan menggunakan *Confusion Matrix* dan metrik evaluasi seperti akurasi, *presisi*, *recall*, dan F1-score. *Confusion Matrix*

memberikan gambaran distribusi hasil klasifikasi model terhadap kelas positif dan negatif. Misalnya, *True Positive* (TP) menunjukkan jumlah data positif yang berhasil diklasifikasi dengan benar, sedangkan *False Positive* (FP) adalah data negatif yang salah diklasifikasi sebagai positif.



Gambar 3. *Confusion matrix model complement naïve bayes*

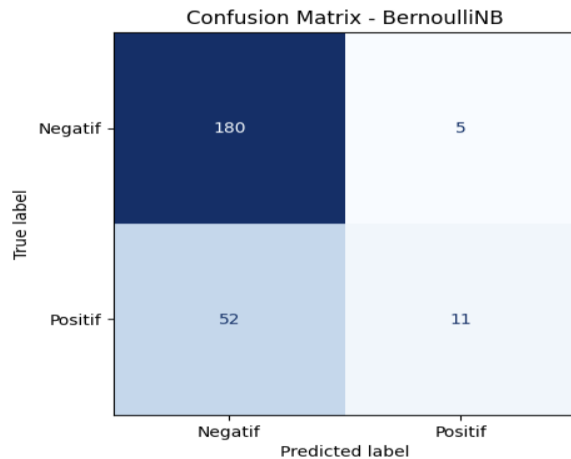
Gambar ini menunjukkan hasil prediksi model *Complement Naïve Bayes* terhadap data uji. Dari total data, model berhasil mengklasifikasikan 126 komentar negatif dan 38 komentar positif secara benar. Namun, masih terdapat kesalahan klasifikasi yaitu 59 komentar negatif diklasifikasikan sebagai positif (*False Positive*) dan 25 komentar positif diklasifikasikan sebagai negatif (*False Negative*). Ini menunjukkan bahwa model *ComplementNB* cenderung lebih baik dalam mengenali komentar negatif dibandingkan komentar positif.



Gambar 4. *Confusion matrix model multinomial naïve bayes*

Confusion matrix untuk model *MultinomialNB* menunjukkan performa yang cukup baik. Model berhasil mengklasifikasikan 173 komentar negatif dan 25 komentar positif secara benar. Hanya terdapat 12

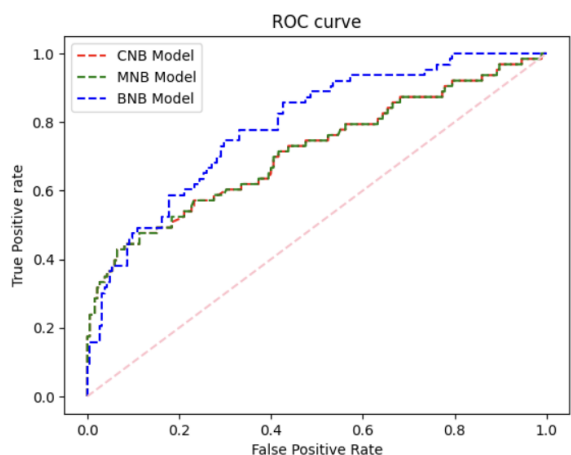
komentar negatif yang salah diklasifikasikan sebagai positif, serta 38 komentar positif yang salah diklasifikasikan sebagai negatif. Meskipun model ini memiliki *True Negative* yang tinggi, performanya dalam mengidentifikasi komentar positif masih dapat ditingkatkan.



Gambar 5. *Confusion matrix model bernolli naïve bayes*

Model *BernoulliNB* memiliki *True Negative* tertinggi di antara ketiga model, dengan 180 komentar negatif yang diklasifikasikan dengan benar. Namun, model ini hanya berhasil mengklasifikasikan 11 komentar positif secara akurat, sementara 52 lainnya salah diklasifikasikan sebagai negatif. Hal ini menunjukkan bahwa model *BernoulliNB* sangat konservatif dalam mendeteksi komentar positif, dan lebih dominan memprediksi negatif.

