



Penerapan Naïve Bayes untuk Analisis Sentimen Publik terhadap Isu Korupsi Pertamina di Media Sosial

Nadeya Ladia Tantri

Informatika, Universitas Sebelas April, Indonesia

*Penulis Korespondensi: nadeyatantri30@gmail.com

Abstract. Corruption remains a major public issue in Indonesia, including the Pertamina corruption case, which has triggered diverse reactions across social media. This study aims to analyze public sentiment toward the Pertamina corruption issue using the Naïve Bayes algorithm as a text classification method. The research process includes data collection, text preprocessing consisting of slang normalization, stopword removal, and stemming, term weighting using Term Frequency-Inverse Document Frequency (TF-IDF), sentiment classification using Naïve Bayes, and model evaluation using a confusion matrix. The dataset used in this study consists of public comments in Indonesian that were processed and classified into negative, neutral, and positive sentiment categories. The evaluation results show that the model achieved an accuracy of 0.7592 with a macro F1-score of 0.4642. The negative sentiment class produced the best performance compared with the neutral and positive classes. Dominant negative terms such as corruption, law, corruptor, and death indicate that public opinion tends to respond negatively to the case. These findings provide a clearer picture of public perception regarding the Pertamina corruption issue and may serve as a reference for stakeholders in understanding public opinion expressed on social media.

Keywords: Corruption; Naïve Bayes; Sentiment Analysis; Social Media; TF-IDF.

Abstrak. Korupsi merupakan salah satu isu penting yang mendapat perhatian luas dari masyarakat Indonesia, termasuk kasus korupsi Pertamina yang memunculkan beragam respons di media sosial. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap isu korupsi Pertamina menggunakan algoritma Naïve Bayes sebagai metode klasifikasi teks. Tahapan penelitian meliputi pengumpulan data, preprocessing teks yang terdiri atas normalisasi slang, penghapusan stopword, dan stemming, pembobotan kata menggunakan Term Frequency-Inverse Document Frequency (TF-IDF), klasifikasi sentimen menggunakan Naïve Bayes, serta evaluasi model menggunakan confusion matrix. Dataset yang digunakan berisi komentar publik berbahasa Indonesia yang telah diproses dan diklasifikasikan ke dalam sentimen negatif, netral, dan positif. Hasil evaluasi menunjukkan bahwa model menghasilkan akurasi sebesar 0,7592 dengan macro F1-score sebesar 0,4642. Kelas sentimen negatif menunjukkan performa terbaik dibandingkan kelas netral dan positif. Kata-kata dominan pada sentimen negatif seperti korupsi, hukum, koruptor, dan mati menunjukkan bahwa opini publik cenderung memberikan respons negatif terhadap kasus tersebut. Hasil penelitian ini diharapkan dapat memberikan gambaran yang lebih jelas mengenai persepsi masyarakat terhadap isu korupsi Pertamina serta menjadi referensi bagi pihak terkait dalam memahami opini publik di media sosial.

Kata Kunci: Analisis Sentimen; Korupsi; Media Sosial; Naïve Bayes; TF-IDF.

1. LATAR BELAKANG

Korupsi masih menjadi salah satu persoalan serius di Indonesia karena berdampak pada menurunnya kepercayaan publik, terganggunya tata kelola pemerintahan, dan melemahnya legitimasi institusi negara. Ketika kasus korupsi melibatkan perusahaan strategis milik negara seperti Pertamina, persoalannya tidak hanya berkaitan dengan aspek hukum, tetapi juga memicu respons publik yang luas di media sosial. Dalam konteks ini, media sosial menjadi ruang utama bagi masyarakat untuk menyampaikan kritik, kecaman, dukungan terhadap penegakan hukum, maupun komentar netral terhadap isu korupsi yang sedang berkembang (Astuti, 2025; Pebrianti, 2025). Kasus korupsi Pertamina yang mencuat pada awal 2025 menarik perhatian besar publik, khususnya di platform X. Penelitian terbaru menunjukkan

bahwa percakapan publik mengenai kasus ini didominasi sentimen negatif, yang menandakan tingginya tingkat kekecewaan dan kekhawatiran masyarakat terhadap integritas perusahaan negara tersebut. Temuan serupa juga ditunjukkan oleh studi lain yang menganalisis kasus korupsi Pertamina menggunakan pendekatan berbasis sentimen, sehingga memperkuat bahwa isu ini memang memunculkan gelombang opini publik yang kuat di ruang digital

Perkembangan media sosial menghasilkan data teks dalam jumlah besar yang tidak terstruktur. Data ini dapat dimanfaatkan untuk memahami persepsi publik secara lebih sistematis melalui pendekatan analisis sentimen. Analisis sentimen adalah salah satu metode dalam text mining dan NLP yang bertujuan mengklasifikasikan opini ke dalam kategori positif, negatif, atau netral. Pendekatan ini relevan untuk isu korupsi Pertamina karena komentar pengguna media sosial umumnya singkat, emosional, dan mengandung variasi bahasa informal yang tinggi (Katya & Rahman, 2024).

