

Analisis Sentimen Publik terhadap Kualitas BBM Pertamina pada Komentar Youtube

Audy Trinugroho Listyawan^{1*}, Nova Tri Romadloni²

^{1,2}Fakultas Sains dan Teknologi Program Studi Informatika Universitas Muhammadiyah
Karanganyar

*Penulis Korespondensi : audytrin@gmail.com

ABSTRACT

Pertamina is a state-owned energy company that plays a vital role in the supply of petroleum products in Indonesia. The quality of the petroleum products it produces is often a topic of public discussion and generates a wide range of reactions on social media. This study aims to analyze public sentiment toward the quality of Pertamina's petroleum products based on user comments on the YouTube platform. A total of 3,775 comments were collected via web scraping and automatically labeled using the Indonesian-Roberta-Base-Sentiment-Classifer model, which categorized the comments into two groups: positive and negative. The analysis was conducted using RapidMiner software with two algorithms: Naïve Bayes and Support Vector Machine (SVM). The analysis stages included text preprocessing (tokenization, stopword removal, stemming) and TF-IDF weighting. Model evaluation was performed using 10-Fold Cross Validation with Accuracy and AUC metrics as indicators of model performance. The results of this study show that SVM performs better with an accuracy of 80.99% and an AUC of 0.76, compared to Naïve Bayes with an accuracy of 67.69% and an AUC of 0.50. Sentiment analysis reveals that the majority of comments are negative, indicating public dissatisfaction with the quality of Pertamina's fuel.

Article History

Received : 10-04-2025
Revised : 12-06-2026
Accepted : 18-06-2026

Keywords

Analisis Sentimen
Naive Bayes
Pertamina
RapidMiner
YouTube

ABSTRAK

Pertamina merupakan perusahaan energi milik negara yang memiliki peran penting dalam penyediaan Bahan Bakar Minyak (BBM) di Indonesia. Kualitas BBM yang dihasilkan sering menjadi perbincangan publik, dan menimbulkan beragam tanggapan di media sosial. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap kualitas BBM Pertamina berdasarkan komentar Pengguna pada platform YouTube. Sebanyak 3.775 komentar dikumpulkan melalui web scraping dan dilabeli otomatis menggunakan model indonesian-roberta-base-sentiment-classifier yang mengelompokkan komentar kedalam dua kategori yaitu positif dan negatif. Proses analisis dilakukan dengan perangkat lunak RapidMiner dengan menggunakan dua algoritma Naïve Bayes dan Support Vector Machine (SVM). Tahapan analisis meliputi preprocessing teks (tokenizing, stopword removal, stemming) serta pembobotan TF-IDF. Evaluasi model dilakukan menggunakan 10-Fold Cross Validation dengan metrik Accuracy dan AUC sebagai indikator performa model. Hasil penelitian ini menunjukkan SVM memiliki performa lebih baik dengan akurasi 80,99% dan AUC 0,76, dibandingkan Naïve Bayes dengan akurasi 67,69% dan AUC 0,50. Analisis Sentimen memperlihatkan bahwa mayoritas komentar bersentimen negatif, yang mengindikasikan ketidakpuasan publik terhadap kualitas BBM Pertamina.

PENDAHULUAN

Bahan Bakar Minyak (BBM) kini menjadi kebutuhan penting bagi hampir semua orang. Setiap hari kita memerlukan BBM untuk berbagai keperluan, terutama untuk transportasi mulai dari kendaraan pribadi, angkutan umum, hingga alat berat di sektor industri. Bahkan BBM memegang peranan yang sangat penting dalam aktivitas ekonomi. BBM juga merupakan faktor penting dalam menentukan perubahan harga bahan pokok atau inflasi di suatu negara [1]. Tanpa BBM, kegiatan masyarakat seperti bekerja, berdagang, atau bahkan bepergian tidak akan berjalan lancar.

Mengingat peran vital BBM bagi kehidupan masyarakat, penyediaan BBM yang berkualitas dan terjangkau menjadi tanggung jawab strategis yang diemban oleh PT Pertamina, sebagai

perusahaan milik negara (BUMN) yang berperan besar dalam memenuhi kebutuhan energi masyarakat Indonesia. Pertamina bertugas mengelola sumber daya minyak dan gas bumi, mengolahnya menjadi berbagai jenis bahan bakar, lalu menyalurkannya ke seluruh wilayah Indonesia. Perusahaan ini tidak hanya memproduksi BBM, tetapi juga memastikan agar masyarakat bisa mendapatkan bahan bakar dengan mudah, baik di kota besar maupun daerah terpencil [2]. Dengan jaringan SPBU yang tersebar luas, Pertamina berusaha menjaga agar pasokan BBM tetap aman dan terjangkau bagi semua kalangan. Keberadaan Pertamina sangat penting bagi kehidupan sehari-hari. Tanpa dukungan dari perusahaan ini, banyak aktivitas masyarakat dan kegiatan ekonomi yang bisa terhenti. Karena itu, Pertamina menjadi salah satu pilar utama dalam menjaga kelancaran mobilitas dan kebutuhan energi di Indonesia sehingga kualitas produk dan layanannya menjadi perhatian publik yang sangat besar.

Namun, peran strategis Pertamina sebagai pemasok energi nasional tidak lepas dari berbagai tantangan dan kritik publik. Di tengah meningkatnya ekspektasi dan tingkat kesadaran masyarakat terhadap transparansi dan integritas layanan publik terutama BUMN, Pertamina sering kali menjadi bahan perbincangan publik, terutama ketika terjadi kenaikan harga, kelangkaan pasokan, atau isu mengenai kualitas BBM [3]. Berbagai keluhan masyarakat terkait dugaan penurunan kualitas BBM, seperti kendaraan yang cepat rusak atau konsumsi bahan bakar yang boros, telah memicu diskusi luas di berbagai platform seperti media sosial.

Di era digital saat ini, Media sosial sebagai media komunikasi yang bersifat interaktif, opini publik tidak lagi terbatas pada ruang fisik, tetapi juga banyak disampaikan melalui media sosial [4]. Salah satu platform yang paling banyak digunakan untuk mengekspresikan pendapat adalah YouTube, di mana masyarakat dapat memberikan komentar langsung pada video yang berkaitan dengan isu energi dan kebijakan pemerintah. Komentar-komentar tersebut mencerminkan persepsi publik terhadap layanan dan produk Pertamina, yang jika dianalisis lebih lanjut, dapat memberikan wawasan penting mengenai tingkat kepuasan dan kepercayaan masyarakat terhadap perusahaan, padahal krisis berbasis persepsi seperti ini memiliki potensi yang besar dalam merusak reputasi secara lebih luas dan dalam, terlebih ketika masyarakat merasa seperti ditipu atau dikecewakan oleh institusi besar yang selama ini mereka percaya dan akan sulit untuk dipulihkan jika tidak ditangani secara strategis dan proaktif [3].