Selain itu, dilakukan pula penggambaran *Receiver Operating Characteristic (ROC) Curve* untuk menilai performa klasifikasi binary. *ROC Curve* memvisualisasikan *trade-off* antara *True Positive Rate (TPR)* dan *False Positive Rate (FPR)*. Luas area di bawah kurva (AUC) memberikan indikator seberapa baik model mampu membedakan antara kelas.



Gambar 6. ROC *curve* model

Kurva ROC menampilkan kinerja ketiga model berdasarkan hubungan antara *True Positive Rate* (TPR) dan *False Positive Rate* (FPR). Model *BernoulliNB* (garis biru) menunjukkan kurva paling mendekati sudut kiri atas grafik, mengindikasikan performa terbaik dalam membedakan kelas positif dan negatif. Sementara itu, *MultinomialNB* dan *ComplementNB* memiliki kurva yang mirip dan berada di bawah *BernoulliNB*, yang berarti performanya sedikit lebih rendah. Semakin tinggi *Area Under Curve* (AUC), semakin baik model dalam klasifikasi biner.

4.7. Word Cloud

Untuk memberikan wawasan visual mengenai kata-kata yang paling sering muncul dalam komentar, dibuat visualisasi *Word Cloud*. Kata-kata yang berukuran lebih besar menunjukkan frekuensi kemunculan yang lebih tinggi. Visualisasi ini membantu dalam mengidentifikasi kata-kata yang sering diasosiasikan dengan sentimen tertentu, seperti "AI", "kurikulum", "anak", "takut", "butuh", dan lainnya.



Gambar 7. *Word cloud* instagram komentar positif

Visualisasi ini menampilkan frekuensi kata yang sering muncul pada komentar berlabel positif. Kata-kata seperti "ajar", "anak", "didik", "baik", dan "bagus" muncul dominan. Hal ini menunjukkan bahwa sentimen positif banyak mengarah pada harapan akan peningkatan kualitas pendidikan, ajakan untuk berpikir cerdas, serta apresiasi terhadap inovasi kurikulum.



Gambar 8. Word cloud instagram komentar negatif

Word cloud ini memvisualisasikan kata-kata yang sering digunakan dalam komentar berlabel negatif. Kata "ai", "pakai", "kurikulum", "anak", dan "ganti" muncul dengan ukuran besar, menandakan seringnya digunakan. Ini mencerminkan kekhawatiran dan resistensi publik terhadap penggunaan AI dalam kurikulum, ketakutan akan dampaknya terhadap anak-anak, serta kritik terhadap pembuatan kebijakan.

Secara keseluruhan, penerapan *Multinomial Naïve Bayes* dengan kombinasi fitur TF-IDF dan n-gram terbukti efektif untuk analisis sentimen komentar media sosial berbahasa Indonesia, dengan akurasi terbaik mencapai 79,84% pada konfigurasi tertentu. Model ini dapat dijadikan acuan dalam pengembangan sistem pemantauan opini publik terkait kebijakan pendidikan di Indonesia, khususnya kurikulum kecerdasan buatan.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian ini, dapat disimpulkan bahwa model *Multinomial Naïve Bayes* (*MultinomialNB*) dengan fitur TF-IDF serta kombinasi n-gram (khususnya bigram dan trigram) terbukti paling efektif dalam mengklasifikasikan sentimen komentar Instagram terkait Kurikulum AI, dengan akurasi tertinggi mencapai 79,84%. Penggunaan teknik ini memperlihatkan peningkatan kinerja signifikan dibandingkan dengan pendekatan lainnya seperti *BernoulliNB* dan *ComplementNB*, serta menegaskan pentingnya pemilihan fitur dan pra-pemrosesan yang sesuai dalam analisis teks berbahasa Indonesia. Penelitian ini juga memberikan kontribusi praktis dengan menyediakan wawasan mengenai persepsi publik yang dapat dimanfaatkan oleh Kemendikbudristek, pengembang kurikulum, dan tenaga pendidik dalam merancang kebijakan dan

strategi implementasi Kurikulum AI yang lebih adaptif dan komunikatif. Untuk penelitian selanjutnya, disarankan agar dilakukan eksplorasi model lain seperti SVM atau *Deep Learning* untuk perbandingan performa, serta memperluas data ke platform media sosial lain guna memperoleh gambaran sentimen publik yang lebih menyeluruh.

