



Opini Publik terhadap Pernyataan Presiden Prabowo tentang Proyek Whoosh: Analisis Sentimen Komentar YouTube Menggunakan Multinomial Naive Bayes dalam Perspektif *Commentary Engagement*

Dimas Adiz Setiawan^{1*}, Dicky Andi Soffan², Harun Al Rosyid³

¹⁻³ Pendidikan Teknologi Informasi, Universitas Negeri Surabaya, Indonesia

Email: dimasadiz.22088@mhs.unesa.ac.id, dicky.22093@mhs.unesa.ac.id, harunrosyid@unesa.ac.id

*Penulis Korespondensi: dimasadiz.22088@mhs.unesa.ac.id

Abstract. YouTube plays a crucial role in shaping and expressing public opinion on national political issues. President Prabowo Subianto's statement regarding the Whoosh High-Speed Rail project in the video "Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Tanggung" (Firm! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Tanggung) sparked a variety of responses in the comments section. This study aims to analyze public comment sentiment and its relationship to user engagement levels. The method used is sentiment analysis based on the Multinomial Naive Bayes algorithm with TF-IDF weighting, through text preprocessing and K-Fold Cross Validation ($k = 2, 5$, and 10) validation stages. The best results were obtained with the 10-fold scheme with an accuracy of 79.45%, a precision of 82.05%, and a recall of 79.45%. The sentiment distribution shows a dominance of neutral sentiment (59.41%), followed by negative (26.81%) and positive (13.78%), indicating a tendency for users to express opinions informatively and critically. The engagement level is relatively low with an average of 2.61 likes, 0.30 replies, and an engagement score of 0.0036. Correlation analysis shows that sentiment does not have a significant effect on engagement, so sentiment analysis is more appropriate for monitoring public opinion than predicting user engagement.

Keywords: Naïve Bayes; Public Opinion; Sentiment Analysis; Whoosh; Youtube.

Abstrak. YouTube berperan penting sebagai ruang pembentukan dan ekspresi opini publik terhadap isu politik nasional. Pernyataan Presiden Prabowo Subianto mengenai proyek Kereta Cepat Whoosh dalam video "Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Jawab" memicu beragam respons di kolom komentar. Penelitian ini bertujuan menganalisis sentimen komentar publik serta hubungannya dengan tingkat engagement pengguna. Metode yang digunakan adalah analisis sentimen berbasis algoritma Multinomial Naive Bayes dengan pembobotan TF-IDF, melalui tahapan preprocessing teks dan validasi K-Fold Cross Validation ($k = 2, 5$, dan 10). Hasil terbaik diperoleh pada skema 10-fold dengan akurasi 79,45%, presisi 82,05%, dan recall 79,45%. Distribusi sentimen menunjukkan dominasi sentimen netral (59,41%), diikuti negatif (26,81%) dan positif (13,78%), yang mengindikasikan kecenderungan pengguna menyampaikan opini secara informatif dan kritis. Tingkat engagement tergolong rendah dengan rata-rata likes 2,61, reply 0,30, dan engagement score 0,0036. Analisis korelasi menunjukkan bahwa sentimen tidak berpengaruh signifikan terhadap engagement, sehingga analisis sentimen lebih tepat digunakan untuk memantau opini publik dibandingkan memprediksi keterlibatan pengguna.

Kata kunci: Analisis Sentimen; Naïve Bayes; Opini Publik; Whoosh; Youtube.

1. LATAR BELAKANG

Perkembangan teknologi informasi dan komunikasi dalam sepuluh tahun terakhir telah secara mendasar mengubah cara orang-orang di Indonesia mendapatkan informasi dan menyampaikan pendapat. Media sosial seperti YouTube kini tidak hanya berfungsi sebagai tempat hiburan, tetapi juga menjadi area yang signifikan untuk komunikasi politik, pembentukan pandangan masyarakat, dan keterlibatan warga dalam masalah-masalah kebijakan. Penelitian oleh (Al Rasyid, Handayani, & Ningsih, 2024) mengindikasikan bahwa media sosial, termasuk YouTube, semakin banyak digunakan sebagai sarana untuk

berpartisipasi dalam politik secara tidak langsung, yaitu melalui konsumsi dan pembuatan komentar mengenai isu-isu yang berkaitan dengan pemerintahan dan politik.

Pernyataan Presiden Prabowo Subianto mengenai kontroversi proyek Kereta Cepat Jakarta–Bandung (Whoosh) yang ditampilkan dalam video YouTube berjudul “Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Jawab”—menjadi menarik untuk dianalisis. Video itu menunjukkan sikap tegas dari presiden, yang mengungkapkan kesediaan untuk bertanggung jawab atas polemik Whoosh serta berusaha menenangkan kekhawatiran masyarakat mengenai dampak keuangan terhadap negara dan Badan Usaha Milik Negara (BUMN) (Argueta & Chen, 2014). Isu yang berkaitan dengan pendanaan, lonjakan biaya, dan restrukturisasi utang proyek Whoosh sebelumnya telah menjadi topik bahasan di berbagai media mainstream hingga sosial, sehingga tidak mengherankan jika video ini menghasilkan reaksi yang kuat di bagian komentar (Misrun, Haerani, Fikry, & Budianita, 2023).