Untuk memahami persepsi publik secara sistematis dan objektif dari data media sosial yang sangat besar, diperlukan pendekatan analisis sentimen. Analisis sentimen merupakan bagian dari Natural Language Processing (NLP) yang digunakan untuk mengidentifikasi opini, emosi, atau sikap pengguna terhadap suatu topik [5][6]. Media sosial, khususnya YouTube, telah menjadi sumber data yang sangat kaya untuk analisis opini publik karena sifatnya yang interaktif dan memungkinkan pengguna menyampaikan pendapat secara spontan [7]. Sebagai salah satu platform video terbesar di dunia, kolom komentar YouTube mencerminkan reaksi emosional dan sudut pandang pengguna secara real-time [8]. Berbeda dengan media sosial lain seperti Twitter yang membatasi karakter, komentar YouTube cenderung lebih panjang dan detail sehingga memberikan konteks sentimen yang lebih mendalam [9]. Penelitian terdahulu menunjukkan bahwa analisis sentimen pada komentar YouTube dapat memberikan wawasan penting bagi organisasi dalam memahami persepsi publik dan mengidentifikasi isu-isu kritis yang memengaruhi reputasi [10][11].

Dalam mengimplementasikan analisis sentimen, pemilihan algoritma klasifikasi yang tepat menjadi kunci keberhasilan. Dalam perkembangan analisis sentimen berbahasa Indonesia, berbagai algoritma machine learning telah diterapkan untuk keperluan mengklasifikasikan sentimen teks, termasuk Naive Bayes dan Support Vector Machine (SVM). Naive Bayes menjadi salah satu algoritma yang paling populer karena kesederhanaan, efisiensi komputasi, dan kemampuannya menangani dataset besar dengan baik [6]. Penelitian oleh Sasmita et al. menunjukkan bahwa Naive

Bayes mencapai akurasi yang baik pada klasifikasi sentimen komentar berbahasa Indonesia dengan waktu pelatihan yang relatif cepat [12]. Di sisi lain, Support Vector Machine (SVM) dikenal memiliki akurasi tinggi dalam klasifikasi teks karena kemampuannya menemukan hyperplane yang ideal [13]. Beberapa studi komparatif menunjukkan bahwa SVM sering menghasilkan akurasi yang lebih baik dibandingkan dengan algoritma Naive Bayes, terutama pada dataset dengan dimensi tinggi [14][15]. Namun, performa kedua algoritma ini sangat bergantung pada teknik preprocessing teks dan metode pembobotan fitur yang digunakan. Salah satu metode pembobotan yang umum dipakai adalah Term Frequency–Inverse Document Frequency (TF-IDF), karena mampu menonjolkan kata-kata penting dan menghasilkan representasi numerik yang efektif untuk proses klasifikasi [16] [18].

Meskipun telah banyak penelitian analisis sentimen pada berbagai domain, penelitian khusus tentang sentimen publik terhadap kualitas BBM Pertamina sebagai BUMN strategis di platform media sosial masih sangat terbatas. Penelitian terdahulu lebih banyak berfokus pada analisis sentimen produk konsumen [17], layanan transportasi [15] atau isu-isu sosial-politik [18], sementara sektor energi yang memiliki implikasi strategis terhadap ekonomi nasional belum banyak dieksplorasi. Padahal, BUMN seperti Pertamina memiliki peran penting dalam menjaga ketahanan energi nasional dan kesejahteraan masyarakat [19]. Analisis sentimen terhadap BUMN strategis seperti PLN telah menunjukkan bahwa opini publik di media sosial dapat menjadi indikator penting untuk evaluasi layanan [20]. Namun, penggunaan dua algoritma seperti Naive Bayes dan SVM untuk menganalisis sentimen pada komentar YouTube berbahasa Indonesia yang berkaitan dengan produk BUMN masih jarang dilakukan secara mendalam.

Oleh karena itu, penelitian ini dilakukan untuk menjawab kebutuhan tersebut dengan menganalisis sentimen masyarakat terhadap kualitas BBM Pertamina melalui komentar yang ada di YouTube. Pada penelitian ini akan digunakan dua algoritma, yaitu Naive Bayes dan SVM, yang diimplementasikan melalui platform RapidMiner. Melalui analisis ini, diharapkan dapat diperoleh gambaran yang lebih jelas mengenai bagaimana publik menilai kualitas BBM Pertamina. Hasilnya juga diharapkan dapat menjadi masukan berharga bagi Pertamina dalam meningkatkan mutu produk, pelayanan, serta merumuskan strategi komunikasi yang lebih adaptif dan sesuai dengan perkembangan opini masyarakat.

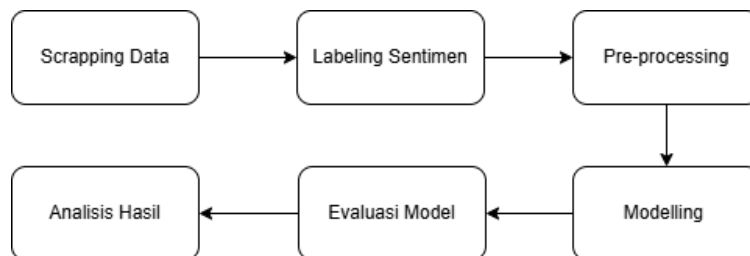
METODE

Penelitian yang saya lakukan ini menggunakan pendekatan kuantitatif berbasis text mining untuk melakukan analisis sentimen publik terhadap kualitas BBM Pertamina melalui komentar pengguna di platform YouTube. Metode penelitian disusun dalam beberapa tahap, mulai dari pengumpulan data, pelabelan sentimen otomatis, praproses data (*text preprocessing*), pemodelan dengan menggunakan dua algoritma yaitu *Naive Bayes* dan *Support Vector Machine (SVM)*, hingga evaluasi model menggunakan *cross validation*.

Alur Penelitian

Secara umum, tahapan penelitian ini dapat dijelaskan melalui diagram alur yang digambarkan pada Gambar 1. Tahapan penelitian ini diawali dengan pengumpulan data yang menggunakan teknik web scraping pada kolom komentar video YouTube yang membahas topik kualitas BBM Pertamina. Data yang diperoleh berupa ribuan komentar mentah kemudian melalui proses pembersihan data (*data cleaning*) untuk menghapus duplikasi, tanda baca, dan juga karakter yang tidak relevan. Selanjutnya dilakukan pelabelan otomatis

yang menggunakan model `w11wo/indonesian-roberta-base-sentiment-classifier` untuk mengelompokkan komentar ke dalam dua kategori sentimen, yaitu positif dan negatif.



Gambar 1. Alur Penelitian.

Tahap berikutnya adalah pra-pemrosesan teks (text preprocessing) yang mencakup proses tokenizing, transform case, stopword removal, stemming, dan pembobotan kata menggunakan metode TF-IDF agar data dapat diubah menjadi bentuk numerik yang dapat diolah oleh algoritma pembelajaran mesin. Setelah tahap praproses selesai, dilakukan pembangunan model klasifikasi menggunakan dua algoritma, yaitu Naïve Bayes dan Support Vector Machine (SVM), yang diimplementasikan dalam perangkat lunak RapidMiner Studio. Selanjutnya dilakukan evaluasi model dengan skema 10-Fold Cross Validation menggunakan metrik Accuracy dan Area Under Curve (AUC) untuk menilai performa model secara kuantitatif. Tahap terakhir adalah analisis hasil dan interpretasi sentimen, yaitu menilai distribusi komentar positif dan negatif untuk mengetahui persepsi publik terhadap kualitas BBM Pertamina serta menarik kesimpulan berdasarkan performa dari kedua algoritma yang dipakai.