Agar data teks dapat diolah dengan baik, diperlukan tahap preprocessing. Tahapan ini umumnya mencakup pembersihan teks, tokenisasi, normalisasi kata tidak baku, penghapusan stopword, dan stemming. Dalam konteks bahasa Indonesia, preprocessing menjadi sangat penting karena komentar media sosial sering menggunakan singkatan, kata gaul, dan bentuk penulisan yang beragam. Preprocessing yang tepat membantu mengurangi noise dan meningkatkan kualitas fitur yang akan digunakan dalam proses klasifikasi (Ardi & Kurniawan, 2024; Lestari & Hutagalung, 2025).

Setelah preprocessing, teks perlu diubah menjadi representasi numerik. Salah satu metode pembobotan yang umum digunakan adalah TF-IDF. Metode ini memberikan bobot lebih tinggi pada kata yang penting dalam suatu dokumen namun jarang muncul di dokumen lain, sehingga cocok digunakan untuk data komentar media sosial yang cenderung pendek. Sejumlah penelitian terbaru menunjukkan bahwa TF-IDF efektif digunakan dalam analisis sentimen bahasa Indonesia, baik dikombinasikan dengan Naïve Bayes maupun algoritma klasifikasi lain (Lestari & Hutagalung, 2025).

Dalam tahap klasifikasi, Naïve Bayes menjadi salah satu algoritma yang banyak digunakan karena sederhana, cepat, dan cukup efektif untuk data teks berdimensi tinggi. Penelitian terbaru menunjukkan bahwa Naïve Bayes tetap relevan untuk analisis sentimen, meskipun pada beberapa kasus masih menghadapi tantangan dalam membedakan sentimen netral dan kategori lainnya. Karena itu, algoritma ini sesuai digunakan untuk mengklasifikasikan komentar publik terhadap kasus korupsi Pertamina yang memiliki pola opini beragam dan konteks emosional yang kuat (Ardi & Kurniawan, 2024).

Berdasarkan uraian tersebut, penelitian ini difokuskan untuk menganalisis sentimen publik terhadap kasus korupsi Pertamina menggunakan algoritma Naïve Bayes dengan tahapan preprocessing, pembobotan TF-IDF, klasifikasi, dan evaluasi model menggunakan confusion matrix. Penelitian ini diharapkan dapat memberikan gambaran yang lebih sistematis mengenai kecenderungan opini masyarakat di media sosial serta menjadi masukan dalam memahami respons publik terhadap isu korupsi di perusahaan strategis milik negara.

2. KAJIAN TEORITIS

Analisis Sentimen

Analisis sentimen merupakan cabang dari text mining dan Natural Language Processing (NLP) yang bertujuan mengidentifikasi serta mengklasifikasikan polaritas opini dalam teks ke dalam kategori positif, negatif, dan netral. Dalam konteks media sosial, analisis sentimen digunakan untuk membaca kecenderungan opini publik secara otomatis dari data teks yang bersifat besar, cepat berubah, dan tidak terstruktur. Pendekatan ini relevan untuk penelitian ini karena media sosial menjadi ruang utama masyarakat dalam mengekspresikan pandangan terhadap isu korupsi Pertamina, sehingga data komentar publik dapat dimanfaatkan untuk memetakan persepsi masyarakat secara kuantitatif (Muhaimin et al., 2024; Ramadhan et al., 2025).

Perkembangan penelitian terkini menunjukkan bahwa analisis sentimen masih menjadi pendekatan yang efektif untuk memahami opini publik pada berbagai isu sosial di media digital. Studi tentang percakapan di media sosial menunjukkan bahwa teks yang dihasilkan pengguna umumnya sarat emosi, opini, dan penilaian subjektif, sehingga perlu diolah dengan metode yang tepat agar dapat menghasilkan informasi yang bermakna. Dalam penelitian ini, analisis sentimen digunakan untuk mengidentifikasi arah opini masyarakat terhadap kasus korupsi Pertamina dan melihat kecenderungan sentimen yang muncul dalam komentar publik (Rodríguez-Ibáñez et al., 2023).

