

Analisis Sentimen Program Makan Bergizi Gratis Menggunakan Random Forest dan Support Vector Machine dengan Penyeimbangan Data Berbasis SMOTE

Slamet Endro Prianto^{*1}, Berlilana², Rujianto Eko Saputro³

¹Magister Ilmu Komputer, Fakultas Ilmu Komputer, Universitas Amikom Purwokerto, Indonesia

^{2,3}Fakultas Ilmu Komputer, Universitas Amikom Purwokerto, Indonesia

Email: ¹slametendrop@gmail.com, ²berlilana@amikompurwokerto.ac.id,

³rujianto@amikompurwokerto.ac.id

Abstrak

Sentimen publik yang diekspresikan melalui media sosial telah menjadi sumber data penting untuk mengevaluasi kebijakan publik, termasuk Program Makan Bergizi Gratis (MBG). Namun, analisis sentimen pada data teks terkait kebijakan menghadapi tantangan utama karena fitur tekstual berdimensi tinggi, kompleksitas linguistik, serta distribusi kelas yang tidak seimbang. Permasalahan ini dapat menyebabkan kinerja model menjadi bias dan menghasilkan interpretasi kebijakan yang kurang akurat apabila tidak ditangani dengan tepat. Penelitian ini bertujuan untuk mengevaluasi kinerja Random Forest (RF) dan Support Vector Machine (SVM) dalam mengklasifikasikan sentimen publik terhadap MBG serta menilai efektivitas SMOTE dalam menangani ketidakseimbangan kelas. Penelitian ini menggunakan pendekatan kuantitatif dengan 2.349 tweet yang dikumpulkan dari platform X (Twitter). Data dipraproses melalui tahapan pembersihan teks (text cleaning), case folding, tokenisasi, normalisasi, penghapusan stopword, dan stemming. Ekstraksi fitur dilakukan menggunakan TF-IDF. Klasifikasi sentimen dilakukan menggunakan model RF dan SVM pada kondisi data tidak seimbang dan data yang telah diseimbangkan. Kinerja model dievaluasi menggunakan metrik akurasi, presisi, recall, dan macro-average F1-score. Hasil penelitian menunjukkan bahwa baik RF maupun SVM menghasilkan kinerja klasifikasi yang relatif rendah, dengan nilai akurasi dan macro-average F1-score berkisar antara 0,32 hingga 0,34. Penerapan SMOTE mampu memperbaiki distribusi kelas dan mengurangi bias prediksi, namun peningkatan kinerja secara keseluruhan masih bersifat marginal. Untuk data teks, SMOTE tidak selalu meningkatkan representasi makna secara signifikan, karena proses oversampling menghasilkan data sintesis yang didasarkan pada ruang fitur numerik tanpa mempertimbangkan konteks semantik kalimat. SVM menunjukkan kinerja yang lebih stabil dibandingkan RF, terutama dalam menangani fitur teks yang berdimensi tinggi dan bersifat sparse. Temuan ini menunjukkan bahwa penyeimbangan data saja belum cukup untuk secara signifikan meningkatkan kinerja klasifikasi sentimen. Karakteristik data teks yang kompleks, seperti variasi bahasa informal, singkatan, dan ekspresi pendapat yang ambigu dalam media sosial, juga menyebabkan keterbatasan peningkatan kinerja. Penelitian ini menekankan pentingnya mengombinasikan teknik penyeimbangan data dengan model semantik kontekstual guna meningkatkan kualitas analisis sentimen dalam evaluasi kebijakan publik. Metode berbasis fitur seperti TF-IDF dapat memberikan representasi makna yang lebih kaya, sehingga dapat meningkatkan kapasitas model untuk memahami konteks opini publik terhadap kebijakan.

Kata kunci: Analisis Kebijakan Publik, Analisis Sentimen, Program MBG, Random Forest, Support Vector Machine.