Penelitian ini bertujuan untuk menganalisis pandangan masyarakat terhadap komentar-komentar yang muncul pada video YouTube berjudul “Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Jawab” (Ciaramita & Johnson, 2024). Semua data komentar akan dikumpulkan secara keseluruhan mengacu pada jumlah yang tertera di video, lalu akan dibagi ke dalam tiga kategori sentimen—positif, negatif, dan netral (Khoirul, Hayati, & Nurdiawan, 2023). Selanjutnya, penelitian ini akan mengevaluasi hubungan antara kategori sentimen dengan jumlah suka dan jumlah balasan pada setiap komentar. Diharapkan hasil dari penelitian ini dapat memberikan insights tentang kecenderungan opini publik terhadap pernyataan Presiden Prabowo mengenai proyek Whoosh yang sudah berjalan.

2. KAJIAN TEORITIS

Analisis Sentimen

Analisis sentimen merupakan bagian dari *Natural Language Processing* (NLP) yang bertujuan untuk mengidentifikasi dan mengklasifikasikan opini atau sikap yang terkandung dalam teks ke dalam kategori tertentu, seperti positif, negatif, atau netral. Analisis sentimen banyak digunakan untuk memahami persepsi publik terhadap suatu isu melalui data teks yang bersumber dari media sosial, forum daring, maupun ulasan pengguna. Dengan kemampuan mengolah data teks dalam jumlah besar secara otomatis, analisis sentimen mampu memberikan gambaran kecenderungan opini masyarakat secara objektif dan terukur (Manoppo et al., 2025).

Media Sosial YouTube sebagai Sumber Data

YouTube merupakan platform media sosial berbasis video yang menyediakan kolom komentar sebagai sarana interaksi pengguna. Kolom komentar ini sering dimanfaatkan untuk

menyampaikan pendapat, kritik, maupun dukungan terhadap suatu konten, termasuk isu politik dan kebijakan publik. Selain komentar teks, YouTube juga menyediakan indikator keterlibatan (*engagement*) berupa jumlah *likes* dan *replies* yang dapat mencerminkan tingkat respons dan intensitas diskusi pengguna. Oleh karena itu, komentar YouTube dianggap sebagai sumber data yang relevan dan representatif dalam menganalisis opini publik secara daring (Newburn, 2020).

Preprocessing Teks

Preprocessing teks merupakan tahap awal yang sangat penting dalam analisis sentimen karena kualitas data pada tahap ini sangat memengaruhi performa model klasifikasi. Tahapan preprocessing bertujuan untuk membersihkan dan menyeragamkan data teks agar lebih mudah diproses secara komputasional. Proses preprocessing umumnya meliputi *cleaning*, *case folding*, *tokenizing*, *stopword removal*, dan *stemming*. Dengan melakukan preprocessing yang tepat, *noise* dalam data dapat dikurangi sehingga fitur teks yang dihasilkan menjadi lebih representatif dan informatif (Adriana, Suarjaya, & Githa, 2023).

Pembobotan TF-IDF

Term Frequency–Inverse Document Frequency (TF-IDF) merupakan metode pembobotan kata yang digunakan untuk mengukur tingkat kepentingan suatu kata dalam dokumen. TF-IDF menggabungkan frekuensi kemunculan kata dalam satu dokumen dan tingkat kelangkaannya dalam keseluruhan dokumen. Metode ini efektif dalam menonjolkan kata-kata yang memiliki kontribusi besar terhadap isi dokumen, sekaligus menekan pengaruh kata yang terlalu umum. Dalam analisis sentimen, TF-IDF digunakan untuk mengubah data teks menjadi representasi numerik agar dapat diproses oleh algoritma klasifikasi (Ningtyas, Solichin, & Pradana, 2023).

Algoritma Multinomial Naive Bayes

Multinomial Naive Bayes merupakan salah satu varian algoritma Naive Bayes yang banyak digunakan dalam klasifikasi teks. Algoritma ini bekerja berdasarkan pendekatan probabilistik dengan asumsi independensi antar fitur dan memperhitungkan frekuensi kemunculan kata dalam dokumen. Multinomial Naive Bayes dikenal memiliki keunggulan dalam hal kesederhanaan, efisiensi komputasi, serta performa yang cukup baik pada permasalahan klasifikasi sentimen berbasis teks. Oleh karena itu, algoritma ini sering digunakan dalam penelitian analisis sentimen pada media sosial (Nurfebria & Sriani, 2024).

K-Fold Cross Validation

K-Fold Cross Validation merupakan metode evaluasi model yang digunakan untuk menguji kestabilan dan keandalan kinerja algoritma klasifikasi. Pada metode ini, dataset dibagi menjadi *k* bagian (*fold*), kemudian proses pelatihan dan pengujian dilakukan secara bergantian

sebanyak k kali. Teknik ini bertujuan untuk mengurangi bias akibat pembagian data tunggal serta memberikan estimasi performa model yang lebih akurat. Penggunaan beberapa nilai k memungkinkan peneliti untuk membandingkan kinerja model secara lebih komprehensif (Pavitha et al., 2022).

