

Analisis Sentimen Kebijakan Pembatasan Subsidi Bahan Bakar Minyak di Indonesia Tahun 2024 Menggunakan Algoritma Klasifikasi

Muh Chaerul^{*1}, Septiadi², Gandung Triyono³

^{1,2,3}Magister Ilmu Komputer, Fakultas Teknologi Informasi, Universitas Budi Luhur, Indonesia
Email: ¹2211602004@student.budiluhur.ac.id, ²2211602095@student.budiluhur.ac.id,
³gandung.triyono@budiluhur.ac.id

Abstrak

Kebijakan pembatasan subsidi Bahan Bakar Minyak (BBM) yang diterapkan pemerintah Indonesia pada tahun 2024 memicu beragam tanggapan masyarakat di media sosial. Penelitian ini bertujuan untuk menganalisis sentimen pengguna X (Twitter) terhadap kebijakan pemerintah tersebut. Dataset yang digunakan terdiri dari 2.011 tweet yang dikumpulkan melalui teknik *scraping* pada tweet periode 1 September 2024 hingga 31 Desember 2024. Data tersebut kemudian melalui tahap *cleansing*, preprocessing dan pelabelan sentimen menggunakan *lexicon-based approach* untuk mengkategorikan tweet ke dalam tiga kelas sentimen yaitu positif, negatif, dan netral. Selanjutnya, dilakukan ekstraksi fitur menggunakan Term Frequency-Inverse Document Frequency (TF-IDF) Vectorizer dan penyeimbangan data dengan teknik Synthetic Minority Oversampling Technique (SMOTE). Pemodelan dilakukan menggunakan dua algoritma klasifikasi yaitu Support Vector Classifier (SVC) dan Random Forest Classifier (RFC). Penalaan parameter dilakukan dengan memanfaatkan Stratified K-Fold Cross-Validation untuk menjaga keseimbangan distribusi kelas selama proses validasi model. Berdasarkan hasil analisis dataset, mayoritas sentimen yang ditemukan adalah negatif sebesar 59%, diikuti oleh sentimen positif sebesar 31%, dan netral sebesar 10%. Selain itu, penelitian juga menghasilkan temuan bahwa model SVC mencapai akurasi lebih baik dibandingkan model RFC yaitu sebesar 77%, sedangkan RFC memiliki nilai akurasi sebesar 73%. Temuan ini mengindikasikan bahwa SVC lebih efektif dalam mengklasifikasikan sentimen terkait kebijakan pembatasan subsidi BBM. Penelitian ini memberikan juga insight mengenai persepsi masyarakat terhadap kebijakan tersebut, yang dapat menjadi pertimbangan bagi pemerintah dalam merumuskan strategi komunikasi dan implementasi kebijakan yang lebih baik di masa depan.

Kata kunci: *analisis sentimen, random forest classifier, subsidi bahan bakar minyak, support vector classifier, twitter*

The Sentiment Analysis of the Fuel Subsidy Limitation Policy Using Support Vector Classifier and Random Forest Classifier Algorithms.

Abstract

The policy of fuel subsidy restrictions implemented by the Indonesian government in 2024 has elicited a wide range of public responses on social media. This study aims to analyze the sentiment of X (Twitter) users towards the government's policy. The dataset utilized consists of 2,011 tweets collected through a scraping technique, covering the period from September 1, 2024, to December 31, 2024. The data underwent a series of cleansing, preprocessing, and sentiment labeling processes using a lexicon-based approach to categorize the tweets into three sentiment classes: positive, negative, and neutral. Feature extraction was subsequently conducted using the Term Frequency-Inverse Document Frequency (TF-IDF) Vectorizer, and data balancing was performed using the Synthetic Minority Oversampling Technique (SMOTE). Modeling was carried out using two classification algorithms: Support Vector Classifier (SVC) and Random Forest Classifier (RFC). Hyperparameter tuning was conducted using Stratified K-Fold Cross-Validation to maintain class distribution balance during model validation. Based on the dataset analysis, the majority of sentiments identified were negative at 59%, followed by positive sentiments at 31%, and neutral sentiments at 10%. Furthermore, the study found that the SVC model achieved a higher accuracy rate compared to the RFC model, with an accuracy of 77% and 73%, respectively. These findings indicate that SVC is more effective in classifying sentiments related to the fuel subsidy restriction policy. This research also provides insights into public perceptions of the policy, which may serve as valuable input for the government in formulating better communication strategies and policy implementations in the future.

Keywords: *fuel subsidy, random forest classifier, sentiment analysis, support vector classifier, twitter*

1. PENDAHULUAN

Kebijakan subsidi Bahan Bakar Minyak (BBM) telah lama menjadi salah satu instrumen fiskal yang digunakan oleh pemerintah Indonesia untuk menjaga stabilitas harga dan melindungi daya beli masyarakat [1]. Namun, dalam beberapa tahun terakhir, kebijakan ini dinilai kurang tepat sasaran, sehingga pemerintah mulai melakukan pembatasan subsidi BBM pada tanggal 1 Oktober 2024 dan mengalihkannya ke program bantuan sosial seperti Bantuan Langsung Tunai, Bantuan Subsidi Upah dan Bantuan Pemerintah Daerah [2]. Perubahan kebijakan ini menimbulkan berbagai reaksi dari masyarakat, yang terekspresikan secara luas melalui platform media sosial, termasuk X (Twitter).

Media sosial telah menjadi wadah penting bagi masyarakat untuk menyampaikan pendapat, kritik, maupun dukungan terhadap kebijakan pemerintah [3]. Sehingga analisis sentimen terhadap diskusi publik di platform ini dapat memberikan gambaran yang komprehensif mengenai persepsi masyarakat terhadap kebijakan pembatasan subsidi BBM. Namun, tantangan utama dalam analisis sentimen adalah mengolah data yang bersifat tidak terstruktur [4], seperti tweet, serta mengklasifikasikannya secara akurat ke dalam kategori sentimen yang relevan.

Analisis sentimen adalah proses yang bertujuan untuk mengidentifikasi dan mengkategorikan emosi atau sentimen yang diungkapkan dalam teks tertulis [5]. Penelitian ini bertujuan untuk menganalisis sentimen pengguna Twitter terhadap rencana pembatasan subsidi BBM menggunakan pendekatan *machine learning*. Dengan memanfaatkan algoritma *Support Vector Classifier* (SVC) dan *Random Forest Classifier* (RFC), penelitian ini berupaya mengidentifikasi pola sentimen yang dominan serta membandingkan performa kedua algoritma dalam mengklasifikasikan data. Hasil penelitian diharapkan dapat memberikan insight yang berguna bagi pemerintah dalam memahami persepsi publik dan merumuskan strategi komunikasi yang lebih efektif terkait kebijakan ini.

Beberapa penelitian terdahulu membandingkan beberapa model *machine learning* dalam memprediksi analisis sentimen media sosial. Model *Support Vector Classifier* (SVC) dan *Random Forest Classifier* (RFC) digunakan oleh Sudianto dkk. pada Tahun 2022 dalam menganalisis sentimen twitter kaburnya selebgram dari karantina [6]. Model kemudian memprediksi sentimen kedalam dua label yaitu sentimen positif dan negative. Hasil penelitian menghasilkan model *Random Forest* memberikan hasil prediksi terbaik dengan nilai akurasi sebesar 94%. Pada penelitian lain, model SVC unggul dalam mengklasifikasi tingkat stres berdasarkan data teks jika dibandingkan dengan model RFC dimana model SVC mendapatkan nilai akurasi sebesar 84% [7]. Model SVC juga unggul pada penelitian dalam penelitian analisis sentimen quick count Pemilu 2024 [8]. Penelitian ini mengklasifikasikan tweet ke dalam tiga label yaitu sentimen positif, negatif dan netral dari dataset yang berjumlah 2000 data. Hasil dari penelitian ini menunjukkan model SVC mendapatkan nilai akurasi tertinggi dengan nilai sebesar 80% dibandingkan model RFC dengan nilai akurasi 78%.