Sentiment Analysis of the Makan Bergizi Gratis Program Using Random Forest and Support Vector Machine with SMOTE-Based Data Balancing

Abstract

Public sentiment expressed through social media has become an important data source for evaluating public policies, including the Makan Bergizi Gratis (MBG) Program. However, sentiment analysis on policy-related text data faces major challenges due to high-dimensional textual features, linguistic complexity, and imbalanced class distribution. These issues may lead to biased model performance and inaccurate policy interpretation if not properly addressed. This study aims to evaluate the performance of RF and SVM in classifying public sentiment toward MBG and to assess the effectiveness of SMOTE in handling class imbalance. This research employed a quantitative approach using 2,349 tweets collected from the X (Twitter) platform. The data were preprocessed through text cleaning, case folding, tokenization, normalization, stopword removal, and stemming. Feature

extraction was performed using TF-IDF. Sentiment classification was conducted using RF and SVM models under imbalanced and balanced data conditions. Model performance was evaluated using accuracy, precision, recall, and macro-average F1-score. The results indicate that both RF and SVM achieved relatively low classification performance, with accuracy and macro-average F1-score values ranging from 0.32 to 0.34. The application of SMOTE improved class distribution and reduced prediction bias; however, overall performance improvements remained marginal. For text data, SMOTE does not always significantly enhance meaning representation, as the oversampling process generates synthetic data based on the numerical feature space without considering the semantic context of the sentences. SVM shows more stable performance compared to RF, especially in handling high-dimensional and sparse text features. These findings indicate that data balancing alone is not sufficient to significantly improve sentiment classification performance. The complex characteristics of text data, such as variations in informal language, abbreviations, and ambiguous expressions of opinion on social media, also limit performance improvement. This research emphasizes the importance of combining data balancing techniques with contextual semantic models to improve the quality of sentiment analysis in public policy evaluation. Feature-based methods such as TF-IDF can provide a richer representation of meaning, thereby enhancing the model's capacity to understand the context of public opinion toward policies.

Keywords: MBG Program, Public Policy Analysis, Random Forest, Sentiment Analysis, Support Vector Machine.

1. PENDAHULUAN

Berbagai industri, termasuk sektor publik, menggunakan data digital sebagai bagian penting dari proses pengambilan keputusan karena kemajuan teknologi informasi dan komunikasi [1]. Pemangku kebijakan dan pemerintah semakin mengandalkan data dari platform digital untuk mengevaluasi kebijakan, meningkatkan perumusan dan implementasi kebijakan publik, dan memahami respons masyarakat [2]. Data berbasis teks memiliki nilai strategis karena menunjukkan opini dan persepsi publik secara langsung [3]. Studi di negara berkembang menunjukkan bahwa analisis data media sosial dapat menjadi alat penting untuk mengevaluasi kebijakan publik dengan cepat dan responsif [4]. Sebagai contoh, penelitian oleh Isnain et al. (2021) menemukan bahwa analisis sentimen Twitter dapat digunakan untuk mengevaluasi respons publik terhadap kebijakan pembelajaran online selama pandemi COVID-19 di Indonesia; pola sentimen publik menunjukkan bagaimana penerimaan dan penerimaan kebijakan publik secara keseluruhan diukur [5].

Terdapat masalah metodologis yang tidak sederhana saat menggunakan data teks dalam konteks kebijakan publik. Karakteristik data yang tidak terstruktur dan memiliki distribusi kelas yang tidak seimbang merupakan masalah utama [3]. Banyak kali, kategori tertentu mendominasi data, sementara kategori lain hanya muncul dalam jumlah yang relatif kecil [6]. Apabila tidak ditangani dengan benar, ketidakseimbangan ini berpotensi menurunkan kualitas pemodelan dan menghasilkan kesimpulan yang tidak akurat [7]. Penelitian oleh Anang Ma'ruf et al. (2026) menunjukkan bahwa, pada data Twitter yang berkaitan dengan kebijakan pemerintah Indonesia, distribusi sentimen seringkali tidak seimbang, dengan sentimen netral dan negatif lebih banyak daripada sentimen positif. Dengan demikian, tanpa teknik penyeimbangan data, model klasifikasi cenderung bias terhadap kelas mayoritas [8].

Untuk mengolah dan mengklasifikasikan data teks, berbagai teknik pembelajaran mesin telah digunakan [9]. Sebaliknya, banyak penelitian masih berfokus pada perbandingan kinerja algoritma tanpa memperhatikan masalah ketidakseimbangan data [10]. Metode ini mungkin menghasilkan model yang memiliki kinerja yang baik secara keseluruhan, tetapi tidak dapat memahami kelas minoritas dengan baik [11]. Dalam kebijakan publik, mengabaikan kelompok minoritas dapat menyebabkan pemahaman yang tidak akurat tentang respons masyarakat secara keseluruhan [12]. Penelitian oleh Rahman et al. (2020) menemukan hasil yang serupa. Penelitian ini menyelidiki persepsi publik terhadap kebijakan kesehatan di Bangladesh dengan menggunakan metode pembelajaran mesin. Meskipun model klasifikasi yang digunakan sangat akurat, namun ketidakseimbangan distribusi data membuatnya tidak dapat mengidentifikasi perasaan minoritas [13].