Confusion Matrix

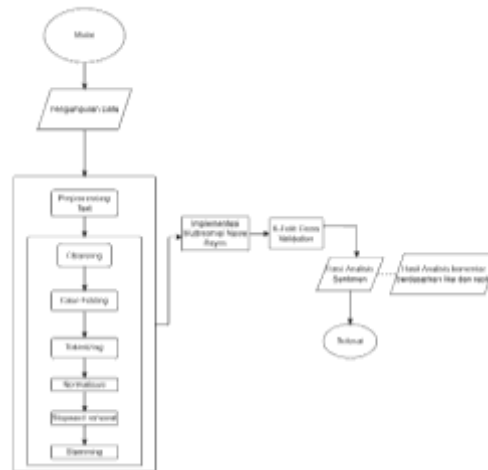
Confusion matrix merupakan alat evaluasi yang digunakan untuk membandingkan hasil prediksi model dengan label sebenarnya. Dari confusion matrix dapat dihitung beberapa metrik evaluasi, yaitu akurasi, presisi, dan recall. Akurasi menunjukkan tingkat ketepatan model secara keseluruhan, presisi menggambarkan ketepatan prediksi pada kelas tertentu, sedangkan recall menunjukkan kemampuan model dalam mengenali data yang benar pada kelas tersebut. Metrik-metrik ini digunakan untuk menilai kinerja algoritma klasifikasi sentimen secara objektif dan terukur (Sains & Teknologi, 2021).

Engagement (Likes dan Replies)

Engagement merupakan indikator yang digunakan untuk mengukur tingkat interaksi pengguna terhadap suatu konten di media sosial. Pada platform YouTube, engagement dapat diamati melalui jumlah *likes* dan *replies* pada komentar. Analisis engagement digunakan untuk memahami sejauh mana suatu opini mendapatkan perhatian dan memicu diskusi lanjutan. Dengan mengkaji hubungan antara sentimen komentar dan engagement, penelitian ini dapat memberikan pemahaman yang lebih mendalam mengenai pola interaksi dan respons publik terhadap isu yang dibahas (Tarigan, 2023).

3. METODE PENELITIAN

Metodologi penelitian ini adalah pendekatan kuantitatif berbasis data mining dengan tahapan pengolahan teks dan klasifikasi (Yutika, Adiwijaya, & Faraby, 2020). Terdiri dari 5 tahap yaitu pengumpulan data, pre-processing, Implementasi Naïve Bayes, K-Fold Cross Validation, dan Hasil analisis sentimen seperti gambar:



Gambar 1. Alur Analisis Sentimen.

Pengumpulan Data

Data dalam penelitian ini dikumpulkan dari kolom komentar pada video YouTube berjudul *“Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Jawab”*. Seluruh komentar yang tampil pada video diambil beserta informasi pendukungnya, yaitu jumlah *like* dan jumlah *reply* pada setiap komentar

Preprocessing Text

Data yang sudah diperoleh kemudian diproses melalui tahapan pre-processing seperti pembersihan teks, tokenizing, penghapusan stopwords, serta stemming untuk mendapatkan teks yang lebih terstruktur.

Sebelum memasuki fase cleaning, data melalui tahap labeling terlebih dahulu. Labeling adalah kegiatan memberi nama atau kategori pada data supaya bisa dipahami dan diproses lebih mudah. Dalam konteks data, memberikan informasi tambahan yang menjelaskan arti atau isi dari data, misalnya menandai label sebagai positif, negatif, dan netral itu merupakan proses labeling.

Tahap *cleansing* bertujuan untuk membersihkan teks dari karakter yang tidak relevan dan tidak memiliki makna semantik terhadap analisis, seperti emoji, simbol khusus, karakter non-ASCII, serta tanda baca yang tidak diperlukan. Proses ini dilakukan untuk mengurangi *noise* pada data sehingga hanya menyisakan teks yang bersifat informatif dan dapat diproses lebih lanjut pada tahap berikutnya.

Tabel 1. Tahap *Cleansing*

Before	After
Kalo tanggung jawab pake uang pribadi si no problem tp kalo mash pake duit rakyat bkn tanggung jawab, tapi tanggung bersama ðŸ˜,	Kalo tanggung jawab pake uang pribadi si no problem tp kalo mash pake duit rakyat bkn tanggung jawab tapi tanggung bersama

Proses case folding bertujuan untuk mentransformasi semua huruf kapital dan kecil dalam teks menjadi format lowercase secara keseluruhan. (Tabel 1). Hal ini dilakukan untuk menghindari perbedaan pengolahan data akibat perbedaan penggunaan huruf besar dan kecil, sehingga memastikan konsistensi dalam analisis teks

Tabel 2. Tahap *Case Folding*.

Before	After
Kalo tanggung jawab pake uang pribadi si no problem tp kalo mash pake duit rakyat bkn tanggung jawab tapi tanggung bersama	kalo tanggung jawab pake uang pribadi si no problem tp kalo mash pake duit rakyat bkn tanggung jawab tapi tanggung bersama

Tokenisasi adalah metode untuk menguraikan teks menjadi elemen-elemen dasar seperti kata, frasa, atau kalimat yang disebut sebagai token. (Tabel 3). Tokenizing membantu dalam mengidentifikasi kata-kata penting yang akan digunakan dalam analisis lebih lanjut

Tabel 3. Tahap *Tokenisasi*.