Giovani Tahun 2020 membandingkan model SVM, K-Nearest Neighbor dan Naïve Bayes dalam memprediksi sentimen positif dan negative pada aplikasi ruang guru di twitter [9]. Penelitian kemudian membandingkan hasil prediksi tanpa seleksi fitur dan dengan seleksi fitur menggunakan algoritma *Particle Swarm Optimazition* (PSO). Hasil penelitian menunjukkan model SVM dengan seleksi fitur memberikan prediksi terbaik dengan nilai akurasi sebesar 78,55%.

Ekstraksi fitur menggunakan Term Frequency-Inverse Document Frequency (TF-IDF) dilakukan pada penelitian Sebastian Tahun 2024 dengan model klasifikasi SVM digunakan untuk memprediksi dua kelas sentimen yaitu positif dan negatif [10]. Hasil menunjukkan bahwa SVM mampu menghasilkan prediksi dengan nilai akurasi sebesar 77%. Larasati Tahun 2022 menggunakan model *Random Forest* dan metode TF-IDF untuk memprediksi tiga kelas sentimen pada ulasan aplikasi Dana dimana SVM memberikan nilai akurasi sebesar 84% [11].

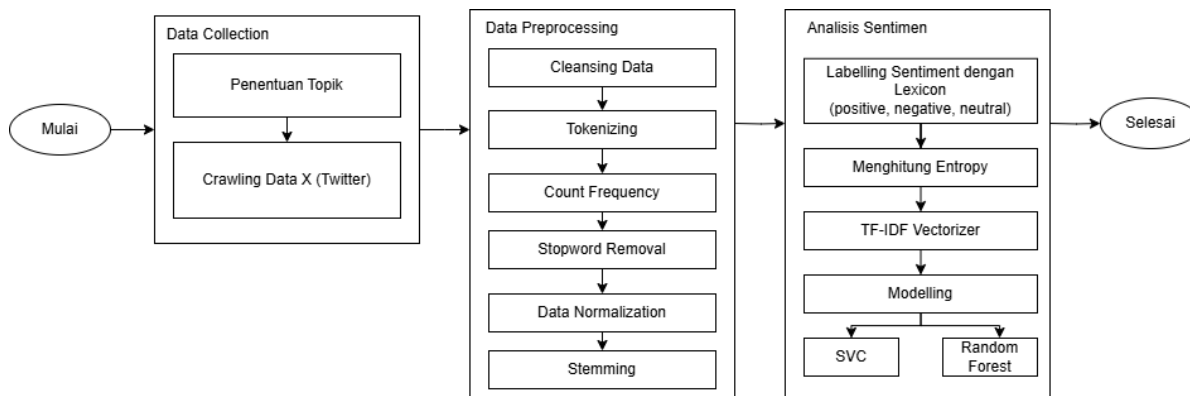
Teknik Synthetic Minority Oversampling Technique (SMOTE) diterapkan pada penelitian Baker Tahun 2023 dalam analisis sentimen tweet terkait agresi Rusia ke Ukraina [12]. Terdapat delapan model Machine Learning yang dibandingkan dalam memprediksi dua kelas sentimen dimana hasil yang didapatkan secara rata-rata model ML mampu memprediksi sentimen dengan sangat baik.

Penelitian ini berbeda dari sebelumnya karena fokus pada domain kebijakan subsidi energi yang memiliki sensitifitas tinggi terhadap kondisi sosial-ekonomi masyarakat. Berbeda dari sebagian besar penelitian terdahulu yang cenderung hanya memprediksi dua kelas sentimen (positif dan negatif), penelitian ini memprediksi tiga kelas sentimen, yaitu positif, negatif, dan netral, guna memperoleh gambaran yang lebih komprehensif terhadap opini publik di media sosial.

Metode penelitian yang digunakan meliputi pengumpulan data melalui teknik *scraping*, pemrosesan data (*cleaning* dan *preprocessing*), serta pelabelan sentimen menggunakan pendekatan berbasis *lexicon*. Selanjutnya, ekstraksi fitur dilakukan dengan *TF-IDF Vectorizer* diikuti oleh penyeimbangan data menggunakan teknik SMOTE. Pemodelan dan evaluasi dilakukan dengan *tuning hyperparameter* menggunakan *Stratified K-Fold Cross-Validation* dan *GridSearchCV* untuk memastikan akurasi yang optimal.

2. METODE PENELITIAN

Tahapan yang dilakukan untuk menyelesaikan penelitian ini dapat terlihat pada Gambar 1. berikut.



Gambar 1. Metode Penelitian

Penelitian menggunakan metode yang terdiri dari tiga tahapan utama yaitu pengumpulan data (*Data Collection*), praproses data (*Data Preprocessing*) dan Analisis Sentimen. Pada tahap *Data Collection*, dilakukan penentuan topik terkait kebijakan subsidi BBM, kemudian data dikumpulkan melalui proses crawling dari platform Twitter. Selanjutnya, pada tahap *Data Preprocessing*, data mentah dibersihkan (*cleansing*), dipecah menjadi token (*tokenizing*), dihitung frekuensinya, dihapus kata-kata umum yang tidak memiliki makna penting (*stopword removal*), dinormalisasi, dan dilakukan *stemming* untuk mengembalikan kata ke bentuk dasarnya. Data selanjutnya masuk ke tahapan Analisis Sentimen dimana tahapan ini dimulai dengan pelabelan sentimen menggunakan pendekatan leksikon (*lexicon-based approach*), dilanjutkan dengan perhitungan entropy dan penerapan TF-IDF vectorizer untuk representasi fitur. Langkah akhir adalah pemodelan menggunakan dua algoritma machine learning, yaitu Support Vector Classifier (SVC) dan Random Forest, untuk mengklasifikasikan sentimen publik terhadap kebijakan pembatasan subsidi BBM tersebut.

2.1. Data Collection

Topik permasalahan dalam penelitian ini adalah analisis sentimen terkait kebijakan pembatasan subsidi Bahan Bakar Minyak (BBM). Oleh karena itu, kata kunci yang digunakan dalam proses data crawling adalah “subsidi BBM”. Periode data dimulai dari tweets tanggal 1 September 2024 hingga 31 Desember 2024, yang mencakup berbagai diskusi publik mengenai kebijakan tersebut terutama setelah pengumuman kebijakan pembatasan BBM. Proses crawling data dilakukan dengan memanfaatkan *tweet-harvest*, sebuah alat berbasis Node.js, yang dijalankan melalui perintah *npx* dalam lingkungan Python [13].

Data yang diperoleh melalui scraping terdiri dari 2.011 tweet yang relevan dengan topik. Tweets tersebut mencakup berbagai pendapat, reaksi, dan opini dari pengguna media sosial yang mengomentari atau memberikan tanggapan terhadap kebijakan pembatasan subsidi BBM.

2.2. Data Preprocessing

Pada tahapan ini, data tweets yang telah dikumpulkan akan melalui serangkaian proses preprocessing untuk memastikan data siap digunakan dalam analisis. Tahapan-tahapan yang dilalui adalah sebagai berikut:

1. Case Folding

Tahapan ini bertujuan untuk mengubah seluruh huruf dalam teks menjadi huruf kecil (*lowercase*) [14].

2. Cleansing Data

Tahap *cleansing* data bertujuan untuk membersihkan data dari elemen-elemen yang tidak relevan atau tidak penting bagi analisis. Pada tahap ini, data diolah dengan menghapus data duplikat, mention (@), hashtag, retweet, URL, serta spasi berlebihan. Elemen-elemen ini biasanya tidak memberikan kontribusi langsung pada analisis sentimen dan dapat mengganggu proses pengolahan data [15].