Salah satu teknik yang banyak dikembangkan untuk mengatasi permasalahan ketidakseimbangan data adalah *Synthetic Minority Over-sampling Technique* (SMOTE) [14]. Metode ini menggunakan pola data yang ada untuk menghasilkan data sintesis pada kelas minoritas, yang membuat distribusi data lebih seimbang tanpa duplikat. Dengan distribusi kelas yang lebih proporsional, diharapkan proses pembelajaran model akan lebih adil dan representatif [6]. Penelitian oleh Vania Agresia dan Ryan Randy Suryono (2025) menunjukkan bahwa SMOTE dapat memperbaiki distribusi kelas dan meningkatkan kemampuan model untuk mengidentifikasi opini minoritas, khususnya dalam hal analisis sentimen kebijakan transportasi publik digital [15].

Algoritma pembelajaran mesin seperti *Random Forest* (RF) dan *Support Vector Machine* (SVM) dikenal memiliki kapasitas yang baik untuk menangani data yang sangat besar dan kompleks [16]. Tetapi cara kedua algoritma tersebut menangani data dengan distribusi yang tidak seimbang berbeda. Oleh karena itu, evaluasi empiris diperlukan untuk memahami bagaimana kedua algoritma tersebut bekerja ketika digunakan bersama dengan teknik penyeimbangan data seperti SMOTE [17]. Dalam penelitian yang dilakukan oleh Gvantsa Gverdtsiteli (2023) tentang analisis sentimen kebijakan publik di Vietnam, SVM lebih stabil dibandingkan RF dalam menangani data teks berukuran besar. Ini terutama berlaku untuk dataset dengan tingkat *sparsity* yang tinggi, seperti data media sosial [18].

Penelitian ini bertujuan untuk mengevaluasi kinerja RF dan SVM saat menggunakan SMOTE sebagai metode penanganan ketidakseimbangan data, seperti yang diuraikan di atas. Diharapkan bahwa penelitian ini akan memberikan bantuan metodologis dalam pengolahan data teks dalam konteks kebijakan publik, khususnya dengan membantu meningkatkan kekuatan model klasifikasi dalam situasi distribusi data yang tidak seimbang. Penelitian ini juga diharapkan dapat memberikan kontribusi empiris untuk pengembangan teknik analisis sentimen untuk evaluasi kebijakan publik di negara berkembang, di mana keterbatasan data survei konvensional seringkali menjadi hambatan untuk memahami respons masyarakat secara *real-time*. Dengan memanfaatkan data opini publik dari media sosial, penelitian ini diharapkan dapat memberikan kontribusi empiris.

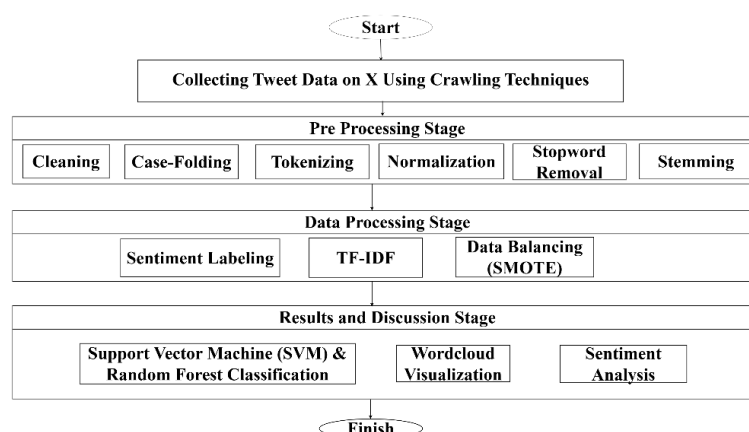
2. METODE PENELITIAN

2.1 Pengumpulan Data

Penelitian ini diawali dengan proses pengumpulan data berupa teks dari platform X (Twitter) menggunakan teknik *crawling* [19]. Data dikumpulkan berdasarkan kata kunci Program MBG yang relevan dengan topik kebijakan publik yang diteliti. Setiap data yang diperoleh berupa teks utuh yang merepresentasikan opini pengguna terhadap isu yang dibahas. Tahap ini bertujuan untuk memperoleh dataset yang mencerminkan respons publik secara alami dan aktual [7].