Before	After
kalo tanggung jawab pake uang pribadi si no problem tp kalo mash pake duit rakyat bkn tanggung jawab tapi tanggung bersama	['kalo', 'tanggung', 'jawab', 'pake', 'uang', 'pribadi', 'si', 'no', 'problem', 'tp', 'kalo', 'mash', 'pake', 'duit', 'rakyat', 'bkn', 'tanggung', 'jawab', 'tapi', 'tanggung', 'bersama']

Tahap normalisasi bertujuan untuk mengubah kata-kata tidak baku, singkatan, atau bahasa percakapan khas media sosial menjadi bentuk kata baku sesuai dengan kaidah Bahasa Indonesia.

Tabel 4. Tahap *Normalisasi*.

Before	After
['kalo', 'tanggung', 'jawab', 'pake', 'uang', 'pribadi', 'si', 'no', 'problem', 'tp', 'kalo', 'mash', 'pake', 'duit', 'rakyat', 'bkn', 'tanggung', 'jawab', 'tapi', 'tanggung', 'bersama']	['kalau', 'tanggung', 'jawab', 'pakai', 'uang', 'pribadi', 'si', 'tidak', 'masalah', 'tapi', 'kalau', 'masih', 'pakai', 'uang', 'rakyat', 'bukan', 'tanggung', 'jawab', 'tapi', 'tanggung', 'bersama']

Stopword merupakan teknik untuk memfilter kata-kata yang tidak mengandung makna penting dalam teks, termasuk kata penghubung dan kata depan. Contoh stopwords dalam bahasa Indonesia adalah "dan", "atau", "yang", "di", dan lain-lain (Tabel 4). Tujuannya adalah untuk mengurangi noise dalam data teks.

Tabel 5. Tahap *Stopword*.

Before	After
['kalau', 'tanggung', 'jawab', 'pakai', 'uang', 'pribadi', 'si', 'tidak', 'masalah', 'tapi', 'kalau', 'masih', 'pakai', 'uang', 'rakyat', 'bukan', 'tanggung', 'jawab', 'tapi', 'tanggung', 'bersama']	['tanggung', 'jawab', 'pakai', 'uang', 'pribadi', 'tidak', 'masalah', 'pakai', 'uang', 'rakyat', 'bukan', 'tanggung', 'jawab', 'tanggung', 'bersama']

Stemming adalah proses mengubah kata berimbuhan (misalnya, kata kerja dengan awalan atau akhiran) menjadi bentuk dasar. Tujuannya yaitu untuk menyederhanakan variasi kata sehingga kata-kata yang memiliki arti sama dapat dianggap sama dalam analisis teks.

Tabel 6. Tahap *Stemming*.

before	After
['tanggung', 'jawab', 'pakai', 'uang', 'pribadi', 'tidak', 'masalah', 'pakai', 'uang', 'rakyat', 'bukan', 'tanggung', 'jawab', 'tanggung', 'bersama']	['tanggung', 'jawab', 'pakai', 'uang', 'pribadi', 'tidak', 'masalah', 'pakai', 'uang', 'rakyat', 'bukan', 'tanggung', 'jawab', 'tanggung', 'sama']

Implementasi Multinomial Naïve Bayes

Multinomial Naïve Bayes ini termasuk bagian dari Naïve Bayes yang digunakan dalam untuk klasifikasi sederhana. Tahapan metode Multinomial Naïve Bayes ini melalui perhitungan probabilitas prior, probability kata-kata dan probabilitas dokumen (Yutika, Adiwijaya, & Faraby, 2021).

K-Fold Cross Validation

K-Fold Cross Validation adalah Teknik validasi yang digunakan untuk mengurangi bias pada data dokumen, Teknik ini digunakan dalam penerapan perulangan model pada setiap data menjadi data latih dan data uji sehingga teruji validitasnya.

Confusion Matrix

Untuk mengevaluasi kinerja model klasifikasi Multinomial Naive Bayes, penelitian ini menggunakan confusion matrix

Pada kasus dengan tiga kelas sentimen (positif, negatif, dan netral), confusion matrix akan berbentuk matriks 3x3 yang berisi:

- 1). True Positive (TP) adalah banyaknya data prediksi positif yang berhasil diklasifikasi dengan benar.
- 2) True Negative (TN) adalah banyaknya data prediksi negatif yang berhasil diklasifikasi dengan benar.
- 3) False Positive (FP) adalah banyaknya data prediksi negatif yang salah diklasifikasikan sebagai data prediksi positif.
- 4) False Negative (FN) adalah banyaknya data prediksi positif yang salah diklasifikasikan sebagai data prediksi negatif.

4. HASIL DAN PEMBAHASAN

Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan melalui proses *crawling* menggunakan Google Collab.



Gambar 2. Data Crawling.