Pengumpulan Data

Data penelitian dikumpulkan dari kolom komentar pada beberapa video resmi di platform YouTube yang membahas topik “Pertamina” dan “BBM”. Sumber data dipilih secara purposive berdasarkan relevansi konten video terhadap isu kualitas bahan bakar dan kebijakan energi nasional. Total data yang berhasil dikumpulkan sebanyak 3.775 komentar dari berbagai video pada kanal resmi dan media pemberitaan nasional yang menayangkan informasi terkait Pertamina.

Proses pengambilan data dilakukan secara otomatis melalui metode web scraping yang menggunakan bahasa pemrograman Python yang dijalankan di Google Colab. Penggunaan Colab memudahkan proses karena mendukung eksekusi kode berbasis cloud tanpa perlu instalasi lokal, serta kompatibel dengan berbagai pustaka analisis teks. Beberapa pustaka utama yang digunakan yaitu `youtube-comment-downloader` untuk mengekstraksi komentar dari URL video YouTube, serta `pandas` untuk mengelola dan menyimpan hasil scraping dalam format CSV agar dapat diolah lebih lanjut. Cuplikan kode yang digunakan dalam proses pengumpulan data, ditunjukkan pada Gambar 2.

Pelabelan Data Sentimen

Setelah proses pengumpulan data dilakukan, tahap selanjutnya adalah menentukan label sentimen untuk setiap komentar agar dapat diketahui apakah komentar tersebut bernada positif atau negatif. Mengingat jumlah data yang cukup besar yaitu 3.775 komentar, pelabelan manual dinilai kurang efisien dan memakan waktu lama, sehingga penelitian ini menggunakan metode pelabelan otomatis (automated sentiment labeling) menggunakan model `"w11wo/indonesian-roberta-base-sentiment-classifier"` yang tersedia pada pustaka Transformers milik Hugging Face. Model ini berbasis RoBERTa (Robustly Optimized BERT Pretraining Approach) yang telah dilatih secara

husus pada korpus teks berbahasa Indonesia, sehingga mampu memahami konteks kalimat dengan lebih akurat dibandingkan pendekatan berbasis leksikon tradisional. Model ini secara otomatis mengklasifikasikan setiap komentar menjadi kategori positif atau negatif berdasarkan makna semantik dari isi teks yang dianalisis. Meskipun demikian, pelabelan otomatis tetap memiliki keterbatasan, terutama pada teks informal yang mengandung kata tidak baku, singkatan, atau slang yang umum ditemukan pada komentar YouTube.

```
!pip install youtube-comment-downloader pandas

from youtube_comment_downloader import YoutubeCommentDownloader
import pandas as pd

# Masukkan link video YouTube di sini
video_url = "https://youtube.com/shorts/svEtAA8PyGg?si=1ZhPko-5rX9C2j8a"

downloader = YoutubeCommentDownloader()
comments = []
for comment in downloader.get_comments_from_url(video_url, sort_by=1): # Changed 'top'
    to 1
    comments.append({
        'author': comment['author'],
        'text': comment['text'],
        'time': comment['time'],
        'votes': comment['votes']
    })

# Simpan ke excel
df = pd.DataFrame(comments)
df.to_excel("Scrapping_BBM.xlsx", index=False)
print("Selesai! Komentar disimpan di: Scrapping_BBM.xlsx")
df.head()
```

Gambar 2. Kode proses pengumpulan data

Perlu ditekankan bahwa RoBERTa dalam penelitian ini berperan murni sebagai alat bantu pelabelan otomatis, bukan sebagai model klasifikasi akhir. Model utama yang dibangun dan diuji kemampuannya adalah Naive Bayes dan SVM, yang dilatih untuk mempelajari pola linguistik dari data komentar secara mandiri menggunakan pendekatan machine learning berbasis TF-IDF. Dengan demikian, akurasi yang dihasilkan mencerminkan kemampuan masing-masing algoritma dalam mengenali pola sentimen dari teks, bukan kemampuan meniru keputusan RoBERTa.

Seluruh proses pelabelan dijalankan menggunakan bahasa pemrograman Python pada lingkungan Google Colaboratory (Colab). Pemilihan Colab dilakukan karena platform ini menyediakan sumber daya komputasi berbasis cloud yang efisien serta mendukung pustaka pemrosesan bahasa alami terkini. Sebelum pelabelan dilakukan, data komentar hasil web scraping dibersihkan terlebih dahulu dengan menghapus nilai kosong (missing values) dan karakter non-teks agar hasil analisis lebih akurat dan stabil. Cuplikan kode Python yang digunakan dalam proses pelabelan otomatis ditunjukkan pada Gambar 3.

```
!pip install transformers torch pandas openpyxl
# Import library
import pandas as pd
from transformers import pipeline
from google.colab import files

# === 1. Upload File Excel ===
uploaded = files.upload()

# === 2. Baca file ===
df = pd.read_excel("Scrapping_BBM.xlsx")

print("Kolom yang tersedia:", df.columns.tolist())
print("\nJumlah data sebelum pembersihan:", len(df))

# Hapus baris kosong di kolom 'text'
df = df.dropna(subset=["text"])
print("Jumlah data setelah pembersihan:", len(df))

# === 3. Load Model Bahasa Indonesia (RoBERTa) ===
sentiment_model = pipeline(
    "sentiment-analysis",
    model="wllwo/indonesian-roberta-base-sentiment-classifier")

# === 4. Fungsi klasifikasi sentimen ===
def indo_sentiment(text):
    try:
        result = sentiment_model(str(text)[:512])[0]
        label = result['label'].upper()
        if "POSITIVE" in label:
            return "Positif"
        else:
            return "Negatif"
    except:
        return "Negatif"

# === 5. Tambahkan kolom Labeling Baru di sebelah kolom lama ===
df["Labeling Baru"] = df["text"].apply(indo_sentiment)

# === 6. Simpan hasil ke file baru ===
output_name = "Scrapping_BBM_Labeling_Baru.xlsx"
df.to_excel(output_name, index=False)
files.download(output_name)

# Lihat 10 data pertama
df.head(10)
```

Gambar 3. Kode Python yang digunakan dalam proses pelabelan otomatis

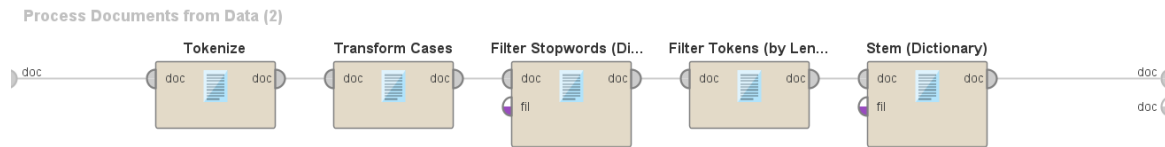
Praprocess Data Teks

Tahapan praproses data teks merupakan salah satu langkah penting dalam analisis sentimen karena menentukan kualitas data yang akan digunakan untuk pelatihan model klasifikasi. Proses ini bertujuan untuk memfilter, membersihkan, menormalisasi, dan menyiapkan data teks mentah agar dapat diubah menjadi representasi numerik yang nantinya dapat diproses oleh algoritma machine learning.