Data yang digunakan dalam penelitian ini mencakup 2.349 *tweet* yang dikirim dalam rentang waktu tertentu ketika masalah Program MBG menjadi subjek diskusi media sosial yang aktif. Jumlah dataset yang dipilih dianggap cukup untuk mewakili variasi opini publik yang muncul secara organik pada platform X dan cukup untuk menguji proses pelatihan dan model klasifikasi berbasis *machine learning*. Pilihan jumlah data juga mempertimbangkan keseimbangan antara proses pembersihan data yang menghasilkan teks yang dapat dianalisis dan ketersediaan data yang relevan dengan topik penelitian [20].

Pemilihan platform X (Twitter) sebagai sumber data didasarkan pada sifatnya sebagai media sosial, yang banyak digunakan untuk menyampaikan pendapat publik secara terbuka dan terkini tentang kebijakan pemerintah. Dengan demikian, data yang diperoleh diharapkan mampu menggambarkan dinamika persepsi masyarakat secara lebih aktual dibandingkan sumber data survei konvensional yang cenderung memiliki keterbatasan waktu dan cakupan responden [21]. Berikut adalah struktur proses data:



Gambar 1. Struktur Proses Data

2.2 Tahap Preprocessing

Data teks yang diperoleh selanjutnya melalui tahap *preprocessing* untuk meningkatkan kualitas data sebelum dianalisis lebih lanjut. Tahapan *preprocessing* yang diterapkan meliputi *cleaning*, *case folding*, *tokenizing*, *normalization*, *stopword removal*, dan *stemming* [22]. Proses *cleaning* dilakukan untuk menghapus

URL, *mention*, tanda baca, dan karakter non-alfabet [15]. *Case folding* bertujuan untuk menyeragamkan huruf menjadi huruf kecil [23]. Selanjutnya, teks dipecah menjadi token kata, dinormalisasi untuk mengatasi variasi penulisan, dihapus kata-kata umum yang tidak bermakna (*stopwords*) [24], serta dilakukan *stemming* untuk mengembalikan kata ke bentuk dasarnya [25]. Tahap ini bertujuan mengurangi *noise* dan meningkatkan konsistensi data teks [26]. Karena teks di platform X sering mengandung bahasa informal, singkatan, *emoji*, dan variasi penulisan yang dapat memengaruhi kualitas representasi fitur, tahapan *preprocessing* ini menjadi sangat penting untuk analisis sentimen berbasis media sosial. Oleh karena itu, sebelum digunakan dalam proses ekstraksi fitur dan klasifikasi, proses pembersihan dan normalisasi dilakukan secara teratur untuk memastikan bahwa data yang dianalisis memiliki kualitas linguistik yang lebih konsisten [27].

2.3 Tahap Data Processing

2.3.1 Sentiment Labeling

Setelah tahap *preprocessing*, data teks diberi label sentimen ke dalam tiga kelas, yaitu *negatif*, *neutral*, dan *positif* [28]. Proses pelabelan dilakukan untuk menyediakan kelas target yang akan digunakan dalam proses pembelajaran mesin. Distribusi kelas pada tahap ini menunjukkan adanya ketidakseimbangan data, yang umum terjadi pada data opini publik berbasis media sosial [29]. Pelabelan sentimen memanfaatkan konteks kalimat secara keseluruhan, sehingga setiap *tweet* dapat dikategorikan berdasarkan kecenderungan pendapat yang disampaikan pengguna. Pada saat ini, ketidakseimbangan distribusi kelas yang ditemukan mencerminkan fenomena umum dalam data media sosial: pendapat tertentu sering lebih dominan daripada pendapat lainnya [25].

2.3.2 Feature Extraction (TF-IDF)

Data teks yang telah diberi label kemudian ditransformasikan ke dalam bentuk numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Metode ini digunakan untuk merepresentasikan tingkat kepentingan suatu kata dalam sebuah dokumen relatif terhadap keseluruhan korpus. Representasi TF-IDF menghasilkan vektor berdimensi tinggi yang digunakan sebagai masukan bagi model klasifikasi [30]. Dalam analisis sentimen, representasi fitur berbasis TF-IDF umumnya digunakan karena metode ini mampu menangkap pola distribusi kata dalam korpus teks secara efektif, selain memberikan bobot yang lebih tinggi pada kata-kata yang berperan penting dalam membedakan dokumen [31].