Labeling Data

Proses pelabelan data dalam penelitian ini mencakup dua kategori sentimen, yaitu sentimen positif dan sentimen negatif di mana angka 0 menunjukkan komentar bernuansa negatif dan angka 1 menunjukkan komentar bernuansa positif.

Preprocessing Text

Tahap ini bertujuan menghilangkan *noise*, membersihkan, serta menyederhanakan teks sehingga memiliki makna yang lebih jelas. Adapun tahapan *preprocessing* yang dilakukan adalah sebagai berikut:

1) Cleansing

Gambar 4 menunjukkan hasil proses cleansing pada komentar YouTube, di mana karakter tidak relevan seperti emoji, simbol khusus, dan tanda baca yang tidak diperlukan telah dihapus. Tahap ini bertujuan mengurangi *noise* sehingga teks menjadi lebih bersih dan siap diproses pada tahap berikutnya.

comment	cleansed_text
Gua kira pinter , dungu juga wkww	Gua kira pinter , dungu juga wkww
120 T pak Hutang buka LabaðŸ™, ðŸ™™	120 T pak Hutang buka Laba

Gambar 3. Hasil Cleansing.

2) Case Folding

Gambar 5 memperlihatkan hasil case folding dengan mengubah seluruh huruf menjadi lowercase. Proses ini dilakukan untuk menyeragamkan bentuk kata dan mencegah perbedaan makna akibat variasi penggunaan huruf kapital.

cleansed_text	case_folded_text
Gua kira pinter , dungu juga wkww	gua kira pinter , dungu juga wkww
120 T pak Hutang buka Laba	120 t pak hutang buka laba

Gambar 4. Hasil Case Folding.

3) Tokenizing

Gambar 6 menampilkan hasil tokenizing, yaitu pemecahan teks menjadi token kata. Tahap ini bertujuan memudahkan identifikasi kata-kata penting yang akan digunakan dalam proses analisis sentimen.

case_folded_text	tokens
gua kira pinter , dungu juga wkww	['gua', 'kira', 'pinter', ',', 'dungu', 'juga', 'wkww']
120 t pak hutang buka laba	['120', 't', 'pak', 'hutang', 'buka', 'laba']

Gambar 5. Hasil Tokenizing.

4) Normalisasi

Gambar 7 menunjukkan hasil normalisasi, di mana kata tidak baku dan singkatan khas media sosial diubah menjadi kata baku. Proses ini dilakukan untuk meningkatkan konsistensi dan keakuratan representasi teks.

tokens	normalized_tokens
['gua', 'kira', 'pinter', ',', 'dungu', 'juga', 'wkww']	['saya', 'kira', 'pinter', ',', 'dungu', 'juga', 'wkww']
['120', 't', 'pak', 'hutang', 'buka', 'laba']	['120', 't', 'pak', 'hutang', 'buka', 'laba']

Gambar 6. Hasil Normalisasi.

5) Stopword Removal

Gambar 8 memperlihatkan hasil penghapusan stopwords, yaitu kata-kata umum yang tidak memiliki kontribusi signifikan terhadap analisis sentimen. Tahap ini membuat teks lebih ringkas dan informatif.

normalized_tokens	stopped_tokens
['saya', 'kira', 'pinter', ',', 'dungu', 'juga', 'wkwwk']	['kira', 'pinter', ',', 'dungu', 'wkwwk']
['120', 't', 'pak', 'hutang', 'buka', 'laba']	['120', 't', 'pak', 'hutang', 'buka', 'laba']

Gambar 7. Hasil Stopword Removal.

6) Stemming

Gf5 Gambar 9 menunjukkan hasil stemming, yaitu pengubahan kata berimbuhan menjadi kata dasar. Proses ini bertujuan menyederhanakan variasi kata agar analisis teks menjadi lebih efektif.

stopped_tokens	stemmed_tokens
['kira', 'pinter', ',', 'dungu', 'wkwwk']	['kira', 'pinter', ',', 'dungu', 'wkwwk']
['120', 't', 'pak', 'hutang', 'buka', 'laba']	['120', 't', 'pak', 'hutang', 'buka', 'laba']

Gambar 8. Hasil Stemming.

Pembobotan TF-IDF

Setelah tahap *preprocessing* teks, dilakukan perhitungan *Term Frequency* (TF) untuk mengetahui frekuensi kemunculan setiap token (term) dalam dataset, serta *Inverse Document Frequency* (IDF) untuk menentukan bobot berdasarkan tingkat keterjadian term.

processed text	pembobotan tfidf
kira pinter dungu wkwwk	['0.7071', '0.7071']
120 t pak hutang buka laba	['0.7968', '0.5131', '0.3190']
buat jalan aja lah pak lebih guna banyak jalan rusak parah	['0.6308', '0.3825', '0.3035', '0.2914', '0.2703', '0.2485', '0.2234', '0.2174', '0.1467']

Gambar 9. Hasil TF IDF.