Preprocessing Teks

Preprocessing teks merupakan tahap penting dalam analisis sentimen karena data media sosial umumnya mengandung noise, seperti tanda baca, URL, simbol, singkatan, kata tidak baku, dan emotikon. Tahapan preprocessing lazim meliputi cleansing, case folding, tokenization, stopword removal, dan stemming. Proses ini bertujuan menstandarkan teks agar lebih siap diproses oleh algoritma klasifikasi serta mengurangi gangguan yang dapat menurunkan kualitas hasil analisis (W et al., 2025)(Hartimar et al., 2025; Lestari & Hutagalung, 2025).

Pada teks berbahasa Indonesia, preprocessing menjadi semakin penting karena penggunaan bahasa informal, slang, dan variasi penulisan kata cukup dominan dalam komentar media sosial. Penelitian terdahulu menunjukkan bahwa normalisasi kata dan stemming dapat membantu model mengenali pola kata yang serupa sehingga representasi fitur menjadi lebih konsisten. Oleh karena itu, preprocessing tidak hanya berfungsi membersihkan data, tetapi juga meningkatkan kualitas informasi yang akan diekstraksi dari teks.

TF-IDF

TF-IDF atau Term Frequency–Inverse Document Frequency merupakan metode pembobotan kata yang digunakan untuk mengubah teks menjadi bentuk numerik berdasarkan tingkat kepentingan suatu kata dalam dokumen dibandingkan dengan keseluruhan korpus. Kata yang sering muncul dalam satu dokumen tetapi jarang muncul pada dokumen lain akan memperoleh bobot lebih tinggi. Teknik ini banyak digunakan dalam analisis sentimen karena efektif untuk menonjolkan kata-kata informatif, terutama pada teks pendek seperti komentar media sosial (Khoerunnisa et al., 2025).

Dalam penelitian ini, TF-IDF digunakan sebagai metode ekstraksi fitur sebelum klasifikasi dengan Naïve Bayes. Pemilihan TF-IDF didasarkan pada kemampuannya menghasilkan representasi teks yang ringkas namun tetap informatif, sehingga sesuai untuk karakteristik data komentar publik yang cenderung pendek dan variatif. Sejumlah studi terbaru juga menunjukkan bahwa TF-IDF masih memberikan performa yang baik dalam berbagai kasus analisis sentimen bahasa Indonesia.

Naïve Bayes

Naïve Bayes merupakan algoritma klasifikasi berbasis probabilitas yang mengasumsikan bahwa setiap fitur bersifat independen terhadap fitur lainnya dalam satu kelas. Walaupun asumsi ini sederhana, Naïve Bayes dikenal cepat, efisien, dan cukup efektif untuk klasifikasi teks, khususnya ketika dipadukan dengan pembobotan TF-IDF. Dalam analisis sentimen, algoritma ini sering digunakan karena mampu menangani data berdimensi tinggi dan cocok untuk teks pendek yang mengandung kata-kata penting (Hartimar et al., 2025; Khoerunnisa et al., 2025).

Pada penelitian ini digunakan Multinomial Naïve Bayes karena sesuai untuk data teks yang direpresentasikan dalam bentuk frekuensi atau bobot kata. Beberapa penelitian sebelumnya membuktikan bahwa Naïve Bayes mampu memberikan hasil yang kompetitif dalam klasifikasi sentimen bahasa Indonesia, meskipun pada kelas yang lebih ambigu seperti netral, performanya sering kali lebih rendah dibandingkan kelas yang polaritasnya jelas (Abdullah Wahid & Darma Andayani, 2025).

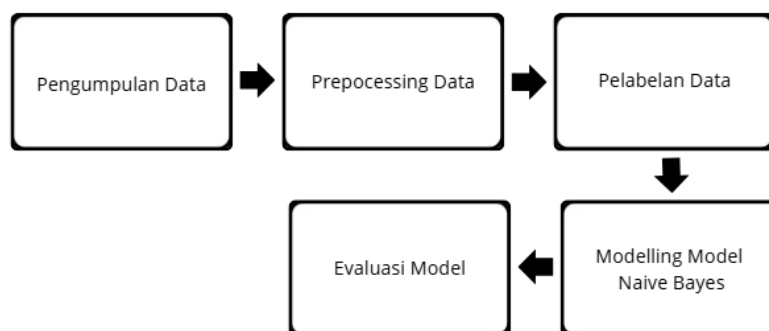
Penelitian Terdahulu

Sejumlah penelitian dalam lima tahun terakhir menunjukkan bahwa kombinasi TF-IDF dan Naïve Bayes efektif digunakan dalam analisis sentimen bahasa Indonesia. Studi pada ulasan aplikasi dan layanan digital membuktikan bahwa pendekatan ini mampu mengklasifikasikan opini pengguna ke dalam kelas positif, negatif, dan netral dengan performa yang cukup baik. Selain itu, penelitian pada isu-isu yang ramai di media sosial juga menunjukkan bahwa data komentar publik cenderung didominasi sentimen tertentu, sehingga cocok dianalisis menggunakan pendekatan klasifikasi teks (Hartimar et al., 2025; Muhaimin et al., 2024; Rasmila et al., 2025).