2.3.3 Data Balancing (SMOTE)

Untuk mengatasi permasalahan ketidakseimbangan kelas, penelitian ini menerapkan SMOTE. SMOTE digunakan untuk menghasilkan data sintetis pada kelas minoritas berdasarkan karakteristik data pada ruang fitur. Penerapan SMOTE dilakukan hanya pada data pelatihan setelah proses pembagian data, guna menghindari *data leakage* dan menjaga validitas evaluasi model [32]. Tujuan implementasi SMOTE pada tahap data pelatihan adalah untuk meningkatkan representasi kelas minoritas sehingga model pembelajaran mesin tidak terlalu bias terhadap kelas mayoritas. Dengan distribusi data yang lebih seimbang, diharapkan model dapat mempelajari pola klasifikasi yang lebih representatif untuk setiap kelas sentimen [27].

2.4 Tahap Klasifikasi dan Analisis

2.4.1 Klasifikasi

Tahap klasifikasi dilakukan menggunakan dua algoritma pembelajaran mesin, yaitu SVM dan RF. SVM digunakan karena kemampuannya dalam menangani data berdimensi tinggi dan membentuk batas keputusan yang optimal [33]. Sementara itu, RF digunakan sebagai metode berbasis *ensemble* yang menggabungkan banyak pohon keputusan untuk meningkatkan stabilitas dan akurasi prediksi [30]. Kedua model dilatih menggunakan data pelatihan hasil SMOTE dan diuji menggunakan data pengujian asli [34]. Perbandingan kinerja antara SVM dan RF diharapkan dapat memberikan gambaran empiris tentang kemampuan masing-masing algoritma dalam menangani distribusi kelas yang tidak seimbang dan data teks yang besar, karena kedua algoritma tersebut digunakan secara luas dalam penelitian klasifikasi teks [35].

2.4.2 Visualisasi dan Analisis Sentimen

Sebagai bagian dari analisis hasil, penelitian ini juga menyajikan *visualisasi wordcloud* untuk menggambarkan kata-kata dominan pada masing-masing kelas sentimen [27]. Selain itu, dilakukan analisis

sentimen untuk menginterpretasikan pola klasifikasi yang dihasilkan oleh model, baik sebelum maupun sesudah penerapan SMOTE [36]. Evaluasi kinerja model dilakukan menggunakan *metrik accuracy, precision, recall, dan F1-score*, dengan penekanan pada *macro-average F1-score* untuk menilai performa model secara seimbang pada seluruh kelas sentimen [37]. Dalam penelitian ini, *macro-average F1-score* dipilih karena metrik tersebut dapat memberikan penilaian yang lebih adil pada dataset dengan distribusi kelas yang tidak seimbang, sehingga kinerja model dapat dievaluasi secara menyeluruh untuk setiap kategori sentimen [38].

3. HASIL DAN PEMBAHASAN

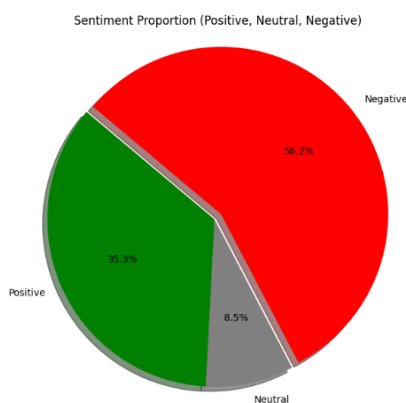
3.1 Karakteristik Data dan Distribusi Sentimen

Data penelitian diperoleh dari platform media sosial X (Twitter) melalui teknik *crawling* menggunakan kata kunci yang berkaitan langsung dengan Program MBG [19]. Setelah proses *preprocessing* dan pelabelan sentimen dari total 2.349 *tweet*, sebanyak 1.644 *tweet* digunakan sebagai data latih, sedangkan 705 *tweet* digunakan sebagai data uji dengan rasio 70:30 [9]. Setiap *tweet* diklasifikasikan ke dalam tiga kelas sentimen, yaitu negatif, netral, dan positif, yang merepresentasikan persepsi publik terhadap kebijakan MBG [39].