Dalam penelitian ini, praproses dilakukan menggunakan operator Process Documents from Data pada RapidMiner Studio, yang menyediakan serangkaian proses otomatis untuk mengolah teks dalam format dokumen. Proses text preprocessing dilakukan melalui lima tahapan berurutan dalam RapidMiner. Tahapan-tahapan yang dilakukan meliputi:

Tahapan pertama adalah tokenize, yaitu memecah kalimat menjadi unit-unit kata individual atau token. Tahapan kedua adalah transform case yang mengubah seluruh teks menjadi huruf kecil (lowercase). Tahapan ketiga adalah filter stopwords yang menghapus kata-kata umum yang tidak bermakna dalam analisis sentimen, seperti "yang", "dan", "ke", dan "di". Tahapan keempat adalah

filter tokens by length yang menghapus token dengan panjang kurang dari tiga karakter untuk mengeliminasi singkatan dan kata kurang informatif. Tahapan kelima adalah stemming menggunakan stem dictionary untuk mengembalikan kata berimbuhan ke bentuk dasarnya.



Gambar 4. Process Documents from Data.

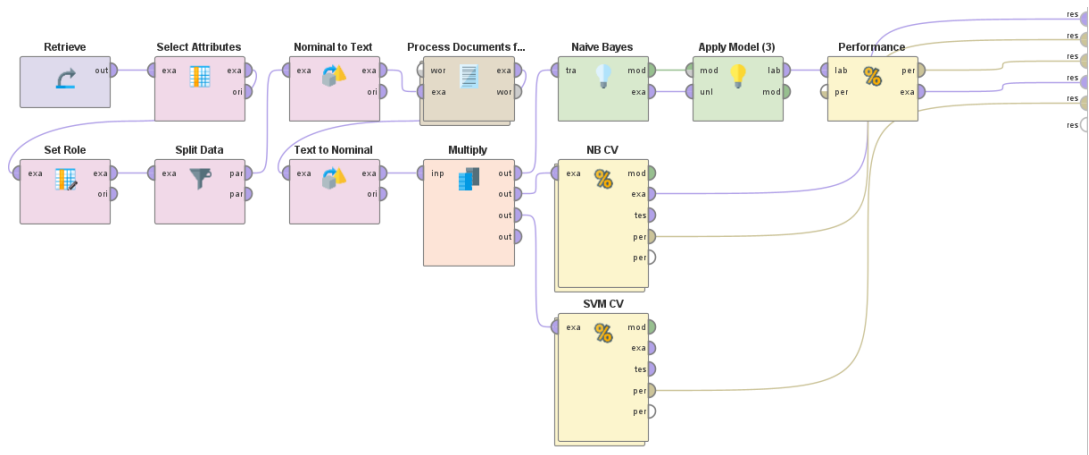
Melalui tahapan ini, teks komentar yang awalnya tidak terstruktur diubah menjadi data yang lebih bersih, seragam, dan bermakna, sehingga dapat diolah lebih lanjut pada tahap pemodelan klasifikasi sentimen yang menggunakan algoritma Naive Bayes dan SVM.

Pemodelan dan Evaluasi

Algoritma yang digunakan adalah Naive Bayes dan Support Vector Machine (SVM). yang diimplementasikan menggunakan perangkat lunak RapidMiner Studio. Proses pembangunan model dilakukan melalui beberapa tahapan operator, seperti terlihat pada Gambar 4, yang menunjukkan keseluruhan alur pemodelan dan evaluasi yang saya lakukan.

Implementasi pemodelan menggunakan RapidMiner Studio dimulai dengan retrieve dataset untuk memanggil data yang telah dilabeli, dilanjutkan dengan select attributes untuk memilih kolom “text” (komentar) dan “Labeling_Baru” (label sentimen). Operator set role kemudian menetapkan Labeling_Baru sebagai label dan text sebagai regular attribute. Tahap preprocessing dimulai dengan konversi nominal to text, kemudian process documents from data untuk melakukan tokenisasi, filtering stopwords, stemming, dan pembobotan TF-IDF yang otomatis ada di operator “Process Documents From Data”. Hasil preprocessing dikonversi kembali ke format nominal menggunakan text to nominal dan digandakan melalui operator multiply untuk memungkinkan evaluasi dua algoritma secara simultan.

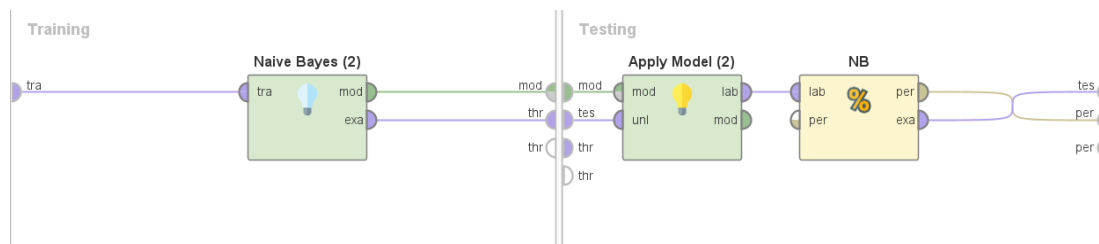
Tahap pemodelan menerapkan Naive Bayes dan Support Vector Machine (SVM) secara paralel pada data hasil preprocessing. Evaluasi dilakukan menggunakan 10-fold cross validation untuk memastikan hasil yang stabil dan tidak bias. Dalam setiap fold, model dilatih pada training set dan diuji pada testing set, dengan operator apply model untuk penerapan model dan operator performance (binomial classification) untuk menghitung accuracy dan AUC sebagai indikator kinerja. Proses ini diulang sepuluh kali dan hasil akhir berupa rata-rata metrik dari seluruh fold. Hasil alur proses RapidMiner dapat dilihat pada Gambar 5, yang menggambarkan hubungan antaroperator dan alur data dari proses pelatihan hingga evaluasi model.