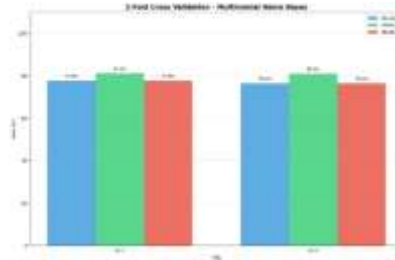
Pembagian dan Pengujian (K-Fold Cross Validation)

1) 2-Fold Cross Validation

Pada skema *2-fold cross validation*, dataset dibagi menjadi dua bagian dengan jumlah yang sama, yaitu masing-masing terdiri atas 2 fold masing-masing 1557 data training dan 1557 data testing. Hasil pengujian menunjukkan bahwa pada langkah kedua diperoleh kinerja model terbaik, dengan akurasi sebesar 76.4%, presisi 80.9%, dan *recall* sebesar 76.4%.

Tabel 7. Hasil 2-Fold Cross Validation.

2-Fold	Multinomial Naïve Bayes	Akurasi	Presisi	Recall
Uji 1		77.6%	81%	77.6%
Uji 2		76.4%	80.9%	76.4%

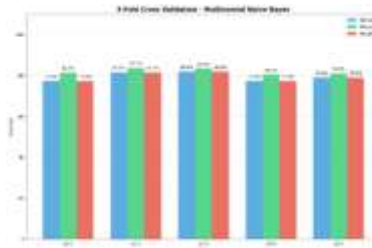
**Gambar 10.** Visual 2-Fold Cross Validation.

2) 5-Fold Cross Validation

Pada skema *5-fold cross validation*, dataset dibagi menjadi dua bagian dengan jumlah yang sama, yaitu masing-masing terdiri atas 5 fold, fold 1 sampai 4 2491 data training dan 623 data testing, fold 5 2492 data training dan 622 data sampel. Hasil pengujian menunjukkan bahwa pada langkah kelima diperoleh kinerja model terbaik, dengan akurasi sebesar 78.9%, presisi 81%, dan *recall* sebesar 78.9%.

Tabel 8. Hasil 5-Fold Cross Validation.

5-Fold	Multinomial Naïve Bayes	Akurasi	Presisi	Recall
Uji 1		77.4%	81.4%	77.4%
Uji 2		81.5%	83.7%	81.5%
Uji 3		81.9%	83.4%	81.9%
Uji 4		77.4%	80.4%	77.4%
Uji 5		78.9%	81%	78.9%

**Gambar 11.** Visual 5-Fold Cross Validation.

3) 10-Fold Cross Validation

Pada skema *10-fold cross validation*, dataset dibagi menjadi dua bagian dengan jumlah yang sama, yaitu masing-masing terdiri atas 10 fold, fold 1 sampai 4 2802 data training dan 312 data testing, fold 5 sampai 10 2803 data training dan 311 data sampel. Hasil pengujian menunjukkan bahwa pada langkah kelima diperoleh kinerja model terbaik, dengan akurasi sebesar 78.9%, presisi 81%, dan *recall* sebesar 78.9%.

Tabel 9. Hasil 10-Fold Cross Validation.

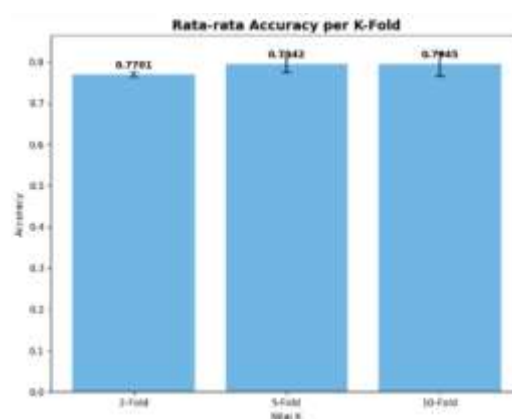
Uji (10-Fold)	Akurasi (%)	Presisi (%)	Recall (%)
Uji 1	78.2	81.6	78.2
Uji 2	76.3	81.1	76.3
Uji 3	81.4	83.1	81.4
Uji 4	82.7	84.4	82.7
Uji 5	85.2	86.4	85.2
Uji 6	77.2	79.3	77.2
Uji 7	78.1	80.7	78.1
Uji 8	76.8	80.2	76.8
Uji 9	80.7	82.5	80.7
Uji 10	77.8	81.3	77.8



Gambar 12. Visual 10-Fold Cross Validation.

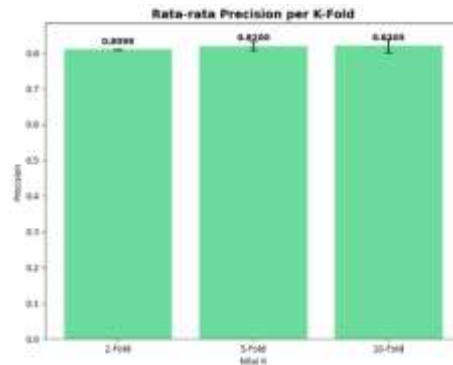
Hasil Analisis K-Fold

Berdasarkan hasil pengujian *k-fold cross validation* dengan nilai *k* sebesar 2, 5, dan 10 menggunakan algoritma *Multinomial Naïve Bayes*, diperoleh nilai keseluruhan untuk metrik akurasi, presisi, dan *recall*.