3. Tokenizing

Proses *tokenizing* adalah memecah teks menjadi unit-unit yang lebih kecil, seperti kata atau token [16]. Hal ini mempermudah komputer dalam memahami teks secara lebih terstruktur. Sebagai contoh, sebuah kalimat akan dipecah menjadi setiap kata yang menyusunnya sehingga dapat diolah lebih lanjut oleh algoritma. Tahap

ini sangat penting untuk mempersiapkan data dalam format yang sesuai untuk analisis sentimen, terutama ketika menggunakan teknik machine learning.

4. Count Frequency

Count frequency adalah proses menghitung jumlah kemunculan setiap token atau kata dalam kumpulan data teks setelah melalui tahap tokenizing. Proses ini bertujuan untuk memahami distribusi kata dalam teks, mengidentifikasi kata-kata yang sering muncul, serta memberikan wawasan tentang tema atau pola dominan dalam dataset.

5. Stopwords Removal

Pada tahap ini, kata-kata umum yang tidak memiliki makna signifikan dalam analisis, seperti "yang", "pada", "dengan", "dan", atau "klo", dihapus dari data [17]. Kata-kata ini, yang dikenal sebagai stopwords, biasanya hanya berfungsi sebagai penghubung dalam kalimat dan tidak memberikan informasi penting untuk analisis sentimen. Dengan menghapus stopwords, model dapat lebih fokus pada kata-kata yang benar-benar mencerminkan sentimen atau opini.

6. Normalisasi Data

Normalisasi data adalah proses menyelaraskan teks ke dalam format yang seragam, termasuk mengubah singkatan, simbol, atau penulisan tidak baku menjadi bentuk standar. Misalnya, simbol "&" akan diubah menjadi "dan", kata "abal2" menjadi "palsu", dan kata informal seperti "akooh" menjadi "aku". Normalisasi juga mencakup mengubah semua teks menjadi huruf kecil untuk menyamakan format. Tahap ini penting untuk mengurangi variasi dalam data yang disebabkan oleh perbedaan penulisan, sehingga model dapat mengenali pola dengan lebih baik.

7. Stemming

Stemming adalah proses mengubah kata menjadi bentuk dasarnya dengan menghilangkan imbuhan seperti awalan, akhiran, atau sisipan [18]. Sebagai contoh, kata "menyalahi" akan diubah menjadi "salah" dan "daerahnya" menjadi "daerah". Tujuan dari proses ini adalah menyederhanakan teks sehingga kata-kata dengan makna yang sama tetapi bentuk berbeda dapat diidentifikasi sebagai satu entitas. Dengan demikian, analisis menjadi lebih konsisten dan akurat.

2.3. Sentiment Labeling

Pendekatan berbasis *lexicon* (lexicon-based approach) dalam analisis sentimen menggunakan *lexicon* sentimen untuk memberikan skor pada token yang dikumpulkan dari teks [19]. Pendekatan ini bersifat unsupervised dan bergantung pada domain (*domain reliance*), karena kata yang sama dapat memiliki sentimen berbeda tergantung pada konteksnya.

2.4. Modeling

Teks yang telah diproses sebelumnya, kemudian diubah menjadi representasi numerik menggunakan dua teknik populer, yaitu Count Vectorizer dan TF-IDF Vectorizer. Count Vectorizer mengubah teks menjadi representasi matriks berdasarkan frekuensi kata dalam dokumen, sedangkan TF-IDF Vectorizer memberikan bobot lebih tinggi pada kata-kata yang lebih relevan dengan mengukur frekuensi kata dalam dokumen relatif terhadap keseluruhan corpus [20].

Tahap berikutnya adalah mapping label, di mana setiap teks diberi label sentimen berdasarkan hasil analisis yang mengklasifikasikan teks ke dalam kategori positif, negatif, atau netral. Selanjutnya, dilakukan train-test split untuk membagi data menjadi dua bagian: 80% data untuk pelatihan model dan 20% data lainnya untuk pengujian. Pada data yang tidak seimbang, Synthetic Minority Over-sampling Technique (SMOTE) digunakan untuk mengatasi masalah ketidakseimbangan kelas dengan menghasilkan contoh sintetis dari kelas minoritas untuk memastikan distribusi kelas yang lebih seimbang dan meningkatkan kinerja model [21].

Algoritma pemodelan yang diterapkan pada penelitian ini Support Vector Classifier dan Random Forest Classifier. Setelah itu, dilakukan hyperparameter tuning untuk menemukan kombinasi parameter terbaik bagi setiap model.

2.5. Random Forest Classifier

RFC merupakan algoritma ansambel acak dari suatu pohon keputusan [7]. Algoritma ini dikembangkan oleh Leo Breiman dan Adèle Cutler pada tahun 2001.

Random Forest adalah algoritma untuk klasifikasi dan regresi yang terdiri dari sekumpulan pohon keputusan (decision trees) yang berfungsi sebagai base classifier. Algoritma ini membangun dan mengombinasikan banyak pohon keputusan untuk meningkatkan akurasi prediksi. Salah satu aspek penting dalam Random Forest adalah *bootstrap sampling*, yaitu teknik pengambilan sampel secara acak dengan pengembalian untuk membangun setiap

pohon keputusan. Setiap pohon dalam Random Forest melakukan prediksi berdasarkan subset fitur yang dipilih secara acak. Untuk klasifikasi, hasil akhir ditentukan berdasarkan mayoritas prediksi dari seluruh pohon (voting), sedangkan untuk regresi, hasil akhir diperoleh dari rata-rata prediksi semua pohon [6].

2.6. Support Vector Classifier

SVC merupakan salah satu metode dalam pembelajaran mesin yang berbasis pada Support Vector Machine (SVM) dan digunakan untuk tugas klasifikasi. Algoritma SVM bekerja dengan menemukan hyperplane terbaik yang memisahkan data ke dalam kelas-kelas yang berbeda.

SVM bertujuan untuk menemukan hyperplane di ruang fitur yang memisahkan data dari dua kelas dengan margin terbesar. Margin adalah jarak terdekat antara data titik mana pun ke hyperplane. Hyperplane terbaik adalah yang memiliki margin terbesar.

Persamaan dasar dalam SVM adalah sebagai berikut:

$$f(x) = w^T x + b \quad (1)$$

dimana:

$f(x)$ = output atau nilai yang akan diprediksi

w^T = vektor bobot

x = vektor fitur dari data input

b = bias

Pada dasarnya, SVC mencari hyperplane terbaik seperti SVM serta juga mendukung pemisahan non-linier menggunakan kernel trick dengan fungsi kernel sebagai berikut:

$$K(x_i, x_j) = \phi(x_i)^T \phi(x_j) \quad (2)$$

dimana:

$K(x_i, x_j)$ = fungsi kernel, yang menghitung kesamaan antara dua titik data tanpa menghitung eksplisit $\phi(x)$

$\phi(x)$ = fungsi pemetaan (transformasi) dari ruang input ke ruang berdimensi lebih tinggi (*feature space*)

2.7. Evaluasi

Tahapan evaluasi digunakan untuk untuk mengetahui sejauh mana model yang dibangun mampu memberikan hasil prediksi yang baik dan akurat. Dalam penelitian ini, evaluasi dilakukan menggunakan empat metrik utama, yaitu akurasi, presisi (precision), recall (sensitivitas), dan F1-score. Seluruh metrik ini dihitung berdasarkan hasil confusion matrix yang merupakan dasar evaluasi dalam permasalahan klasifikasi multi kelas [22]. Struktur dari confusion matrix dapat dilihat pada Tabel 1.