Table 1. Kriteria Pelabelan

Sampel Data	Full_text	Sentimen
8	trnyata yg janji libur nataru hnyalah jajan pasargk sampek gocap	Negatif
94	gizi aman pelan-pelan cerah mbg anakanak asupan sehat bikin badan kuat otak jalan	Netral
55	dukung program makan gizi gratis bentuk sdm indonesia unggul sehat bangga	Positif

Hasil analisis distribusi sentimen menunjukkan adanya ketidakseimbangan kelas yang cukup signifikan. Sentimen negatif mendominasi sebesar 56,2%, diikuti oleh positif sebesar 35,3%, sementara sentimen netral hanya sebesar 8,5%. Dominasi sentimen negatif mengindikasikan bahwa wacana publik terkait MBG lebih banyak diwarnai oleh kritik, kekhawatiran, dan penilaian negatif terhadap implementasi program [6] [40]. Kondisi ini berpotensi menimbulkan bias pada proses klasifikasi sentimen dan menjadi tantangan utama dalam analisis sentimen multikelas berbasis data media sosial [41].



Gambar 2. Proporsi Sentimen

3.2 Data Preprocessing

Tahap *data preprocessing* bertujuan untuk meningkatkan kualitas data teks sebelum dilakukan ekstraksi fitur dan klasifikasi. Proses ini diawali dengan *text cleaning* untuk menghapus elemen yang tidak relevan, seperti URL, *mention*, *hashtag*, angka, simbol, dan tanda baca berlebih yang tidak memiliki kontribusi semantik terhadap analisis sentimen MBG [11] [42].

Selanjutnya dilakukan *case folding* untuk menyeragamkan teks menjadi huruf kecil, diikuti dengan tokenisasi, normalisasi kata, *stopword removal*, dan *stemming*. Tahapan ini menghasilkan representasi teks yang lebih ringkas dan berfokus pada kata-kata bermakna, seperti “gizi”, “makan”, “gratis”, dan “program”, yang merupakan istilah kunci dalam diskursus kebijakan MBG [4] [23]. Dengan demikian, proses *preprocessing*

berperan penting dalam memastikan bahwa fitur yang diekstraksi benar-benar merepresentasikan opini publik terhadap substansi program.

Table 2. Data Preprocessing

	cleaned_text	case_folded_text	Normalized	final_text
	<pre>def cleaning(text): if not isinstance(text, str): return "" text = re.sub(r'http\S+ www\S+ https\S+', '', text) text = re.sub(r'@w+ #w+', '', text) text = re.sub(r'd+', '', text) text = re.sub(r'^a-zA-Z\s ', '', text) text) text.strip().lower()</pre>	<pre>Def case_folded_text(text): if not isinstance(text, str): return "" text = re.sub(r'http\S+ www\S+ https\S+', '', text) text = re.sub(r'@w+ #w+', '', text) text = re.sub(r'd+', '', text) text = re.sub(r'^a-zA-Z\s ', '', text) text) text.strip().lower()</pre>	<pre>def normalize_text(text, kamus): if not isinstance(text, str): return "" words = text.split() return ' '.join([kamus.get(word, word) for word in words])</pre>	<pre>def final_text(text, stopwords_set, stemmer_obj): if not words = [word for word in text.split() if word not in stopwords_set] stemmed_words = [stemmer_obj.stem(word) for word in words] return ' '.join(stemmed_words)</pre>
Sampel 1	Hasil Cleaning test Gizi seimbang mendukung pertumbuhan kecerdasan dan kesehatan generasi muda Papua	Hasil case_folded_text gizi seimbang mendukung pertumbuhan kecerdasan dan kesehatan generasi muda papua	Hasil Normalized gizi seimbang mendukung pertumbuhan kecerdasan dan kesehatan generasi muda papua	Hasil final_text gizi imbang dukung tumbuh cerdas sehat generasi muda papua
Sampel 2	Program MBG aman dan transparan di Kamar karena Presiden Prabowo tuduhan korupsi dipatahkan hak murid tetap terjaga	program mbg aman dan transparan di kamar karena presiden prabowo tuduhan korupsi dipatahkan hak murid tetap terjaga	program mbg aman dan transparan di kamar karena presiden prabowo tuduhan korupsi dipatahkan hak murid tetap terjaga	program mbg aman transparan kamar presiden prabowo tuduh korupsi patah hak murid jaga

3.3 Ekstraksi Fitur Menggunakan TF-IDF

Data teks hasil *preprocessing* kemudian ditransformasikan ke dalam bentuk numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Metode ini digunakan untuk mengukur tingkat kepentingan suatu kata dalam sebuah *tweet* relatif terhadap keseluruhan korpus diskursus MBG [7].