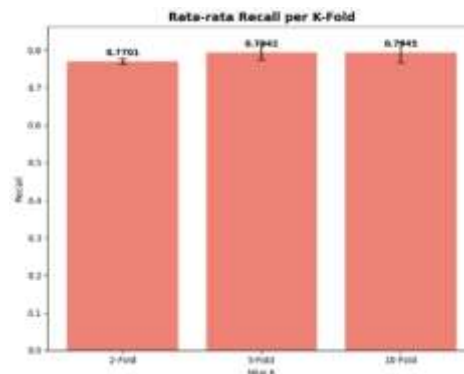
Gambar 13. Hasil Analisis K-Fold.

Berdasarkan gambar tersebut, nilai akurasi tertinggi pada pengujian *k-fold cross validation* dengan algoritma *Multinomial Naïve Bayes* diperoleh pada skema 10-fold, dengan akurasi mencapai 0.7945 atau 79.45%.



Gambar 14. Hasil Akurasi.

Berdasarkan gambar tersebut, nilai presisi tertinggi pada pengujian *k-fold cross validation* menggunakan algoritma *Multinomial Naïve Bayes* diperoleh pada skema 10-fold, dengan capaian presisi sebesar 0.8205 atau 82.05%.



Gambar 15. Hasil Presisi.

Berdasarkan gambar tersebut, nilai *recall* tertinggi pada pengujian *k-fold cross validation* dengan algoritma *Multinomial Naïve Bayes* terlihat pada skema 10-fold, dengan nilai *recall* mencapai 0.7945 atau 79.45%.

a) Hasil Analisis Sentimen

YouTube plays a crucial role in shaping and expressing public opinion on national political issues. President Prabowo Subianto's statement regarding the Whoosh High-Speed Rail project in the video "Tegas! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Tanggung" (Firm! Prabowo: Apa Itu Ribut-Ribut Whoosh, Saya Tanggung Tanggung) sparked a variety of responses in the comments section. This study aims to analyze public comment sentiment and its relationship to user engagement levels. The method used is sentiment analysis based on the Multinomial Naive Bayes algorithm with TF-IDF weighting, through text preprocessing and K-Fold Cross Validation ($k = 2, 5$, and 10) validation stages. The best results were obtained with the 10-fold scheme with an accuracy of 79.45%, a precision of 82.05%, and a recall of 79.45%. The sentiment distribution

shows a dominance of neutral sentiment (59.41%), followed by negative (26.81%) and positive (13.78%), indicating a tendency for users to express opinions informatively and critically. The engagement level is relatively low with an average of 2.61 likes, 0.30 replies, and an engagement score of 0.0036. Correlation analysis shows that sentiment does not have a significant effect on engagement, so sentiment analysis is more appropriate for monitoring public opinion than predicting user engagement.

b) Tingkat Engagement Pengguna

Hasil pengukuran engagement menunjukkan nilai rata-rata *likes* sebesar 2,61, *reply* sebesar 0,30, dan *engagement score* sebesar 0,0036. Nilai tersebut mengindikasikan bahwa interaksi pengguna berada pada tingkat yang relatif rendah. Mayoritas komentar tidak menghasilkan respons lanjutan, sehingga pola interaksi cenderung bersifat pasif dan terbatas pada pemberian *like*.

c) Hubungan Sentimen dengan Engagement

Analisis engagement berdasarkan kategori sentimen menunjukkan bahwa komentar positif memiliki nilai engagement score tertinggi, meskipun perbedaannya sangat kecil dibandingkan sentimen netral dan negatif. Hasil uji ANOVA menghasilkan nilai p sebesar 0,9807, yang menunjukkan tidak adanya perbedaan engagement yang signifikan antar kategori sentimen. Dengan demikian, sentimen komentar tidak dapat dianggap sebagai faktor utama yang memengaruhi tingkat engagement pengguna.

Temuan ini diperkuat oleh analisis korelasi yang menunjukkan hubungan sangat lemah antara sentimen dengan likes, reply, maupun engagement score. Sebaliknya, korelasi tinggi ditemukan antara likes dan reply terhadap engagement score, yang mengindikasikan bahwa engagement lebih dipengaruhi oleh dinamika interaksi kuantitatif antar pengguna dibandingkan oleh muatan emosional teks

d) Pola Interaksi Like dan Reply

Rasio *like* terhadap *reply* menunjukkan bahwa komentar positif cenderung memperoleh lebih banyak *like* tanpa memicu diskusi lanjutan. Sebaliknya, komentar netral relatif lebih sering menghasilkan balasan, sehingga berperan sebagai pemicu interaksi. Pola ini mencerminkan kecenderungan pengguna untuk mengekspresikan persetujuan secara pasif, sementara keterlibatan aktif melalui balasan terjadi pada konteks yang lebih informatif.

e) Pembahasan Singkat

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa sentimen komentar tidak memiliki pengaruh signifikan terhadap engagement pengguna. Engagement lebih

dipengaruhi oleh faktor non-linguistik, seperti karakteristik topik dan pola interaksi sosial pengguna. Oleh karena itu, analisis sentimen lebih relevan digunakan sebagai alat pemantauan opini publik dibandingkan sebagai indikator utama dalam memprediksi tingkat engagement.