DAFTAR PUSTAKA

- [1] R. Kurniawaty, "Mempersiapkan Generasi Masa Depan: Pentingnya Digitalisasi Pembelajaran, Coding, dan Artificial Intelligence di Sekolah," BPMP Jakarta – LPMP DKI Jakarta. Accessed: Jun. 29, 2025. [Online]. Available: <https://lpmpdki.kemdikbud.go.id/mempersiapkan-generasi-masa-depan-pentingnya-digitalisasi-pembelajaran-coding-dan-artificial-intelligence-di-sekolah/>
- [2] I. H. Y. Yim and J. Su, "Artificial intelligence literacy education in primary schools: a review," *Int J Technol Des Educ*, pp. 1–30, Apr. 2025, doi: 10.1007/S10798-025-09979-W/TABLES/8.
- [3] K. Safitri and D. Damarjati, "Tahun Depan Kelas 5 SD Bakal Mulai Belajar 'Coding' dan AI," Kompas.com. Accessed: Jun. 29, 2025. [Online]. Available: <https://nasional.kompas.com/read/2025/04/29/17303501/tahun-depan-kelas-5-sd-bakal-mulai-belajar-coding-dan-ai>
- [4] Brilio, "Wapres Gibran sebut kurikulum AI diterapkan di sekolah mulai tahun ajaran baru SD-SMA," Brilio.net. Accessed: Jun. 29, 2025. [Online]. Available: <https://www.brilio.net/serius/gibran-sebut-kurikulum-ai-diterapkan-di-sekolah-mulai-tahun-ajaran-baru-sd-sma-250503a.html>
- [5] C. Yulianti, "Pakar Sebut Mahasiswa Perlu Dilibatkan Ngajar Coding & AI di Sekolah, Kenapa?," Detik (DetikEdu). Accessed: Jun. 29, 2025. [Online]. Available: <https://www.detik.com/edu/sekolah/d-7777432/pakar-sebut-mahasiswa-perlu-dilibatkan-ngajar-coding-ai-di-sekolah-kenapa>
- [6] C. M. Annur, "Indonesia Masuk 5 Besar Negara dengan Pengguna Instagram Terbanyak di Dunia," Databoks – Katadata Insight Center. Accessed: Jun. 29, 2025. [Online]. Available: <https://databoks.katadata.co.id/media/statistik/e775c90d0a853ab/indonesia-masuk-5-besar-negara-dengan-pengguna-instagram-terbanyak-di-dunia>
- [7] K. Alahmadi, S. Alharbi, J. Chen, and X. Wang, "Generalizing sentiment analysis: a review of progress, challenges, and emerging directions," *Soc Netw Anal Min*, vol. 15, no. 1, pp. 1–28, Dec. 2025, doi: 10.1007/S13278-025-01461-8/FIGURES/2.
- [8] K. Khairunnisa, S. K. Dewi, D. D. Rahmawati, and A. P. Sari, "ANALISIS SENTIMEN KOMENTAR PADA POSTINGAN