Gambar 5. Alur proses Rapid Miner

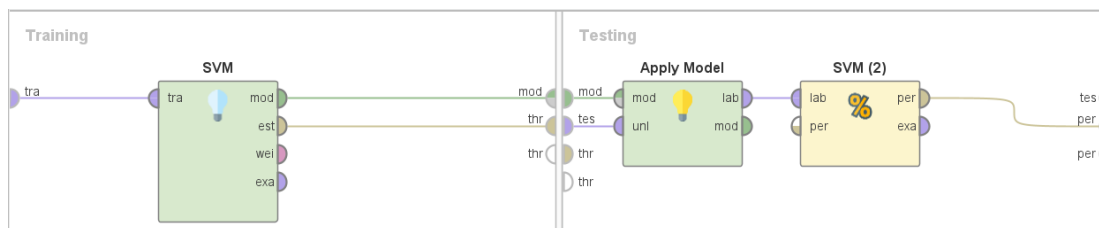
Kedua model diuji menggunakan skema 10-Fold Cross Validation, yaitu dengan membagi seluruh data menjadi sepuluh bagian secara acak namun tetap menjaga proporsi kelas yang seimbang (stratified). Proses pengujian dilakukan sebanyak sepuluh kali, di mana pada setiap putaran, sembilan bagian data digunakan untuk melatih model dan satu bagian sisanya digunakan untuk menguji performa model. Hasil evaluasi diukur berdasarkan akurasi (Accuracy) dan Area Under Curve (AUC).

Cross Validation Naïve Bayes:



Gambar 6. Cross Validation Naïve Bayes

Cross Validation Support Vector Machine :



Gambar 7. Cross Validation Support Vector Machine

Proses evaluasi menggunakan metrik akurasi yang didefinisikan dengan Persamaan (1)

$$Accuracy = \frac{TP+TN}{TP+FP+FN+TN} \quad (1)$$

Dimana :

TP = *True Positive (TP)* yaitu jumlah komentar positif yang berhasil diprediksi dengan benar

TN = *True Negative (TN)* yaitu jumlah komentar negatif yang diprediksi dengan tepat.

FP = *False Positive (FP)* yaitu komentar negatif yang salah diprediksi sebagai positif.

FN = *False Negative (FN)* yaitu komentar positif yang salah diprediksi sebagai negatif [21].

Algoritma ini bekerja berdasarkan teorema Bayes dengan asumsi bahwa setiap kata (fitur) dalam teks bersifat independen atau tidak saling mempengaruhi satu sama lain, sebagaimana diformulasikan dalam Persamaan (2):

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)} \quad (2)$$

Keterangan:

$P(C|X)$ $P(C|X)$ $P(C|X)$ = probabilitas hipotesis kelas CCC diberikan data XXX,

$P(X|C)$ $P(X|C)$ $P(X|C)$ = probabilitas kemunculan data XXX dalam kelas CCC,

$P(C)$ $P(C)$ $P(C)$ = probabilitas awal kelas,

$P(X)$ $P(X)$ $P(X)$ = probabilitas data secara keseluruhan [21].

HASIL DAN PEMBAHASAN

Penelitian ini bertujuan guna menganalisis sentimen publik terhadap kualitas BBM Pertamina berdasarkan komentar pengguna di platform YouTube menggunakan pendekatan text mining. Tahapan hasil penelitian disusun secara berurutan berdasarkan proses yang dilakukan, yaitu: pengumpulan data (scraping), pelabelan otomatis, praproses teks, pemodelan menggunakan dua algoritma Naive Bayes dan SVM, serta evaluasi performa model. Setiap tahap menghasilkan data dan temuan yang saling berkaitan untuk menjawab tujuan penelitian.

Hasil Pengumpulan Data (Scraping Data)

Data dikumpulkan dari kolom komentar pada video YouTube yang membahas topik “Pertamina” dan “BBM”. Proses pengambilan data dilakukan menggunakan bahasa pemrograman Python melalui pustaka youtube-comment-downloader yang dijalankan di Google Colaboratory. Dari hasil scraping, diperoleh total 3.775 komentar dari beberapa video yang relevan dengan isu kualitas bahan bakar dan kebijakan energi nasional.

Data disimpan dalam format excel dengan beberapa kolom utama, yaitu:

- author : nama pengguna yang memberikan komentar,
- text : isi dari komentar,
- time : waktu komentar diposting, dan
- votes : jumlah suka (likes).

Contoh hasil data awal disajikan pada Gambar 8.

	author	text	time	votes
0	@Raviel-v	Kocak setelah shell swasta berhenti jualan , k...	33 minutes ago	0
1	@sahafoficial9288	Menurut saya cuci otak dan hati para pengurus ...	1 hour ago	0
2	@LFarizal-x3t	Cepat pecat Bahlil mau maunya kita Nerima di g...	1 hour ago	0
3	@MujahidHarris	Wajah dan kelakuan si bahlil binin enek	1 hour ago	0
4	@adityarafaal8738	Udah harga bbm paling mahal ke2 di ASEAN kual...	2 hours ago	0

Gambar 8. Contoh Hasil Scraping Data.

Data mentah ini kemudian menjadi bahan dasar untuk tahap pelabelan sentimen otomatis.

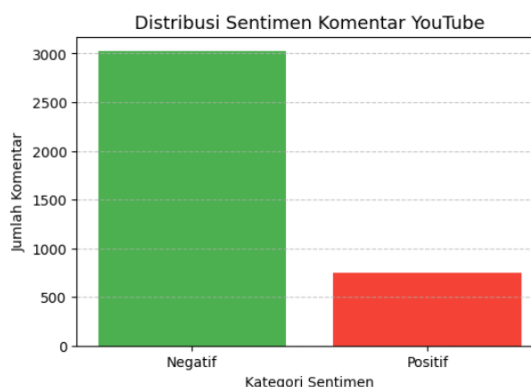
Hasil Pelabelan Otomatis

Tahap pelabelan dilakukan menggunakan model bahasa Indonesia `w11wo/indonesian-roberta-base-sentiment-classifier` dari pustaka Hugging Face Transformers. Model ini dikonfigurasi untuk memberikan dua label utama, yaitu Positif dan Negatif. Proses labeling dilakukan otomatis di Google Colaboratory dengan menggunakan fungsi `pipeline("sentiment-analysis")`. Sebelum dilakukan labeling, dilakukan pembersihan data dengan menghapus nilai kosong dan karakter non-alfabet agar model dapat membaca teks dengan benar. Hasil labeling menunjukkan distribusi sentimen seperti yang ditunjukkan pada Gambar 9.

	author	text	time	votes	label	Labeling Baru
0	@Raviel-v	Kocak setelah shell swasta berhenti jualan , k...	19 minutes ago	0	Negative	Negatif
1	@sahafofficial9288	Menurut saya cuci otak dan hati para pengurus ...	1 hour ago	0	Negative	Negatif
2	@LFarizal-x3t	Cepat pecat Bahlil mau maunya kita Nerima di g...	1 hour ago	0	Negative	Negatif
3	@MujahidHarris	Wajah dan kelakuan si bahlil binin enek	1 hour ago	0	Negative	Negatif
4	@adityarafeal8738	Udah harga BBM paling mahal ke2 di ASEAN kualiti...	1 hour ago	0	Negative	Negatif
5	@tiktube7278	Orang begituan jadi menteri, rusak negara ini. \n...	2 hours ago	0	Negative	Negatif
6	@anejr4	Udah mahal kualitas nya buruk lg 🤬	2 hours ago	0	Negative	Negatif
7	@ahmaddewantara4512	Pecat saja menteri ESDM nya Nggak becus kerja ...	4 hours ago	0	Negative	Negatif
8	@modifikasiroda255	Bang katanya Pertamina bakal kasih uang 7 juta ...	5 hours ago	0	Negative	Negatif
9	@musniad4600	gda ya org yg lebih baik jadi menteri di bandi...	6 hours ago	0	Negative	Negatif

Gambar 9. Contoh Hasil Labelling Otomatis

Pada gambar hasil terlihat ada 2 label yaitu “label” dan “Labeling_baru”. “label” saya tambahkan dengan cara manual, sementara “Labeling_baru” menggunakan python dengan model RoBERTa yang juga terlihat berbeda dengan kolom “label” jika dilihat full dari data pertama hingga data terakhir.