Pada topik yang berkaitan dengan Pertamina, penelitian terbaru menunjukkan bahwa kasus ini memunculkan respons publik yang kuat dan cenderung negatif di media sosial. Studi tentang kasus korupsi Pertamina memperlihatkan bahwa isu tersebut memicu perhatian luas dan menjadi objek analisis sentimen yang relevan untuk memahami persepsi masyarakat. Dengan demikian, penelitian ini memiliki dasar akademik yang kuat karena menggabungkan konteks isu publik yang aktual dengan metode analisis teks yang telah terbukti efektif.

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode klasifikasi teks untuk menganalisis sentimen publik terhadap isu korupsi Pertamina. Tahapan penelitian meliputi pengumpulan data, preprocessing teks, ekstraksi fitur menggunakan TF-IDF, klasifikasi dengan Multinomial Naïve Bayes, serta evaluasi model menggunakan confusion matrix dan metrik evaluasi klasifikasi (Arya et al., 2026).



Gambar 1. Alur Penelitian.

Alur penelitian dimulai dari pengumpulan dataset komentar media sosial yang telah diberi label sentimen, kemudian dilakukan preprocessing untuk membersihkan dan menstandarkan teks. Setelah itu, data diubah menjadi representasi numerik menggunakan TF-

IDF agar dapat diproses oleh algoritma klasifikasi. Tahap berikutnya adalah pelatihan model Multinomial Naïve Bayes, lalu evaluasi dilakukan untuk menilai performa model dalam mengklasifikasikan sentimen komentar ke dalam kelas positif, negatif, dan netral (Ardi & Kurniawan, 2024).

Pengumpulan Data

Data penelitian diperoleh dari dataset publik berformat CSV berjudul *dataset-sentimen-korupsi-pertamina.csv* yang berisi komentar pengguna media sosial terkait isu korupsi Pertamina. Penggunaan dataset publik dipilih karena topiknya sesuai dengan fokus penelitian, tersedia dalam format terstruktur, dan telah memiliki label sentimen yang dapat digunakan sebagai dasar proses klasifikasi. Setelah proses pembersihan data, jumlah data yang digunakan dalam pemodelan sebanyak 952 data (Arya et al., 2026).

Preprocessing

Tahap preprocessing dilakukan untuk mengurangi noise pada teks media sosial dan menyelaraskan bentuk kata agar lebih mudah diproses oleh model. Pada penelitian ini, preprocessing meliputi *cleaning*, *case folding*, normalisasi slang, *stopword removal*, dan stemming. Proses *cleaning* dilakukan untuk menghapus URL, tanda baca, simbol, angka, dan karakter yang tidak diperlukan, sedangkan *case folding* digunakan untuk menyeragamkan teks menjadi huruf kecil (Lestari & Hutagalung, 2025).

Normalisasi slang dilakukan untuk mengubah kata tidak baku atau singkatan ke bentuk baku agar variasi penulisan tidak mengganggu proses klasifikasi. Selanjutnya, *stopword removal* menghapus kata umum yang tidak membawa informasi sentimen, dan stemming diterapkan untuk mengubah kata berimbuhan menjadi bentuk dasar. Tahapan ini penting karena komentar media sosial berbahasa Indonesia cenderung informal, variatif, dan mengandung banyak bentuk kata yang tidak baku (Ardi & Kurniawan, 2024).

TF-IDF

Setelah preprocessing, setiap komentar diubah menjadi bentuk numerik menggunakan TF-IDF. Metode ini memberikan bobot lebih tinggi pada kata yang sering muncul dalam satu dokumen tetapi relatif jarang muncul dalam dokumen lain, sehingga kata yang bersifat informatif lebih menonjol dalam proses klasifikasi. TF-IDF dipilih karena efektif untuk data teks pendek seperti komentar media sosial dan telah banyak digunakan dalam penelitian analisis sentimen berbahasa Indonesia (Khoerunnisa et al., 2025).

Pelabelan Data

Dataset yang digunakan dalam penelitian ini telah memiliki label sentimen sehingga dapat langsung dimanfaatkan sebagai *ground truth*. Kelas sentimen dibagi menjadi tiga

kategori, yaitu positif, negatif, dan netral. Kelas negatif merepresentasikan kritik atau ketidaksetujuan terhadap isu korupsi Pertamina, kelas positif merepresentasikan dukungan atau respons yang bersifat apresiatif, sedangkan kelas netral menunjukkan komentar yang informatif atau tidak memuat kecenderungan emosi yang kuat. Ketersediaan label awal memudahkan proses pembelajaran model dan evaluasi hasil klasifikasi (Arya et al., 2026).