Tabel 1. Confusion Matrix Tiga Kelas Sentimen

	Prediksi Positif	Prediksi Netral	Prediksi Negatif
Aktual Positif	TP_Pos	FN_Pos_Neu	FP_Pos_Neg
Aktual Netral	FP_Neu_Pos	TP_Neu	FP_Neu_Neg
Aktual Negatif	FP_Neg_Pos	FP_Neg_Neu	TP_Neg

dimana:

TP_Pos = jumlah tweet positif yang benar diprediksi sebagai positif

TP_Neu = jumlah tweet netral yang benar diprediksi sebagai netral

TP_Neg = jumlah tweet negatif yang benar diprediksi sebagai negative

FN_Pos_Neu = jumlah tweet positif yang salah diprediksi sebagai netral

FN_Pos_Neg = jumlah tweet positif yang salah diprediksi sebagai negative

FP_Neu_Pos = jumlah tweet netral yang salah diprediksi sebagai positif

FN_Neu_Neg = jumlah tweet netral yang salah diprediksi sebagai negative

FP_Neg_Pos = jumlah tweet negatif yang salah diprediksi sebagai positif

FP_Neg_Neu = jumlah tweet negatif yang salah diprediksi sebagai netral

Akurasi digunakan untuk mengukur seberapa sering model benar dalam prediksi, dengan rumus:

$$Akurasi = \frac{TP_{Pos} + TP_{Neu} + TP_{Neg}}{(Total\ Semua\ Prediksi)} \quad (3)$$

Precision (Presisi) digunakan untuk mengukur seberapa akurat prediksi positif model. Dihitung dengan rumus:

$$Precision_{Positive} = \frac{TP_{Pos}}{TP_{Pos} + FP_{Neu_Pos} + FP_{Neg_Pos}} \quad (4)$$

$$Precision_{Neutral} = \frac{TP_{Neu}}{TP_{Neu} + FN_{Pos_Neu} + FP_{Neg_Neu}} \quad (5)$$

$$Precision_{Negative} = \frac{TP_{Neg}}{TP_{Neg} + FP_{Pos_Neg} + FN_{Neu_Neg}} \quad (6)$$

Recall mengukur kemampuan model untuk mengidentifikasi semua kasus aktual dari kelas tertentu. Recall dihitung per kelas dengan rumus:

$$Recall_{Positive} = \frac{TP_{Pos}}{TP_{Pos} + FN_{Pos_Neu} + FN_{Pos_Neg}} \quad (7)$$

$$Recall_{Neutral} = \frac{TP_{Neu}}{TP_{Neu} + FP_{Neu_Pos} + FN_{Neu_Neg}} \quad (8)$$

$$Recall_{Negative} = \frac{TP_{Neg}}{TP_{Neg} + FP_{Neg_Pos} + FP_{Neg_Neu}} \quad (9)$$

F1-score merupakan rata-rata harmonis antara presisi dan recall. F1-Score dihitung per kelas dihitung dengan rumus:

$$F1_{Positive} = 2 \times \frac{Precision_{Positive} \times Recall_{Positive}}{Precision_{Positive} + Recall_{Positive}} \quad (10)$$

$$F1_{Neutral} = 2 \times \frac{Precision_{Neutral} \times Recall_{Neutral}}{Precision_{Neutral} + Recall_{Neutral}} \quad (11)$$

$$F1_{Negative} = 2 \times \frac{Precision_{Negative} \times Recall_{Negative}}{Precision_{Negative} + Recall_{Negative}} \quad (12)$$

Untuk mendapatkan satu nilai yang merepresentasikan performa keseluruhan pada semua kelas, selanjutnya dihitung nilai agregat dengan menggunakan nilai *Macro Average* dan nilai *Weighted Average* pada matriks *Precision*, *Recall* dan *F1-Score*.

2.8. Perangkat dan Lingkungan Pengembangan

Penelitian ini dilakukan dengan menggunakan bahasa pemrograman Python versi 3.11 dan dijalankan pada platform Google Colab yang menyediakan lingkungan komputasi berbasis cloud. Beberapa pustaka dan library yang digunakan dalam proses pengolahan data meliputi *pandas* dan *numpy* untuk manipulasi data, *nlTK* dan *Sastrawi* untuk proses data preprocessing seperti tokenisasi, *stopword removal*, dan *stemming*. Visualisasi data dilakukan menggunakan *matplotlib* dan *seaborn*, sementara *wordcloud* digunakan untuk menampilkan frekuensi kata dalam bentuk visual. Untuk proses klasifikasi sentimen, digunakan *scikit-learn* yang menyediakan algoritma SVC dan RFC, serta *imblearn* untuk menangani ketidakseimbangan kelas dalam data.

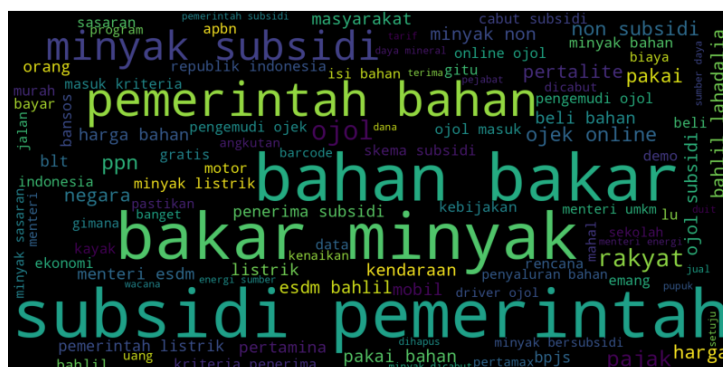
3. HASIL DAN PEMBAHASAN

3.1. Data Collection

Berdasarkan hasil crawling, jumlah tweet yang berhasil dikumpulkan adalah 2.011 *tweet* dengan periode tweet dari tanggal 1 September 2024 hingga 31 Desember 2024. Struktur dataset awal hasil crawling serta deskripsi dari masing-masing kolom dapat dilihat pada Tabel 2. Setiap kolom memiliki 2.011 nilai non-null, yang berarti tidak terdapat data yang hilang atau kosong dalam dataset ini.

Tabel 2. Struktur Dataset Awal

Kolom	Non-Null Count	Tipe	Deskripsi
conversation_id_str	2.011 non-null	object	ID percakapan yang menghubungkan tweet dengan thread atau diskusi yang lebih luas.
id_str	2.011 non-null	object	ID unik untuk setiap tweet yang dikumpulkan.
full_text	2.011 non-null	object	Isi lengkap dari tweet yang mengandung kata kunci "subsidi BBM".
username	2.011 non-null	object	Nama pengguna Twitter yang memposting tweet.
created_at	2.011 non-null	object	Waktu dan tanggal ketika tweet dipublikasikan.



Gambar 2. Visualisasi *Wordcloud* Data Awal

Berdasarkan Gambar 2. terlihat bahwa kata-kata yang dominan berkaitan dengan "subsidi," "pemerintah," "bahan bakar," dan "minyak." Hal ini menunjukkan bahwa diskusi publik berfokus pada peran pemerintah dalam kebijakan subsidi BBM serta dampaknya terhadap harga dan masyarakat. Kata-kata seperti "ojol," "pengemudi," dan "rakyat" mengindikasikan bahwa kelompok yang paling terdampak adalah pengemudi ojek online dan masyarakat umum. Selain itu, munculnya istilah seperti "pajak," "bansos," dan "PPN" menandakan bahwa diskusi juga mencakup aspek fiskal dan bantuan sosial yang terkait dengan subsidi BBM. Sentimen terhadap kebijakan ini kemungkinan besar beragam, dengan adanya kata-kata seperti "cabut" dan "harga" yang dapat mencerminkan kekhawatiran terhadap kenaikan harga BBM akibat pembatasan subsidi.

3.2. Data Preprocessing

Proses data preprocessing mencakup serangkaian tahapan untuk membersihkan dan meningkatkan kualitas data dengan menghilangkan duplikasi, membersihkan teks dari karakter yang tidak relevan, serta memastikan bahwa data dalam format yang sesuai untuk pemodelan.