Hasil ekstraksi fitur menghasilkan sebanyak 4.149 fitur, yang mencerminkan karakteristik data teks media sosial yang berdimensi tinggi dan bersifat *sparse* [11]. Representasi TF-IDF memungkinkan model klasifikasi untuk menangkap pola linguistik yang membedakan antara sentimen negatif, netral, dan positif terhadap kebijakan MBG, meskipun belum mampu sepenuhnya merepresentasikan konteks dan nuansa bahasa secara semantik [27] [43].

	aamiin	aang	abad	abadi	abah	abai	abalabal	abang	abcdz	abis	...	\
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	
3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	
4	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	...	

	yulia	yulianto	yusuf	zaman	zodiak	zombie	zonasi	zonauang	zulhas	\
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	
4	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	

Gambar 3. TF-IDF

3.4 Hasil Klasifikasi

3.4.1 RF dan SVM

Hasil pengujian menunjukkan bahwa model RF menghasilkan akurasi sebesar 33% (0,33) pada data uji. Kelas netral (1) mencatat nilai *recall* tertinggi sebesar 0,44, mengindikasikan bahwa model cenderung memetakan opini publik MBG ke dalam kategori netral. Namun, nilai *precision* yang rendah (0,35) dan *F1-score* sebesar 0,39 menunjukkan tingginya kesalahan klasifikasi lintas kelas [6] [7] [11].

Sebaliknya, kelas positif (2) menunjukkan performa terlemah dengan nilai *recall* 0,25 dan *F1-score* 0,27, yang menandakan bahwa banyak dukungan terhadap MBG salah diklasifikasikan sebagai sentimen lain. Kelas negatif (0) juga memperlihatkan kinerja yang rendah dengan *recall* 0,29 dan *F1-score* 0,32. Nilai *macro-average F1-score* sebesar 0,33 menegaskan bahwa RF belum mampu mengklasifikasikan sentimen MBG secara seimbang [28] [29].

Model SVM menghasilkan akurasi sedikit lebih tinggi, yaitu 34% (0,34). Kinerja SVM relatif lebih merata pada seluruh kelas sentimen MBG, dengan nilai *F1-score* berada pada rentang 0,33–0,36. Kelas Netral kembali menjadi kelas dengan performa paling stabil, sementara kelas negatif dan positif masih menunjukkan keterbatasan dalam pengenalan pola sentimen secara konsisten. Nilai *macro-average F1-score* sebesar 0,34 menunjukkan bahwa SVM memiliki kemampuan generalisasi yang sedikit lebih baik dibandingkan RF, meskipun tetap belum memadai [28] [41].

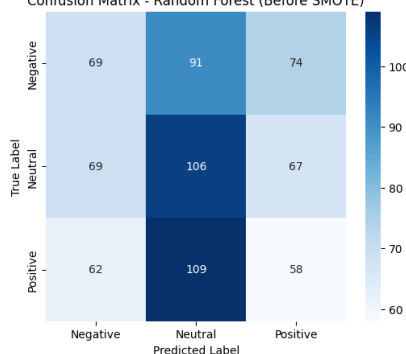
Table 3. RF Sebelum SMOTE

	Precision	Recall	F1-Score
0	0.34	0.29	0.32
1	0.35	0.44	0.39
2	0.29	0.25	0.27
Accuracy			0.33
macro avg	0.33	0.33	0.33
Weighted avg	0.33	0.33	0.33

Table 4. SVM Sesudah SMOTE

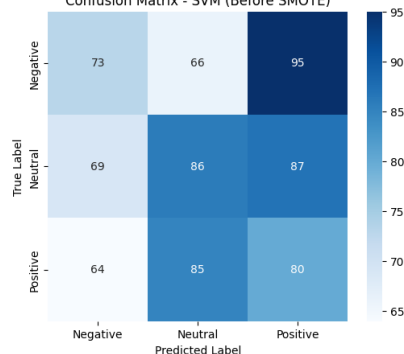
	Precision	Recall	F1-Score
0	0.35	0.31	0.33
1	0.36	0.36	0.36
2	0.31	0.35	0.33
Accuracy			0.34
macro avg	0.34	0.34	0.34
weighted avg	0.34	0.34	0.34