Klasifikasi Model

Tahap klasifikasi merupakan inti dari penelitian ini, yaitu mengelompokkan komentar ke dalam kategori sentimen menggunakan algoritma Naïve Bayes. Algoritma ini dipilih karena sederhana, cepat, dan cukup efektif untuk menangani data teks dengan fitur yang besar seperti hasil TF-IDF.

Naïve Bayes bekerja dengan menghitung probabilitas kemunculan fitur kata terhadap masing-masing kelas sentimen. Dalam penelitian ini digunakan Multinomial Naïve Bayes, yang umum dipakai untuk klasifikasi teks berbasis frekuensi kata. Parameter smoothing alpha dioptimasi menggunakan GridSearchCV, dan nilai terbaik yang diperoleh adalah alpha 0.1.

Data dibagi menjadi 80% data latih dan 20% data uji secara stratified untuk menjaga proporsi kelas pada masing-masing subset data. Skema ini sesuai dengan praktik umum dalam penelitian klasifikasi sentimen berbasis teks. Model kemudian dilatih menggunakan data yang telah melalui preprocessing dan vectorization, lalu diuji pada data yang belum pernah dilihat sebelumnya untuk mengukur kemampuan generalisasi model (Ardi & Kurniawan, 2024).

Evaluasi

Evaluasi model dilakukan menggunakan confusion matrix serta metrik klasifikasi, yaitu accuracy, precision, recall, dan F1-score. Confusion matrix digunakan untuk melihat kesesuaian antara label aktual dan label prediksi, sedangkan precision, recall, dan F1-score digunakan untuk menilai performa model pada masing-masing kelas sentimen. Dalam klasifikasi multikelas, khususnya pada data dengan distribusi kelas yang tidak seimbang, penggunaan accuracy saja dapat memberikan gambaran yang kurang lengkap sehingga evaluasi perlu dilengkapi dengan metrik lain yang lebih sensitif terhadap performa tiap kelas. Oleh karena itu, evaluasi multikelas menjadi penting agar kemampuan model dalam membedakan sentimen positif, negatif, dan netral dapat dinilai secara lebih objektif (Abidin et al., 2025).

4. HASIL DAN PEMBAHASAN

Hasil Preprocessing Teks

Tahap preprocessing memiliki peranan penting dalam analisis sentimen karena kualitas teks masukan sangat memengaruhi performa model klasifikasi. Dalam penelitian ini, preprocessing dilakukan melalui beberapa tahapan utama, yaitu cleaning, normalisasi slang, stopword removal, stemming, dan vectorization menggunakan TF-IDF. Tahapan tersebut diterapkan pada dataset komentar publik mengenai isu korupsi Pertamina agar data lebih bersih, seragam, dan siap digunakan dalam proses klasifikasi.

Pada tahap cleaning, karakter yang tidak diperlukan seperti tanda baca, angka, URL, simbol, dan karakter non-alfabet dihapus, kemudian teks diubah menjadi huruf kecil agar variasi penulisan menjadi seragam. Hasil dari tahap ini ditampilkan pada Tabel 2 sebagai contoh teks asli yang telah dibersihkan.

Tabel 1. Hasil Cleaning.

<i>Text Asli</i>	<i>Clean Text</i>
“data itu tetap gk bisa memungkiri kualitas RON92 plat merah lebih buruk dari si kerang atau si biru”	data itu tetap gk bisa memungkiri kualitas ron92 plat merah lebih buruk dari si kerang atau si biru
“Aint this kinda sus too? Hasil tesnya oleh lembaga negara yang sama aja plat merah.”	aint this kinda sus too hasil tesnya oleh lembaga negara yang sama aja plat merah

Tahap normalisasi slang dilakukan untuk mengubah kata-kata informal dan singkatan yang sering muncul pada komentar media sosial menjadi bentuk yang lebih baku. Tahap ini penting karena bahasa media sosial cenderung bervariasi dan tidak selalu mengikuti kaidah baku. Setelah normalisasi, kata-kata tidak baku diseragamkan agar lebih mudah diolah oleh model.

Selanjutnya, stopwords removal digunakan untuk menghapus kata-kata umum yang tidak memiliki muatan sentimen signifikan. Penghapusan stopwords membantu mengurangi noise dan membuat model lebih fokus pada kata yang benar-benar membawa makna sentimen. Contoh hasil tahap ini ditunjukkan pada Tabel 2.

Tabel 2. Hasil Stopword Removal.