Hasil dari proses data preprocessing mereduksi jumlah data dari 2.011 menjadi 1.995 baris data yang siap digunakan untuk proses analisis selanjutnya. Tabel 3. memperlihatkan hasil dari data preprocessing pada setiap tahapan.

Tabel 3. Data Hasil Preprocessing

Tabel 3. Data Teks Preprocessing		
Data Preprocessing	Sebelum	Sesudah
<i>Tokenizing</i>	cuman di indonesia pejabatnya gila gaya hidup mewah ala monaco rakyat diperas dg berbagai kewajiban ppn 12 pajak kendaraan naik bbm subsidi dihapus biaya kuliah naik bpjs naik negeri para para bedebah	['cuman', 'di', 'indonesia', 'pejabatnya', 'gila', 'gaya', 'hidup', 'mewah', 'ala', 'monaco', 'rakyat', 'diperas', 'dg', 'berbagai', 'kewajiban', 'ppn', '12', 'pajak', 'kendaraan', 'naik', 'bbm', 'subsidi', 'dihapus', 'biaya', 'kuliah', 'naik', 'bpjs', 'naik', 'negeri', 'para', 'para', 'bedebah']
<i>Stopword Removal</i>	['cuman', 'di', 'indonesia', 'pejabatnya', 'gila', 'gaya', 'hidup', 'mewah', 'ala', 'monaco', 'rakyat', 'diperas', 'dg', 'berbagai', 'kewajiban',	['cuman', 'indonesia', 'pejabatnya', 'gila', 'gaya', 'hidup', 'mewah', 'ala', 'monaco', 'rakyat', 'diperas', 'kewajiban', 'ppn', '12',

Data Preprocessing	Sebelum	Sesudah
<i>Normalisasi Data</i>	'ppn', '12', 'pajak', 'kendaraan', 'naik', 'bbm', 'subsidi', 'dihapus', 'biaya', 'kuliah', 'naik', 'bpjs', 'naik', 'negeri', 'para', 'para', 'bedebah']	'pajak', 'kendaraan', 'bbm', 'subsidi', 'dihapus', 'biaya', 'kuliah', 'bpjs', 'negeri', 'bedebah']
<i>Stemming</i>	['cuman', 'indonesia', 'pejabatnya', 'gila', 'gaya', 'hidup', 'mewah', 'ala', 'monaco', 'rakyat', 'diperas', 'kewajiban', 'ppn', '12', 'pajak', 'kendaraan', 'bbm', 'subsidi', 'dihapus', 'biaya', 'kuliah', 'bpjs', 'negeri', 'bedebah']	['cuman', 'indonesia', 'pejabatnya', 'gila', 'gaya', 'hidup', 'mewah', 'ala', 'monaco', 'rakyat', 'diperas', 'kewajiban', 'ppn', '12', 'pajak', 'kendaraan', 'bbm', 'subsidi', 'dihapus', 'biaya', 'kuliah', 'bpjs', 'negeri', 'bedebah']

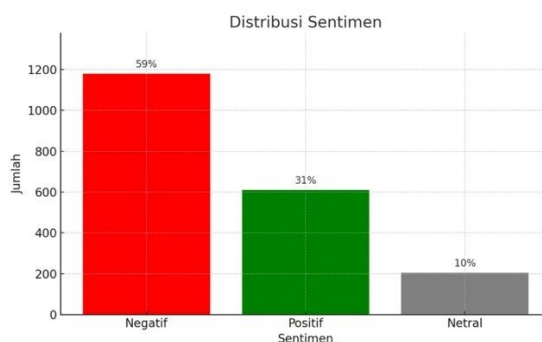
3.3. Sentiment Labeling

Perhitungan entropi dilakukan untuk mengukur tingkat ketidakpastian dalam distribusi data. Menghitung entropi berdasarkan probabilitas kemunculan setiap elemen dalam dataset. Hasil perhitungan entropi menunjukkan sejauh mana variasi dalam distribusi data. Jika nilai entropi yang diperoleh tinggi, berarti distribusi data lebih acak dan bervariasi, yang menandakan keberagaman dalam dataset. Sebaliknya, jika nilai entropi rendah, berarti terdapat pola tertentu yang lebih dominan dalam data, sehingga tingkat ketidakpastian lebih rendah.

Tabel 4. Hasil Perhitungan Entropi dan Sentiment

<i>Full_text</i>	<i>Tweet_Entropy</i>	<i>Tweet_Sentiment</i>	<i>Tweet_polarity_score</i>	<i>Tweet_polarity_Positive_inset</i>	<i>Tweet_polarity_Negative_inset</i>
kenapa wowo selalu berjargon belum bekerja sudah 2 bln lebih yg dihasilkan hanya kontroversi dan kegaduhan mulai dari memilih begundal2 utk kabinet sindir2 org ppn naik tarif air naik subsidi bbm dialihkan jd blt plan makan siang gratis makin ngawur polri makin bobrok	5,5213	negative	-9	{'sudah': 3, 'lebih': 1, 'mulai': 1, 'naik': 1, 'makan': 1, 'gratis': 4, 'makin': 2}	{'sudah': -2, 'hanya': -3, 'kontroversi': -4, 'kegaduhan': -3, 'dari': -3, 'kabinet': -3, 'air': -1, 'makan': -1, 'bobrok': -5}
3 dalam rapat kabinet bayangan stafnya khawatir pak subsidi bbm harus segera kita bahas sandiara menjawab sambil pasang ring light nanti upload dulu algoritma tak pernah menunggu	4,7548	negative	-7	{'dalam': 3, 'rapat': 1, 'bayangan': 3, 'khawatir': 1, 'segera': 3, 'bahas': 2, 'pasang': 4}	{'rapat': -3, 'kabinet': -3, 'khawatir': -5, 'harus': -5, 'segera': -3, 'nanti': -2, 'menunggu': -3}
harga bbm kan yg nentuinnya juga pemerintah coba lebih vokalnya ke yg atur harga bbm lagian pertamina juga ada yg subsidi jd masih cukup membantu sih istilahnya buat org kecil	5,1424	positive	6	{'harga': 3, 'coba': 2, 'atur': 4, 'bantu': 3}	{'coba': -1, 'atur': -4, 'bantu': -4}

Tabel 4 menampilkan hasil perhitungan entropi dan analisis sentimen terhadap beberapa tweet dimana dapat dilihat bahwa tweet dengan entropi tinggi tidak selalu bersentimen positif, melainkan tergantung pada bobot kata-kata bernuansa negatif atau positif dalam tweet tersebut. Misalnya, tweet dengan entropi 5,5213 memiliki sentimen negatif dengan skor polaritas -9, yang dipengaruhi oleh kata-kata seperti "kontroversi", "kegaduhan", dan "bobrok". Sebaliknya, tweet dengan entropi 5,1424 diklasifikasikan sebagai positif dengan skor polaritas +6, didominasi oleh kata-kata seperti "harga", "bantu", dan "atur". Hasil ini memperkuat relevansi pendekatan berbasis entropi dan polaritas kata dalam mendukung algoritma klasifikasi untuk analisis sentimen kebijakan publik, khususnya dalam konteks pembatasan subsidi BBM.



Gambar 3. Distribusi Label Sentimen

Gambar 3 memperlihatkan distribusi label sentimen dimana berdasarkan perhitungan polaritas sentimen dari tweet yang dianalisis menunjukkan bahwa mayoritas tweet memiliki sentimen negatif, dengan total 1.179 tweet atau sekitar 59% dari keseluruhan data. Hal ini mengindikasikan bahwa sebagian besar opini yang terekam dalam dataset cenderung mengungkapkan ketidakpuasan dan kritik terkait topik yang dibahas. Selain itu, terdapat 611 tweet (31%) yang memiliki sentimen positif, menunjukkan sejumlah pengguna yang memberikan pandangan atau respons yang bersifat mendukung atau optimis. Sisanya, sebanyak 205 tweet (10%), dikategorikan sebagai netral, yang berarti tidak mengandung kecenderungan emosi yang kuat ke arah positif maupun negatif.