- INSTAGRAM AKUN ‘STANDWITHUS’ MENGGUNAKAN KLASIFIKASI NAIVE BAYES,” *JURNAL ILMIAH INFORMATIKA*, vol. 12, no. 02, pp. 191–199, Sep. 2024, doi: 10.33884/JIF.V12I02.9263.
- [9] J. Homepage, D. Pramudita, Y. Akbar, and T. Wahyudi, “Analisis Sentimen Terhadap Program Kartu Indonesia Pintar Kuliah pada Media Sosial X Menggunakan Algoritma Naive Bayes,” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 4, pp. 1420–1430, Aug. 2024, doi: 10.57152/MALCOM.V4I4.1565.
- [10] D. S. Arum, S. Butsianto, R. Astuti, and U. Pelita Bangsa, “ANALISIS SENTIMEN MASYARAKAT INDONESIA TERHADAP SEA GAMES 2023 DI TWITTER DENGAN METODE NAÏVE BAYES,” *Journal of Information System, Applied, Management, Accounting and Research*, vol. 7, no. 3, pp. 728–738, Jul. 2023, doi: 10.52362/JISAMAR.V7I3.1150.
- [11] M. A. Prihardana and M. T. Lestari, “Analysis of the Utilization of Social Media Monitoring @Roseallday.Co Instagram to Prevent Cancel Culture Crisis,” *Daengku: Journal of Humanities and Social Sciences Innovation*, vol. 4, no. 6, pp. 1089–1094, Dec. 2024, doi: 10.35877/454RI.daengku2965.
- [12] P. Pandunata, C. K. Ananta, and Y. Nurdiansyah, “Analisis Sentimen Opini Publik Terhadap Pekan Olahraga Nasional Pada Instagram Menggunakan Metode Naive Bayes Classifier,” *INFORMAL: Informatics Journal*, vol. 7, no. 2, p. 146, Aug. 2022, doi: 10.19184/ISJ.V7I2.33928.
- [13] Y. Akbar and T. Sugiharto, “Analisis Sentimen Pengguna Twitter di Indonesia Terhadap ChatGPT Menggunakan Algoritma C4.5 dan Naive Bayes,” *Jurnal Sains dan Teknologi*, vol. 5, pp. 155–122, Aug. 2023, Accessed: Jun. 30, 2025. [Online]. Available: <https://ejournal.sisfokomtek.org/index.php/saintek/article/view/1368>
- [14] A. R. Fahky and S. Zulaikha, “Indonesia upayakan langkah kolaborasi pengembangan AI,” *ANTARA News*. Accessed: Jun. 29, 2025. [Online]. Available: <https://www.antaranews.com/berita/4924217/indonesia-upayakan-langkah-kolaborasi-pengembangan-ai>
- [15] NapoleonCat, “Social Media Demographics by Country – Monthly Updates,” NapoleonCat (Social Media Statistics). Accessed: Jun. 29, 2025. [Online]. Available: <https://napoleontcat.com/stats/instagram-users-in-indonesia/2024/01/>
- [16] P. S. Ghatora, S. E. Hosseini, S. Pervez, M. J. Iqbal, and N. Shaukat, “Sentiment Analysis of Product Reviews Using Machine Learning and Pre-Trained LLM,” *Big Data and Cognitive Computing 2024, Vol. 8, Page 199*, vol. 8, no. 12, p. 199, Dec. 2024, doi: 10.3390/BDCC8120199.
- [17] P. S. Ghatora, S. E. Hosseini, S. Pervez, M. J. Iqbal, and N. Shaukat, “Sentiment Analysis of Product Reviews Using Machine Learning and Pre-Trained LLM,” *Big Data and Cognitive Computing 2024, Vol. 8, Page 199*, vol. 8, no. 12, p. 199, Dec. 2024, doi: 10.3390/BDCC8120199.
- [18] A. N. Zhafira *et al.*, “Analisis Sentimen Ulasan Film pada Dataset IMDB menggunakan Algoritma Naive Bayes,” *Jurnal Manajemen Informatika, Sistem Informasi dan Teknologi Komputer (JUMISTIK)*, vol. 4, no. 1, pp. 373–383, Jun. 2025, doi: 10.70247/JUMISTIK.V4I1.139.
- [19] C. Hayati, S. Shofiah Hilabi, A. Lia Hananto, S. Informasi, and U. Buana Perjuangan Karawang, “KLASIFIKASI DAN PREDIKSI ULASAN APLIKASI DANA PADA GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA NAÏVE BAYES,” *Jurnal Informatika Teknologi dan Sains (Jinteks)*, vol. 7, no. 2, pp. 596–605, May 2025, doi: 10.51401/JINTEKS.V7I2.5691