<i>No Stopwords</i>
kualitas ron92 plat merah buruk kerang biru hasil tes lembaga negara sama plat merah

Tahap stemming diterapkan untuk mengubah kata berimbuhan menjadi bentuk dasarnya. Proses ini menyatukan variasi kata yang memiliki akar makna sama sehingga jumlah fitur unik berkurang dan model lebih efisien dalam mengenali pola teks. Contoh hasil stemming ditampilkan pada Tabel 3.

Tabel 3. Hasil Stemming.

Stemming
kualitas ron92 plat merah buruk kerang biru
hasil tes lembaga negara sama plat merah

Setelah preprocessing selesai, data diubah menjadi representasi numerik menggunakan TF-IDF agar dapat diproses oleh algoritma Naïve Bayes. Tahap ini penting karena algoritma klasifikasi tidak dapat langsung membaca teks mentah, melainkan membutuhkan fitur numerik sebagai input.summary_naive_bayes_preprocessed.csv+1

Hasil Klasifikasi dan Evaluasi Model

Tahap pemodelan dilakukan menggunakan Multinomial Naïve Bayes yang dilatih dengan data yang telah melalui preprocessing dan pembobotan TF-IDF. Jumlah data yang digunakan dalam penelitian ini adalah 952 data. Model dilatih menggunakan skema pembagian data latih dan data uji, kemudian dilakukan optimasi parameter smoothing. Parameter terbaik yang diperoleh adalah alpha 0.1.

Hasil evaluasi model menunjukkan bahwa Naïve Bayes memperoleh accuracy sebesar 0.759162 dan macro F1-score sebesar 0.464155. Nilai tersebut menunjukkan bahwa model cukup baik dalam mengklasifikasikan komentar secara keseluruhan, namun masih memiliki keterbatasan dalam membedakan kelas yang lebih kecil dan lebih ambigu.

Tabel 4. Hasil Evaluasi Model Naïve Bayes.

Kelas	Precision	Recall	F1-score	Support
Negative	0.7697	0.9648	0.8562	142
Neutral	0.5714	0.1333	0.2162	30
Positive	0.6667	0.2105	0.32	19
Accuracy	-	-	0.7592	191
Macro avg	0.6693	0.4362	0.4642	191
Weighted avg	0.7283	0.7592	0.7024	191

Berdasarkan Tabel 5, kelas Negative menunjukkan performa terbaik dengan precision 0.7697, recall 0.9648, dan F1-score 0.8562. Nilai recall yang sangat tinggi menunjukkan bahwa hampir seluruh komentar negatif pada data uji berhasil dikenali oleh model.

Sebaliknya, kelas Neutral dan Positive masih menunjukkan performa yang rendah, terutama pada recall yang masing-masing hanya 0.1333 dan 0.2105. Hal ini menunjukkan bahwa model masih kesulitan menangkap komentar yang bersifat netral atau positif secara eksplisit.

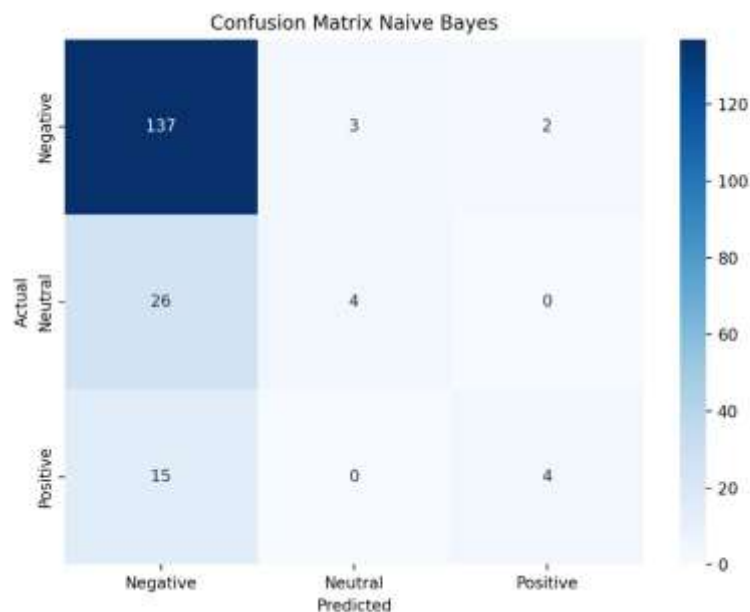
Pola Sentimen, Fitur Dominan, dan Pembahasan

Distribusi label aktual pada dataset menunjukkan bahwa sentimen Negative mendominasi data dengan jumlah 710 data, diikuti Neutral sebanyak 148 data, dan Positive sebanyak 94 data. Dominasi kelas negatif ini menunjukkan bahwa opini publik terhadap isu korupsi Pertamina lebih banyak bernada kritik dan penolakan.