3.4. Penalaan Parameter

Dataset yang telah dilakukan pemodelan, kemudian dilakukan penalaan parameter (*hyperparameter tuning*) yang bertujuan untuk menemukan kombinasi parameter terbaik yang dapat meningkatkan akurasi dan stabilitas model dalam melakukan klasifikasi. Nilai parameter SVC dan RFC dapat dilihat pada Tabel 5.

Tabel 5. Hyperparameter pada SVC dan RFC

Model	Hyperparameter	Nilai
Random Forest Classifier	Bootstrap	100, 200, 300
	Max depth	none, 10, 20, 30
	Min samples leaf	1, 2, 4
	Min samples split	2, 5, 10
	N_estimator	scale, auto
Support Vector Classifier	C	0.1, 1, 10
	Kernel	linear, rbf
	Gamma	scale, auto

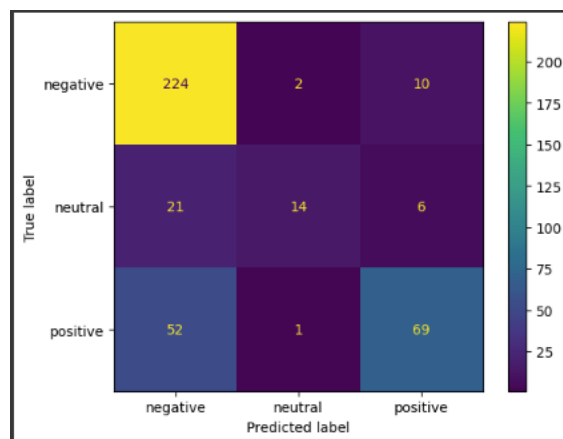
Berdasarkan proses penalaan parameter diperoleh kombinasi parameter terbaik pada model RFC, yaitu $n_estimators = 200$, $max_depth = None$, $min_samples_split = 2$, $min_samples_leaf = 1$, dan $bootstrap = False$, dengan akurasi terbaik sebesar 89,73%. Parameter ini menunjukkan bahwa model bekerja optimal dengan 200 pohon keputusan tanpa batasan kedalaman, sehingga setiap pohon dapat belajar dari seluruh data hingga mencapai pemisahan terbaik. Selain itu, penggunaan $min_samples_split = 2$ dan $min_samples_leaf = 1$ menunjukkan bahwa model diperbolehkan untuk membuat split sesering mungkin, yang dapat meningkatkan akurasi tetapi juga berisiko overfitting. Namun, dengan $bootstrap = False$, model belajar dari seluruh dataset tanpa pengambilan sampel acak, yang mungkin lebih efektif untuk dataset dengan jumlah sampel yang cukup besar.

Sedangkan pada model SVC diperoleh kombinasi parameter terbaik yaitu nilai $C = 10$ yang membuat model lebih ketat dalam memisahkan kelas, sementara pemilihan $gamma = 'scale'$ memungkinkan model untuk

menyesuaikan pengaruh tiap titik data berdasarkan karakteristik dataset. Skor terbaik didapatkan pada nilai 92,13% yang menandakan bahwa model sudah cukup optimal untuk dilakukan klasifikasi berdasarkan parameter terbaik.

3.5. Hasil Evaluasi

Hasil evaluasi dari kedua model klasifikasi sentimen, yaitu Support Vector Classifier (SVC) dan Random Forest Classifier (RFC), menunjukkan performa yang bervariasi dalam mendeteksi tiga kelas sentimen, terlihat dari *Confusion Matrix* dan tabel evaluasi model dibawah ini.



Gambar 4. Confusion Matrix Support Vector Classifier

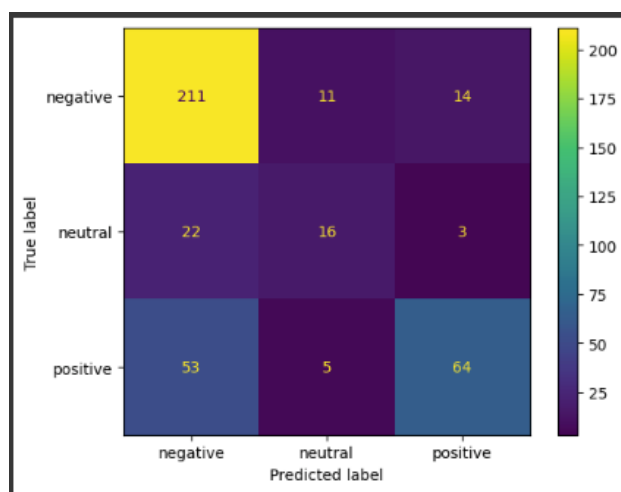
Gambar 4 menampilkan *confusion matrix* pada model Support Vector Classifier (SVC), dari matriks tersebut terlihat performa klasifikasi paling baik terdapat pada kelas negatif (*negative*), di mana model mampu mengklasifikasikan 224 dari total 236 data secara benar. Jumlah kesalahan klasifikasi pada kelas ini tergolong rendah, yaitu hanya 12 data yang salah diklasifikasikan sebagai kelas lain (2 ke netral dan 10 ke positif). Hal ini sejalan dengan nilai *recall* kelas negatif yang sangat tinggi (0.95), menunjukkan bahwa hampir semua data berlabel negatif berhasil dikenali dengan benar oleh model. Nilai *precision* juga cukup baik di angka 0.75, menandakan bahwa mayoritas prediksi kelas negatif memang benar adanya.

Sebaliknya, model menunjukkan kelemahan dalam mengenali kelas netral (*neutral*). Hanya 14 dari 41 data netral yang berhasil diklasifikasikan dengan tepat, sementara sisanya salah diklasifikasikan ke kelas negatif 21 data dan positif (*positive*) 6 data. Ini terlihat jelas dari nilai *recall* yang rendah (0.34), mencerminkan bahwa sebagian besar data berlabel netral tidak terdeteksi dengan baik. Meskipun *precision* kelas ini cukup tinggi (0.82), artinya prediksi ke kelas netral cukup akurat, tapi jumlah prediksi tersebut sangat sedikit menunjukkan bahwa model terlalu "ragu" atau "konservatif" dalam mengklasifikasikan ke kelas netral.

Tabel 6. Evaluasi Model Support Vector Classifier

Sentimen	Precision	Recall	F1-Score	Support
Negatif	0.75	0.95	0.84	236
Netral	0.82	0.34	0.48	41
Positif	0.81	0.57	0.67	122
Accuracy			0.77	399
Macro Avg	0.80	0.62	0.66	399
Weighted Avg	0.78	0.77	0.75	399

Secara keseluruhan, model SVC memiliki akurasi sebesar 0.77 seperti terlihat pada Tabel 6. evaluasi model *Support Vector Classifier*, yang menunjukkan performa cukup baik secara umum. Namun, jika melihat *metrik macro average* yang memperlakukan setiap kelas secara setara didapatkan nilai precision 0.80, recall 0.62, dan F1-score 0.66. Ini mengindikasikan bahwa ketidakseimbangan performa antar kelas masih cukup terasa, terutama karena buruknya kinerja pada kelas netral. Sedangkan *weighted average* yang memperhitungkan jumlah data per kelas memberikan gambaran lebih seimbang, dengan F1-score 0.75, mendekati akurasi keseluruhan.