Gambar 10. Visualisasi hasil labeling otomatis

Visualisasi hasil labeling dapat dilihat pada Gambar 10, yang memperlihatkan bahwa mayoritas komentar bersifat negatif dengan jumlah menyentuh angka 3000, menunjukkan persepsi publik yang cukup buruk terhadap Pertamina.

Hasil Preprocessing

Tahap preprocessing dilakukan untuk mengubah teks komentar mentah menjadi teks yang lebih bersih, terstruktur, dan siap digunakan pada proses analisis sentimen. Preprocessing dilakukan

didalam rapid miner memnggunakan operator “Process Documents from Data”, serta beberapa tahapan didalamnya seperti tokenizing, transform cases, stopword removal, filter tokens (by length) dan yang terakhir ada stemming. Proses ini bertujuan mengambil kata , menghilangkan karakter-karakter yang tidak relevan, menyederhanakan bentuk kata, dan memastikan setiap komentar berada dalam format yang konsisten.

Sebagai contoh, Tabel 1 berikut menunjukkan beberapa sampel komentar sebelum dan sesudah dilakukan preprocessing. Terlihat bahwa teks awal yang masih mengandung huruf kapital, tanda baca, kata tidak baku, serta kata-kata umum yang tidak memiliki nilai analitis dibersihkan dan diubah ke bentuk dasar. Hasil ini menjadi dasar bagi tahap ekstraksi fitur TF-IDF pada pemodelan di RapidMiner.

Tabel 1. Hasil Preprocessing

No	Teks Awal	Teks Hasil Preprocessing
1	<i>“Bang katanya pertamina bakal kasih uang 7 juta kalo posting hal baik dari pertamina bener ga si”</i>	bang pertamina kasih uang juta kalo posting pertamina bener
2	<i>“Bahlii MENDUNIA”</i>	bahlil dunia
3	<i>“Tenang..bentar lagi bubar indonesia”</i>	tenang bentar bubar indonesia
4	<i>“Jual barang busuk dengan harga tinggi.. Pertamina kalau ngga monopoli bakal kalah saing kayak pos Indonesia 🏆”</i>	jual barang busuk tidak pertamina tidak monopoli kalah saing iya pos indonesia
5	<i>“Ini si bahlil lulusan apa sih sekolahnya? Kyak yg gk prnah bljr kimia sama sekali.”</i>	bahlil lulusan sih iya iya prnah bljr kimia

Dari hasil preprocessing di atas, terlihat bahwa komentar menjadi lebih ringkas dan relevan untuk dianalisis oleh algoritma klasifikasi. Proses ini juga menghilangkan noise seperti tanda baca, emoji, huruf kapital, serta kata-kata yang tidak berkontribusi pada penentuan sentimen. Oleh karena itu, teks hasil preprocessing menjadi lebih representatif dan memudahkan pembentukan fitur numerik menggunakan metode TF-IDF pada tahap pemodelan selanjutnya. Meskipun demikian, hasil preprocessing tidak selalu menghasilkan bentuk kata yang sepenuhnya akurat, terutama pada komentar yang mengandung bahasa slang, singkatan, atau kesalahan penulisan. Beberapa kata dapat mengalami overstemming ataupun understemming karena keterbatasan kamus dan aturan linguistik pada algoritma Sastrawi. Namun, secara keseluruhan proses ini tetap memberikan peningkatan signifikan terhadap kualitas data sehingga model dapat belajar pola sentimen dengan lebih efektif.

Hasil Pemodelan dan Evaluasi

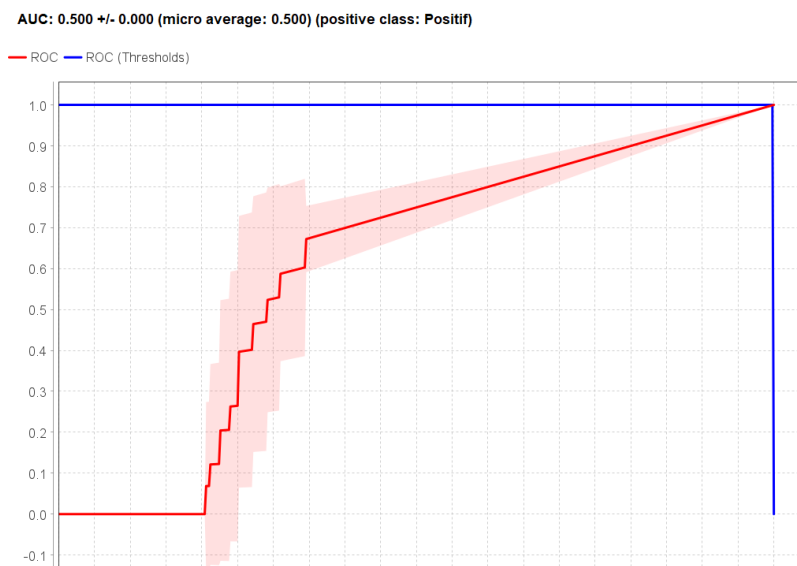
Tahap pemodelan dilakukan dengan dua algoritma: Naive Bayes dan Support Vector Machine (SVM). Kedua model diuji menggunakan 10-Fold Cross Validation agar hasil pengujian tidak bias dan lebih stabil. Evaluasi performa model dalam penelitian ini mengukur kemampuan Naive Bayes dan SVM dalam mengenali pola sentimen secara independen dari data yang telah

dilabeli, sehingga hasil akurasi mencerminkan kekuatan masing-masing algoritma dalam memahami teks komentar berbahasa Indonesia. Hasil pengujian pada penelitian ini ditunjukkan pada Gambar 11.

accuracy: 67.69% +/- 3.75% (micro average: 67.69%)			
	true Negatif	true Positif	class precision
pred. Negatif	1035	134	88.54%
pred. Positif	476	243	33.80%
class recall	68.50%	64.46%	

Gambar 11. Hasil Accuracy Naïve Bayes

Model Naive Bayes mencapai akurasi 67.69% ($\pm 3.75\%$) dengan confusion matrix yang menunjukkan 1.035 True Negative, 243 True Positive, 476 False Positive, dan 134 False Negative (Gambar 9). Class recall untuk sentimen negatif (68.50%) dan positif (64.46%) menunjukkan performa yang relatif seimbang. Namun, class precision menunjukkan perbedaan signifikan dengan nilai 88.54% untuk negatif dan hanya 33.80% untuk positif. Tingginya False Positive (476 kasus) mengindikasikan kecenderungan model untuk overpredict kelas positif, yang dapat disebabkan oleh ketidakseimbangan kelas atau kompleksitas linguistik komentar YouTube berbahasa Indonesia.