Pada hasil prediksi model terhadap data uji, kelas Negative diprediksi sebanyak 178 data, sedangkan kelas Neutral hanya 7 data dan kelas Positive sebanyak 6 data. Hasil ini memperlihatkan bahwa model memiliki kecenderungan kuat untuk mengklasifikasikan komentar ke kelas negatif, sesuai dengan dominasi kata-kata bernuansa keras pada data.

Tabel 5. Distribusi Label Aktual dan Prediksi.

Kelas	Aktual	Prediksi
Negative	710	178
Neutral	148	7
Positive	94	6



Gambar 2. Confusion Matrix Naïve Bayes.

Berdasarkan confusion matrix pada Gambar 2, dari 142 data aktual Negative, sebanyak 137 berhasil diprediksi benar sebagai Negative, 3 salah diprediksi sebagai Neutral, dan 2 salah diprediksi sebagai Positive. Dari 30 data aktual Neutral, hanya 4 yang berhasil diprediksi benar sebagai Neutral, sedangkan 26 lainnya salah diprediksi sebagai Negative. Pada kelas Positive, dari 19 data aktual, 4 berhasil diprediksi benar sebagai Positive dan 15 salah diprediksi sebagai

Negative. Analisis fitur kata teratas menunjukkan bahwa kelas Negative didominasi oleh kata-kata seperti korupsi, mati, indonesia, negara, hukum, koruptor, ya, rakyat, pertamina, dan hukum mati. Kata-kata tersebut memperlihatkan bahwa komentar negatif banyak berkaitan dengan isu penegakan hukum, kemarahan publik, serta kritik terhadap korupsi.

Pada kelas Neutral, kata dominan yang muncul antara lain danantara, ya, ron, bagaimana, shell, jenis, dulu, berapa, sepertinya, dan alhamdulillah. Kata-kata ini cenderung menggambarkan komentar yang bersifat informatif, tanya-jawab, atau tidak menunjukkan polaritas emosi yang kuat.

Sementara itu, kelas Positive didominasi oleh kata pak, undang, prabowo, rakyat, koruptor, mantap, kejaksaan, presiden, shell, dan korupsi. Dominasi kata-kata tersebut menunjukkan bahwa komentar positif lebih banyak muncul dalam bentuk dukungan terhadap tindakan tegas atau harapan terhadap penegakan hukum.

Keterbatasan dan Implikasi

Penelitian ini memiliki beberapa keterbatasan. Pertama, model masih mengandalkan pendekatan tradisional berbasis TF-IDF, sehingga belum menangkap konteks semantik yang lebih dalam seperti model berbasis transformer. Kedua, distribusi kelas masih belum seimbang, terutama pada kelas Neutral dan Positive, sehingga berdampak pada rendahnya recall pada dua kelas tersebut. Ketiga, komentar media sosial yang sangat informal dan sarat slang masih menyisakan tantangan meskipun sudah dilakukan normalisasi.

Meskipun demikian, hasil penelitian ini tetap memberikan gambaran yang jelas mengenai respons publik terhadap isu korupsi Pertamina. Informasi ini dapat dimanfaatkan untuk membaca arah opini publik, membantu pihak terkait memahami persepsi masyarakat, dan menjadi dasar bagi penelitian lanjutan dengan metode klasifikasi yang lebih kompleks.

5. KESIMPULAN DAN SARAN

Kesimpulan Berdasarkan hasil penelitian, dapat disimpulkan bahwa algoritma Naïve Bayes mampu digunakan untuk menganalisis sentimen publik terhadap isu korupsi Pertamina dengan hasil yang cukup baik, terutama dalam mengenali sentimen negatif. Dari total 952 data yang digunakan, model memperoleh accuracy sebesar 0.7592 dan macro F1-score sebesar 0.4642 dengan parameter terbaik alpha 0.1. Hasil evaluasi menunjukkan bahwa kelas Negative memiliki performa terbaik dengan recall 0.9648 dan F1-score 0.8562, sedangkan kelas Neutral dan Positive masih memiliki performa yang rendah, khususnya pada nilai recall. Temuan ini menunjukkan bahwa opini publik terhadap isu korupsi Pertamina cenderung didominasi oleh sentimen negatif, yang tercermin pula dari kata-kata dominan seperti korupsi, hukum, koruptor,

dan mati pada hasil analisis fitur. Meskipun hasil penelitian ini telah memberikan gambaran yang jelas mengenai kecenderungan sentimen publik, penelitian ini masih memiliki keterbatasan, terutama karena pendekatan yang digunakan masih berbasis TF-IDF dan Naïve Bayes sehingga belum sepenuhnya mampu menangkap konteks bahasa yang lebih kompleks, terutama pada komentar netral dan positif. Selain itu, distribusi kelas yang tidak seimbang turut memengaruhi performa model dalam mengenali kelas minoritas. Oleh karena itu, penelitian selanjutnya disarankan untuk menggunakan teknik penyeimbangan data atau metode klasifikasi yang lebih kompleks, seperti Support Vector Machine atau model berbasis transformer, agar hasil klasifikasi pada kelas netral dan positif dapat ditingkatkan. Penelitian lanjutan juga dapat memperluas sumber data dari berbagai platform media sosial agar hasil analisis sentimen menjadi lebih representatif terhadap persepsi publik secara umum.