Gambar 5. Confusion Matrix Random Forest Classifier

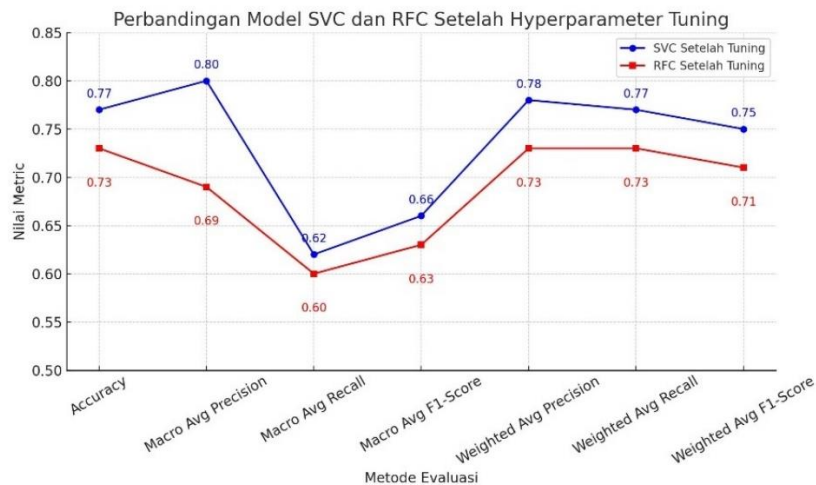
Pada Gambar 5 terlihat bahwa model memiliki performa terbaik dalam mendeteksi sentimen negatif. Sebanyak 211 dari total 236 data berlabel negatif berhasil diklasifikasikan dengan benar. Hal ini tercermin dari *recall* sebesar 0.89 dan *precision* sebesar 0.74, yang berarti model mampu mengenali sebagian besar data negatif dan cukup akurat dalam memprediksinya. Meskipun demikian, masih terdapat sejumlah kesalahan klasifikasi, yaitu 11 data diprediksi sebagai netral dan 14 sebagai positif.

Untuk sentimen netral, model menunjukkan kelemahan yang signifikan. Dari 41 data berlabel netral, hanya 16 yang diklasifikasikan dengan benar, sementara 22 salah diklasifikasikan sebagai negatif dan 3 sebagai positif. Ini menghasilkan nilai *recall* yang rendah (0.39) dan *precision* yang juga terbatas (0.50), yang mengindikasikan bahwa model kesulitan membedakan data netral dari kelas lain. Masalah ini dapat disebabkan oleh kurangnya jumlah data netral dalam pelatihan atau fitur yang tumpang tindih antara kelas netral dengan lainnya.

Tabel 7. Evaluasi Model Random Forest Classifier

Sentimen	Precision	Recall	F1-Score	Support
Negatif	0.74	0.89	0.81	236
Netral	0.50	0.39	0.44	41
Positif	0.79	0.52	0.63	122
Accuracy			0.73	399
Macro Avg	0.68	0.60	0.63	399
Weighted Avg	0.73	0.73	0.72	399

Secara keseluruhan, akurasi model RFC adalah 0.73 seperti terlihat pada Tabel 7, sedikit lebih rendah dari model sebelumnya (SVC). Nilai macro average F1-score sebesar 0.63 menunjukkan bahwa jika ketiga kelas diperlakukan sama penting, performa model masih kurang optimal, terutama karena buruknya klasifikasi pada kelas netral. Sementara *weighted average* F1-score sebesar 0.72 memperlihatkan bahwa secara keseluruhan, model bekerja cukup stabil karena tertolong oleh performa pada kelas dengan jumlah data lebih banyak, yaitu negatif dan positif.



Gambar 6. Perbandingan Model SVC dan RFC

Berdasarkan perbandingan antara model Support Vector Classifier (SVC) dan Random Forest Classifier (RFC) seperti yang terlihat pada Gambar 6, dapat disimpulkan bahwa SVC memiliki performa keseluruhan yang lebih baik, terutama dalam hal akurasi (0.77 vs 0.73) dan macro F1-score (0.66 vs 0.63), serta menunjukkan hasil yang lebih seimbang antara precision dan recall di setiap kelas. Kedua model unggul dalam mengklasifikasikan sentimen negatif, namun sama-sama kesulitan dalam mengenali netral, dengan RFC mencatat kinerja yang lebih rendah pada kelas ini. SVC terbukti lebih efektif dalam menangani ketidakseimbangan kelas, seperti ditunjukkan oleh nilai weighted average F1-score yang lebih tinggi (0.75 vs 0.71). Oleh karena itu, meskipun keduanya masih perlu perbaikan pada deteksi sentimen netral, SVC lebih direkomendasikan sebagai model klasifikasi sentimen dalam konteks ini.

Hasil dua penelitian sebelumnya dengan dengan topik penelitian Klasifikasi Tingkat Stres [7] dan Sentiment Quick Count Pemilu 2024 [8] Model SVC secara konsisten menunjukkan performa lebih baik dibandingkan Model RFC dalam tugas klasifikasi teks berbasis sentimen. Pada kedua penelitian tersebut, SVC unggul dalam akurasi masing-masing 80% dan 84%, mengindikasikan kemampuannya menggeneralisasi pola-pola sentimen dari data media sosial yang relatif bervariasi. Sedangkan Penelitian ini tetap menunjukkan bahwa SVC lebih baik dalam menjaga keseimbangan *precision* dan *recall* antar kelas dibanding RFC dengan tingkat akurasi sebesar 77%. Ketiga penelitian ini menyoroti bahwa SVC lebih tahan terhadap ketidakseimbangan data, terutama pada data dengan dominasi sentimen negatif seperti dalam kasus sentimen subsidi BBM, sedangkan RFC cenderung mengalami kesulitan, khususnya dalam mendeteksi sentimen netral.

Secara metodologis, faktor kunci yang membedakan kinerja di setiap penelitian adalah kompleksitas preprocessing data dan teknik pembobotan fitur. Pada klasifikasi Tingkat Stres [7], penggunaan TF-IDF dan pengolahan afeksi teks meningkatkan performa SVC secara signifikan, sedangkan pada penelitian ini, meskipun telah dilakukan pembersihan dan normalisasi data, masih ada tantangan dalam membedakan sentimen netral akibat dominasi opini negatif. Sementara dataset pada penelitian sentiment Quick Count Pemilu 2024 [8] relatif lebih seimbang baik pada sentimen positif, netral maupun negatif, sehingga model dapat belajar lebih baik. Analisis ini menunjukkan bahwa keberhasilan klasifikasi teks tidak hanya bergantung pada pilihan model, tetapi juga sangat dipengaruhi oleh kualitas preprocessing, teknik feature engineering, serta distribusi sentimen dalam dataset. Dengan demikian, SVC dapat dikatakan lebih stabil dalam berbagai variasi dataset sosial, sedangkan RFC memerlukan data yang lebih terstruktur dan seimbang untuk mencapai performa optimal.

4. KESIMPULAN

Penelitian ini berhasil melakukan analisis sentimen menggunakan dua algoritma pembelajaran mesin, yaitu *Support Vector Classifier* (SVC) dan *Random Forest Classifier* (RFC), untuk mengklasifikasikan sentimen dari data yang telah dikumpulkan. Berdasarkan hasil distribusi sentimen, mayoritas opini masyarakat cenderung menolak kebijakan pembatasan subsidi BBM ini. Hal ini terlihat dari tingginya sentimen negatif yaitu sebesar 59% dari total seluruh sentimen. Berdasarkan hasil klasifikasi sentimen, model SVC memberikan performa lebih tinggi dibandingkan RFC dengan nilai akurasi sebesar 77%. Temuan ini memberikan kontribusi dalam pengembangan sistem pemantauan opini publik berbasis data, yang dapat dimanfaatkan oleh pemerintah sebagai alat bantu dalam merumuskan dan mengevaluasi kebijakan secara responsif terhadap suara masyarakat. Penelitian selanjutnya disarankan untuk dapat menggunakan metode representasi teks yang lebih kontekstual seperti Word Embedding

(Word2Vec, GloVe) atau transformer-based models (BERT) agar model lebih sensitif terhadap nuansa kalimat yang kompleks seperti sarkasme dan ironi yang umum ditemukan di media sosial.