5. KESIMPULAN DAN SARAN

Penelitian ini menunjukkan bahwa sentimen netral mendominasi komentar publik (59,41%), diikuti sentimen negatif (26,81%) dan positif (13,78%), yang mengindikasikan bahwa pengguna cenderung menyampaikan pendapat secara informatif dan lebih aktif mengekspresikan kritik dibandingkan apresiasi. Tingkat engagement pengguna secara umum tergolong rendah, dengan rata-rata likes sebesar 2,61, reply 0,30, dan engagement score 0,0036, serta mayoritas komentar tidak memicu interaksi lanjutan. Analisis statistik dan korelasi menunjukkan bahwa sentimen komentar tidak berpengaruh signifikan terhadap engagement, sementara keterlibatan pengguna lebih dipengaruhi oleh dinamika interaksi kuantitatif, khususnya hubungan antara likes dan reply. Dengan demikian, analisis sentimen lebih tepat digunakan untuk memantau opini publik dibandingkan sebagai indikator utama dalam memprediksi engagement pengguna.

DAFTAR REFERENSI

- Adriana, N. M. T. O., Suarjaya, I. M. A. D., & Githa, D. P. (2023). Analisis sentimen publik terhadap aksi demonstrasi di Indonesia menggunakan support vector machine dan random forest. *Decode: Jurnal Pendidikan Teknologi Informasi*, 3(2), 257–267. <https://doi.org/10.51454/decode.v3i2.187>
- Al Rasyid, R., Handayani, D., & Ningsih, U. (2024). Penerapan algoritma TF-IDF dan cosine similarity untuk query pencarian pada dataset destinasi wisata. *Jurnal Teknologi Informasi dan Komunikasi*, 8(1). <https://doi.org/10.35870/jti>
- Argueta, C., & Chen, Y.-S. (2014). Multi-lingual sentiment analysis of social data based on emotion-bearing patterns.
- Ciaramita, M., & Johnson, M. (2024). Supersense tagging of unknown nouns in WordNet.
- Khoirul, M., Hayati, U., & Nurdiawan, O. (2023). Analisis sentimen aplikasi BRImo pada ulasan pengguna di Google Play menggunakan algoritma Naive Bayes. *Jurnal Mahasiswa Teknik Informatika*, 7(1).
- Manoppo, M. R., Kolang, I. C., Nur Fiat, D. N., Mawara, R. M. C., Sumarno, A. D. P., Yusupa, A., & Tarigan, V. (2025). Analisis sentimen publik di media sosial terhadap kenaikan

- PPN 12% di Indonesia menggunakan IndoBERT. Jurnal Kecerdasan Buatan dan Teknologi Informasi, 4(2), 152–163. <https://doi.org/10.69916/jkbti.v4i2.322>
- Misrun, C. A., Haerani, E., Fikry, M., & Budianita, E. (2023). Analisis sentimen komentar YouTube terhadap Anies Baswedan sebagai bakal calon presiden 2024 menggunakan metode Naive Bayes classifier. Jurnal CoSciTech (Computer Science and Information Technology), 4(1), 207–215. <https://doi.org/10.37859/coscitech.v4i1.4790>
- Newburn, J. (2021). Keyword extraction in BERT-based models for reviewer system.
- Ningtyas, A. A., Solichin, A., & Pradana, R. (2023). Analisis sentimen komentar YouTube tentang prediksi resesi ekonomi tahun 2023 menggunakan algoritma Naive Bayes. 20(1).
- Nurfefia, K., & Sriani, S. (2024). Sentiment analysis of skincare products using the Naive Bayes method. Journal of Information Systems and Informatics, 6(3), 1663–1676. <https://doi.org/10.51519/journalisi.v6i3.817>
- Pavitha, N., Pungliya, V., Raut, A., Bhonsle, R., Purohit, A., Patel, A., & Shashidhar, R. (2022). Movie recommendation and sentiment analysis using machine learning. Global Transitions Proceedings, 3(1), 279–284. <https://doi.org/10.1016/j.gltp.2022.03.012>
- Sains, F., & Teknologi, D. (2022). Aplikasi sistem ekstraksi kata kunci berbahasa Indonesia menggunakan algoritma TextRank: Studi kasus data Wikipedia Indonesia.
- Tarigan, V. (2023). Pembuatan aplikasi data mining untuk memprediksi masa studi mahasiswa menggunakan algoritma Naive Bayes. 11(1).
- Yutika, C. H., Adiwijaya, A., & Faraby, S. A. (2020). Analisis sentimen berbasis aspek pada review komentar Instagram menggunakan TF-IDF dan Naïve Bayes. Jurnal Media Informatika Budidarma, 65(2), 432.
- Yutika, C. H., Adiwijaya, A., & Faraby, S. Al. (2021). Analisis sentimen berbasis aspek pada review Female Daily menggunakan TF-IDF dan Naïve Bayes. Jurnal Media Informatika Budidarma, 5(2), 422. <https://doi.org/10.30865/mib.v5i2.2845>