Gambar 12. Hasil AUC Naïve Bayes

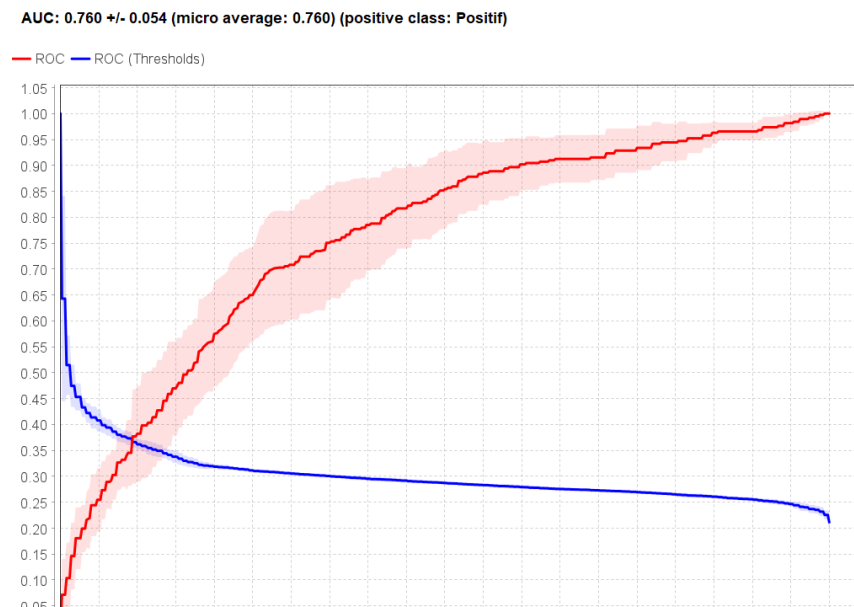
Gambar 12 menampilkan kurva ROC dengan nilai AUC sebesar 0.500 (± 0.000), setara dengan performa random classifier. Kurva yang hampir berimpitan dengan garis diagonal menunjukkan bahwa model tidak memiliki kemampuan diskriminatif dalam memisahkan kelas positif dan negatif secara probabilistik. Kontras antara akurasi (67,69%) dan AUC rendah (0.500). Kondisi ini dapat disebabkan oleh beberapa faktor, antara lain distribusi data yang tidak seimbang antara kelas positif dan negatif, proses preprocessing yang belum sepenuhnya optimal, serta karakteristik algoritma Naive Bayes yang mengasumsikan independensi antar fitur. Pada teks komentar yang kompleks dan banyak mengandung gaya bahasa yang informal, asumsi ini menjadi sulit terpenuhi, sehingga model kehilangan kemampuan untuk mengenali konteks yang sebenarnya.

accuracy: 80.99% +/- 1.12% (micro average: 80.99%)

	true Negatif	true Positif	class precision
pred. Negatif	1497	345	81.27%
pred. Positif	14	32	69.57%
class recall	99.07%	8.49%	

Gambar 13. Hasil Accuracy Support Vector Machine

Model Support Vector Machine (SVM) mencapai akurasi 80.99% ($\pm 1.12\%$), meningkat 13.30 poin persentase dibandingkan Naive Bayes dengan stabilitas yang lebih baik. Confusion matrix menunjukkan 1.497 True Negative, 32 True Positive, 14 False Positive, dan 345 False Negative. Class recall untuk sentimen negatif sangat tinggi (99.07%), namun recall untuk sentimen positif sangat rendah (8.49%), menunjukkan bahwa model sangat konservatif dalam memprediksi kelas positif. Class precision mencapai 81.27% untuk negatif dan 69.57% untuk positif, jauh lebih baik dibandingkan Naive Bayes. Pola ini mengindikasikan bahwa SVM cenderung severely underpredict kelas minoritas untuk memaksimalkan akurasi keseluruhan, dengan hanya membuat 46 prediksi positif dari total 1.888 sampel.



Gambar 14. Hasil AUC Support Vector Machine

Nilai AUC SVM sebesar 0.760 (± 0.054), tergolong “acceptable” hingga “good” dan mengindikasikan kemampuan diskriminatif yang solid. Kurva ROC membentuk lengkungan signifikan ke kiri atas, menunjukkan bahwa model mampu membedakan sentimen positif dan negatif secara probabilistik dengan baik. Pada Gambar 14, tampak dua kurva ROC yang ditampilkan oleh RapidMiner, yaitu ROC (kurva utama yang dihitung berdasarkan probabilitas prediksi model pada setiap fold) dan ROC (Threshold) (kurva yang menunjukkan hubungan antara nilai threshold keputusan dengan True Positive Rate dan False Positive Rate). Kedua kurva ini saling melengkapi: ROC digunakan untuk menghitung nilai AUC keseluruhan, sementara ROC (Threshold) membantu mengidentifikasi nilai threshold optimal yang memaksimalkan performa klasifikasi. Nilai AUC 0.760 berarti model memiliki probabilitas 76% untuk memberikan skor lebih tinggi pada komentar positif dibandingkan negatif ketika diberikan sepasang sampel acak.

Dari beberapa gambar di atas, dapat disimpulkan bahwa algoritma SVM memiliki performa yang lebih tinggi dan lebih baik dibandingkan Naive Bayes, baik pada akurasi maupun nilai AUC. Hal ini karena SVM mampu membedakan kelas dengan batas margin yang lebih optimal dibandingkan dengan pendekatan probabilistik sederhana pada Naive Bayes.

KESIMPULAN

Berdasarkan hasil penelitian diatas, dapat disimpulkan bahwa algoritma Support Vector Machine (SVM) memberikan hasil performa yang lebih baik dibandingkan dengan algoritma Naive Bayes dalam menganalisis sentimen komentar publik berbahasa Indonesia terhadap kualitas BBM Pertamina di platform YouTube. Model SVM memperoleh akurasi sebesar 80,99% dan nilai AUC sebesar 0,760, sedangkan Naive Bayes menghasilkan akurasi 67.69% dengan AUC 0,500. Perbedaan ini menunjukkan bahwa algoritma SVM dinilai lebih efektif dalam mengklasifikasikan teks berdimensi tinggi yang dihasilkan dari proses representasi TF-IDF. Analisis distribusi sentimen juga menunjukkan bahwa mayoritas komentar publik bersifat negatif, menandakan adanya persepsi kurang puas terhadap kualitas dan pelayanan BBM Pertamina. Penelitian selanjutnya disarankan untuk menguji algoritma lain seperti Random Forest atau Deep Learning guna memperoleh hasil klasifikasi yang lebih optimal dan akurat terhadap data teks berbahasa Indonesia.