DAFTAR REFERENSI

- Abdullah Wahid, Y., & Darma Andayani, D. (2025). Performance comparison of SVM and naïve Bayes for Indonesian-language sentiment analysis on Free Fire reviews using TF-IDF and SMOTE. *Journal of Embedded System Security and Intelligent Systems*, 6(4), 674–689.
- Abidin, D. Z., Afuan, L., Toscani, A. N., & Nurhadi, N. (2025). A comprehensive benchmarking pipeline for transformer-based sentiment analysis using cross-validated metrics. *Jurnal Teknik Informatika (JUTIF)*, 6(4), 1797–1810. <https://doi.org/10.52436/1.jutif.2025.6.4.4894>
- Ardi, A., & Kurniawan. (2024). Optimasi metode naïve Bayes classifier menggunakan pendekatan term frequency–inverse document frequency (TF-IDF) pada analisis sentimen. *JSAI (Journal Scientific and Applied Informatics)*, 7(3), 458–463. <https://doi.org/10.36085/jsai.v7i3.7153>
- Arya, M., Nugraha, A., Budiarto, S., & Supriyadi, S. (2026). Implementasi metode naïve Bayes untuk analisis sentimen terhadap program makan siang gratis pada media sosial X. *XVI*.
- Astuti, Y. D. (2025). Online media and narrative hegemony: Discourse network analysis of the 2025 Pertamina case. *10*(2), 347–372.
- Hartimar, L., Manza, Y., & Putriani Siregar, K. (2025). Text classification using TF-IDF and naïve Bayes: Case study of MyXL app user review data. *Journal of Technology and Computer (JOTECHCOM)*, 2(2), 100–108.
- Katya, E., & Rahman, S. R. (2024). Applications of natural language processing in social media sentiment analysis.
- Khoerunnisa, S., Shiddieq, D. F., & Nurhayati, D. (2025). Penerapan algoritma naïve Bayes dengan teknik TF-IDF dan cross validation untuk analisis sentimen terhadap Starlink. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 5(2), 566–577. <https://doi.org/10.57152/malcom.v5i2.1852>

- Lestari, V. B., & Hutagalung, C. A. (2025). Evaluation of TF-IDF extraction techniques in sentiment analysis of Indonesian-language marketplaces using SVM, logistic regression, and naïve Bayes. *Journal of Computer Science and Applications*, 8(1), 22–35. <https://doi.org/10.21009/j->
- Muhaimin, A., Fahrudin, T. M., Alamiyah, S. S., & Arviani, H. (2024). Sentiment analysis in social media: Case study in Indonesia. 2024, 27–30. <https://doi.org/10.11594/nstp.2024.4106>
- Pebrianti, R. D. (2025). Analisis sentimen masyarakat platform X terhadap korupsi PT Pertamina (Persero) menggunakan SVM. *JITET (Jurnal Informatika dan Teknik Elektro Terapan)*, 13(2).
- Ramadhan, M. F., Panjaitan, F., & Oktafiandi, H. (2025). Analisis sentimen kutipan media sosial berbahasa Indonesia menggunakan convolutional neural network. 2, 1–17.
- Rasmila, R., Saputri, Y., Syaki, F., & Hadinata, N. (2025). Sentiment analysis of trending topics on social media X using natural language processing and LSTM. *Journal of Applied Informatics and Computing*, 9(6), 3034–3041. <https://doi.org/10.30871/jaic.v9i6.10931>
- Rodríguez-Ibáñez, M., Casáñez-Ventura, A., Castejón-Mateos, F., & Cuenca-Jiménez, P. M. (2023). A review on sentiment analysis from social media platforms. *Expert Systems with Applications*, 223, Article 119862. <https://doi.org/10.1016/j.eswa.2023.119862>
- W, A. W., Andana, H. H., Zeniarja, J., & Febriyanto, A. (2025). Sentiment analysis of the 2024 general election through Twitter using long short-term memory algorithm. *Journal of Informatics and Web Engineering*, 4(2), 387–400. <https://doi.org/10.33093/jiwe.2025.4.2.25>