DAFTAR PUSTAKA

- [1] S. Budiantoro, "Policy Brief: Subsidi Dalam Penguatan Kebijakan Fiskal Pro Kemiskinan," 2013. <https://repository.theprakarsa.org/media/publications/661-policy-brief-04-subsidi-dalam-penguatan-7763bafc.pdf>. (accessed: Mar. 04, 2025).
- [2] Tempo, "Kenapa Subsidi BBM Dianggap Tidak Tepat Sasaran , Rencana Prabowo Mengubah Menjadi BLT," 2024. <https://www.tempo.co/ekonomi/kenapa-subsidi-bbm-dianggap-tidak-tepat-sasaran-rencana-prabowo-mengubah-menjadi-blt-1164466>. (accessed: Mar. 03, 2025).
- [3] R. N. Muhammad, L. Wulandari, and B. Tanggahma, "Pengaruh Media Sosial Pada Persepsi Publik Terhadap Sistem Peradilan : Analisis Sentimen di Twitter," *UNES Law Rev.*, vol. 7, no. 1, pp. 507–516, 2024, doi: 10.31933/unesrev.v7il.
- [4] D. Purnamasari *et al.*, *Pengantar Metode Analisis Sentimen*. Indonesia: Penerbit Gunadarma Indonesia, 2023.
- [5] K. L. Tan, C. P. Lee, and K. M. Lim, "A Survey of Sentiment Analysis: Approaches, Datasets, and Future Research," *Appl. Sci.*, vol. 13, no. 7, pp. 1–21, 2023, doi: 10.3390/app13074550.
- [6] Sudianto, P. Wahyuningtias, H. Warih Utami, U. Ahda Raihan, H. Nur Hanifah, and Y. Nicholas Adanson, "Perbandingan Metode Random Forest Dan Support Vector Machine Pada Analisis Sentimen Twitter (Studi Kasus: Kaburnya Selebgram Rachel Vennya Dari Karantina)," *J. Tek. Inform.*, vol. 3, no. 1, pp. 141–145, 2022, doi: 10.20884/1.jutif.2022.3.1.168.
- [7] N. Fathirachman Mahing, A. Lazuardi Gunawan, A. Foresta Azhar Zen, F. Abdurrachman Bachtiar, and S. Agung Wicaksono, "Klasifikasi Tingkat Stress dari Data Berbentuk Teks dengan Menggunakan Algoritma Support Vector Machine (SVM) dan Random Forest," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 7, pp. 1527–1536, 2023, doi: 10.25126/jtiik.1078010.
- [8] I. Septiana and D. Alita, "Perbandingan Random Forest dan SVM dalam Analisis Sentimen Quick Count Pemilu 2024," *J. Inform. J. Pengemb. IT*, vol. 9, no. 3, pp. 224–233, 2024, doi: 10.30591/jpit.v9i3.6640.
- [9] A. P. Giovani, A. Ardiansyah, T. Haryanti, L. Kurniawati, and W. Gata, "Analisis Sentimen Aplikasi Ruang Guru Di Twitter Menggunakan Algoritma Klasifikasi," *J. Teknoinfo*, vol. 14, no. 2, pp. 116–124, 2020, doi: 10.33365/jti.v14i2.679.
- [10] D. F. Sebastian, H. Sulistiani, and A. R. Isnain, "Analisis Sentimen Opini Masyarakat Mengenai Hak Angket di Indonesia Tahun 2024 Menggunakan Metode Support Vector Machine (SVM)," *J. Tek. Inform.*, vol. 5, no. 4, pp. 1025–1034, 2024, doi: 10.52436/1.jutif.2024.5.4.1968.
- [11] F. A. Larasati, D. E. Ratnawati, and B. T. Hanggara, "Analisis Sentimen Ulasan Aplikasi Dana dengan Metode Random Forest," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 6, no. 9, pp. 4305–4313, 2022.
- [12] M. R. Baker, Y. N. Taher, and K. H. Jihad, "Prediction of People Sentiments on Twitter Using Machine Learning Classifiers During Russian Aggression in Ukraine," *Jordanian J. Comput. Inf. Technol.*, vol. 9, no. 3, pp. 189–206, 2023, doi: 10.5455/jjcit.71-1676205770.
- [13] H. Satria, "Crawl Data Twitter Menggunakan Tweet Harvest." 2023. <https://helimisatria.com/blog/crawl-data-twitter-menggunakan-tweet-harvest/>. (accessed: Apr. 03, 2025).
- [14] A. Maulana, Inayah Khasnaputri Afifah, Asghafi Mubarrak, Kiagus Rachmat Fauzan, Ardhan Dwintara, and B. P. Zen, "Comparison of Logistic Regression, Multinomialnb, SVM, and K-NN Methods on Sentiment Analysis of Gojek App Reviews on the Google Play Store," *J. Tek. Inform.*, vol. 4, no. 6, pp. 1487–1494, 2023, doi: 10.52436/1.jutif.2023.4.6.863.
- [15] W. Y. Chong, B. Selvaretnam, and L. K. Soon, "Natural Language Processing for Sentiment Analysis: An Exploratory Analysis on Tweets," in *Proceedings - 2014 4th International Conference on Artificial Intelligence with Applications in Engineering and Technology, ICAIET 2014*, 2014, pp. 212–217. doi: 10.1109/ICAIET.2014.43.
- [16] A. Bayhaqy, S. Sfenrianto, K. Nainggolan, and E. R. Kaburuan, "Sentiment Analysis about E-Commerce from Tweets Using Decision Tree, K-Nearest Neighbor, and Naïve Bayes," in *2018 International Conference on Orange Technologies, ICOT 2018*, 2018, pp. 1–6. doi: 10.1109/ICOT.2018.8705796.
- [17] R. Faiz Ananda, A. Syahri, and F. N. Hasan, "Sentiment Analysis of Customer Satisfaction in Gojek and Grab Application Reviews Using the Naive Bayes Algorithm," *J. Tek. Inform.*, vol. 5, no. 1, pp. 233–241,

- 1680, doi: 10.52436/1.jutif.2024.5.1.1680.
- [18] S. N. Listyarini and D. A. Anggoro, "Analisis Sentimen Pilkada di Tengah Pandemi Covid-19 Menggunakan Convolution Neural Network (CNN)," *J. Pendidik. dan Teknol. Indones.*, vol. 1, no. 7, pp. 261–268, 2021, doi: 10.52436/1.jpti.60.
- [19] M. Wankhade, A. C. S. Rao, and C. Kulkarni, "A survey on sentiment analysis methods, applications, and challenges," *Artif. Intell. Rev.*, vol. 55, no. 7, pp. 5731–5780, 2022, doi: 10.1007/s10462-022-10144-1.
- [20] J. E. Br Sinulingga and H. C. K. Sitorus, "Analisis Sentimen Opini Masyarakat terhadap Film Horor Indonesia Menggunakan Metode SVM dan TF-IDF," *J. Manaj. Inform.*, vol. 14, no. 1, pp. 42–53, 2024, doi: 10.34010/jamika.v14i1.11946.
- [21] R. Obiedat *et al.*, "Sentiment Analysis of Customers' Reviews Using a Hybrid Evolutionary SVM-Based Approach in an Imbalanced Data Distribution," *IEEE Access*, vol. 10, pp. 22260–22273, 2022, doi: 10.1109/ACCESS.2022.3149482.
- [22] Y. Mao, Q. Liu, and Y. Zhang, "Sentiment analysis methods , applications , and challenges : A systematic literature review," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 36, no. 4, pp. 1–16, 2024, doi: 10.1016/j.jksuci.2024.102048.