DAFTAR PUSTAKA

- [1] A. Amelia, L. N. Hayati, and H. Darwis, "Analisis Sentimen Masyarakat Terhadap Sistem Pembayaran Mypertamina dengan Metode Random Forest, SVM, dan Naive Bayes," *Literatur Informatika & Komputer*, vol. 1, no. 1, pp. 28–44, 2024, doi: 10.33096/linier.v1i1.2269.
- [2] R. A. Softina, F. M. Amin, and N. Wahyudi, "Analisis Faktor yang Mempengaruhi Innovation Resistance dan Intention to Use Terhadap Penerapan Pembayaran Non Tunai," *JURNAL SISTEM INFORMASI BISNIS*, vol. 12, no. 1, pp. 26–35, Sep. 2022, doi: 10.21456/vol12iss1pp26-35.
- [3] O. : Arinda, S. Efendi, L. Ladzidza Vidiananda, F. P. Utami, A. Fitriyah, and D. I. Mardiningsih, "PT. Media Akademik Publisher REPUTASI PERTAMINA DAN MANAJEMEN KRISIS PASCA ISU PENGOPLOSAN BBM DI INDONESIA," *JMA*, vol. 3, no. 6, pp. 3031–5220, 2025, doi: 10.62281.
- [4] E. Nazalatus Sa, adiyah Sy, F. Yunanto, and R. Kasanova, "Peran Media Sosial Dalam Pembentukan Opini Mahasiswa Fkip Universitas Madura: Analisis Interaksi Di Era Digital," *Jurnal Cendekia Ilmiah*, vol. 3, no. 6, 2024.
- [5] W. Jurnal, J. Husna, M. Dyah Anggraini Widirahayu, S. Tinggi Multi Media, and J. Magelang Km, "Pemodelan Topik dan Analisis Sentimen pada 'Voices of History: 50 Iconic Speeches' Menggunakan Pendekatan Natural Language Processing," *Jurnal Ilmiah Manajemen Informasi dan Komunikasi*, vol. 8, no. 1, 2024, doi: 10.20895/jimik.v8i1.page.
- [6] Rizal Chandra Rivaldi and T.D. Wismarini, "Analisis Sentimen Pada Ulasan Produk Dengan Metode Natural Language Processing (NLP)," *Elkom: Jurnal Elektronika dan Komputer*, vol. 17, no. 1, pp. 120–128, Jul. 2024, doi: 10.51903/elkom.v17i1.1680.
- [7] P. Studi Ilmu Komunikasi, "TREN PENGGUNAAN MEDIA SOSIAL SELAMA PANDEMI DI INDONESIA," 2020.
- [8] S. Suwanto, A. Muzaki, and M. Muhtarom, "Pemanfaatan Media YouTube sebagai Media Pembelajaran pada Siswa Kelas XII MIPA di SMA Negeri 1 Tawangsari," *Media Penelitian Pendidikan : Jurnal Penelitian dalam Bidang Pendidikan dan Pengajaran*, vol. 15, no. 1, pp. 26–30, Jun. 2021, doi: 10.26877/mpp.v15i1.7531.
- [9] P. Tutiasri, K. Lamionto, and K. Nazri, "Pemanfaatan Youtube Sebagai Media Pembelajaran Bagi Mahasiswa di Tengah Pandemi Covid-19," vol. 2, 2020.
- [10] T. Muhayat, A. Fauzi, and D. J. Indra, "Analisis Sentimen Terhadap Komentar Video Youtube Menggunakan Support Vector Machines," vol. 19, pp. 231–240, 2023.
- [11] C. Andrew, A. Budiyantara, and I. Lewenusu, "ANALISIS+SENTIMEN+KOMENTAR+PENGGUNA+YOUTUBE+UNTUK+PONSEL+DI+INDONESIA+BERBASIS+WEB," 2024.

- [12] A. Bagus Sasmita, B. Rahayudi, and L. Muflikhah, "Analisis Sentimen Komentar pada Media Sosial Twitter tentang PPKM Covid-19 di Indonesia dengan Metode Naïve Bayes," 2022. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [13] A. Mudya Yolanda and R. Tri Mulya, "Implementasi Metode Support Vector Machine untuk Analisis Sentimen pada Ulasan Aplikasi Sayurbox di Google Play Store," *VARIANSI: Journal of Statistics and Its Application on Teaching and Research*, vol. 6, no. 2, pp. 76–83, 2024, doi: 10.35580/variansiunm258.
- [14] M. Ma'rufudin and A. Yudhistira, "Analisis Sentimen Petani Milenial Pada Media Sosial X Menggunakan Algoritma Support Vector Machine (SVM)," *Jurnal Pendidikan dan Teknologi Indonesia*, vol. 5, no. 3, pp. 845–857, Mar. 2025, doi: 10.52436/1.jpti.717.
- [15] B. Budiman, Z. Silvana Anggraeni, C. Habibi, and N. Alamsyah, "Analisis Sentimen Publik pada Media Sosial Twitter Terhadap Tiket.com Menggunakan Algoritma Klasifikasi," *Jurnal Informatika*, vol. 11, no. 1, pp. 1–10, Apr. 2024, doi: 10.31294/inf.v11i1.17988.
- [16] N. Silalahi and G. L. Ginting, "Analisa Sentimen Masyarakat Dalam Penggunaan Vaksin Sinovac Dengan Menerapkan Algoritma Term Frequence – Inverse Document Frequence (TF-IDF) dan Metode Deskripsi," *Journal of Information System Research (JOSH)*, vol. 3, no. 3, pp. 206–217, Apr. 2022, doi: 10.47065/josh.v3i3.1441.
- [17] N. V. Loca and D. Abdullah, "Analisis Sentimen terhadap Dedy Mulyadi Berdasarkan Komentar Youtube Menggunakan Metode Naïve Bayes," *Progresif: Jurnal Ilmiah Komputer*, vol. 21, no. 2, p. 874, Aug. 2025, doi: 10.35889/progresif.v21i2.3102.
- [18] T. Widyanto, I. Ristiana, and A. Wibowo, "Komparasi Naïve Bayes dan SVM Analisis Sentimen RUU Kesehatan di Twitter," *SINTECH Journal*, vol. 6, 2023, [Online]. Available: <https://doi.org/10.31598>
- [19] N. D. Japari *et al.*, "PERAN PT. PERTAMINA SEBAGAI PENYEDIA PASOKAN BAHAN BAKAR MINYAK DI INDONESIA."
- [20] H. B. Tambunan and T. W. D. Hapsari, "Analisis Opini Pengguna Aplikasi New PLN Mobile Menggunakan Text Mining," *PETIR*, vol. 15, no. 1, pp. 121–134, Dec. 2021, doi: 10.33322/petir.v15i1.1352.
- [21] S. T. Andini, A. Eviyanti, H. Setiawan, and C. Taurusta, "Analisis Sentimen Pengguna Aplikasi Tantan: Perbandingan Kinerja Metode Naive Bayes dan SVM Sentiment Analysis of Tantan Application Users: A Performance Comparison Between Naive Bayes and SVM."