Confusion Matrix - Random Forest (Before SMOTE)



Gambar 4. Confusion Matrix RF Sebelum SMOTE

Confusion Matrix - SVM (Before SMOTE)



Gambar 5. Confusion Matrix SVM Sesudah SMOTE

Setelah penerapan SMOTE, distribusi prediksi antar kelas sentimen MBG menjadi lebih seimbang. Pada model RF, akurasi tercatat sebesar 32% (0,32), relatif stabil dibandingkan sebelum SMOTE. Peningkatan paling nyata terlihat pada kelas Netral (1) dengan *recall* meningkat menjadi 0,43. Namun, kelas Positif (2) tetap menunjukkan performa terendah dengan *F1-score* sebesar 0,27, yang menandakan bahwa SMOTE belum sepenuhnya mampu memperbaiki pengenalan sentimen dukungan terhadap MBG [11] [4].

Model SVM setelah SMOTE menunjukkan kinerja yang lebih stabil dengan akurasi tetap sebesar 34% (0,34) dan *macro-average F1-score* sebesar 0,34. Distribusi prediksi antar kelas menjadi lebih proporsional, meskipun nilai *precision* dan *recall* masih berada pada kisaran rendah. Temuan ini menunjukkan bahwa SMOTE efektif dalam mengurangi bias distribusi data, tetapi belum cukup untuk meningkatkan kemampuan diskriminatif model secara signifikan dalam konteks sentimen MBG [10].

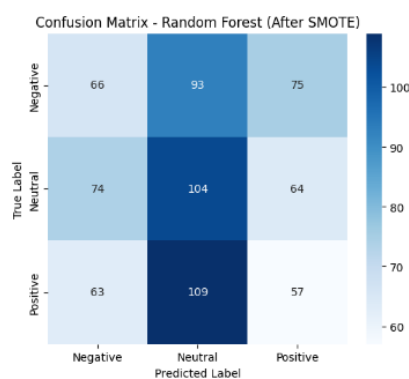
SMOTE juga memiliki keterbatasan konseptual saat menganalisis sentimen kebijakan publik. Dengan menggunakan interpolasi pada ruang fitur numerik, data sintetis SMOTE tidak mempertimbangkan konteks semantik atau makna linguistik teks. Akibatnya, representasi makna opini publik terhadap kebijakan MBG mungkin tidak meningkat secara signifikan meskipun distribusi kelas secara statistik menjadi lebih seimbang. Kondisi ini menunjukkan bahwa untuk menangkap nuansa opini publik dengan lebih akurat, pendekatan penyeimbangan data berbasis numerik harus dikombinasikan dengan teknik pemodelan bahasa yang lebih kontekstual [27].

Table 5. RF Sebelum SMOTE

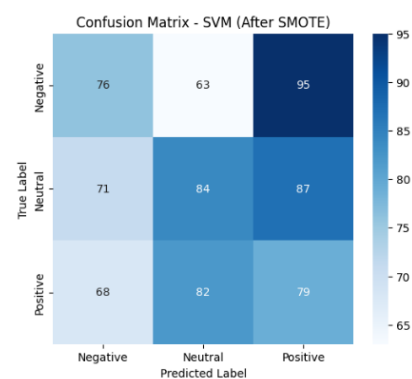
	Precision	Recall	F1-Score
0	0.33	0.28	0.30
1	0.34	0.43	0.38
2	0.29	0.25	0.27
Accuracy			0.32
macro avg	0.32	0.32	0.32
weighted avg	0.32	0.32	0.32

Table 6. SVM Setelah SMOTE

	Precision	Recall	F1-Score
0	0.35	0.32	0.34
1	0.37	0.35	0.36
2	0.30	0.34	0.32
Accuracy			0.34
macro avg	0.34	0.34	0.34
weighted avg	0.34	0.34	0.34



Gambar 6. Confusion Matrix RF Sebelum SMOTE



Gambar 7. Confusion Matrix SVM Setelah SMOTE

Perbandingan Kinerja Model

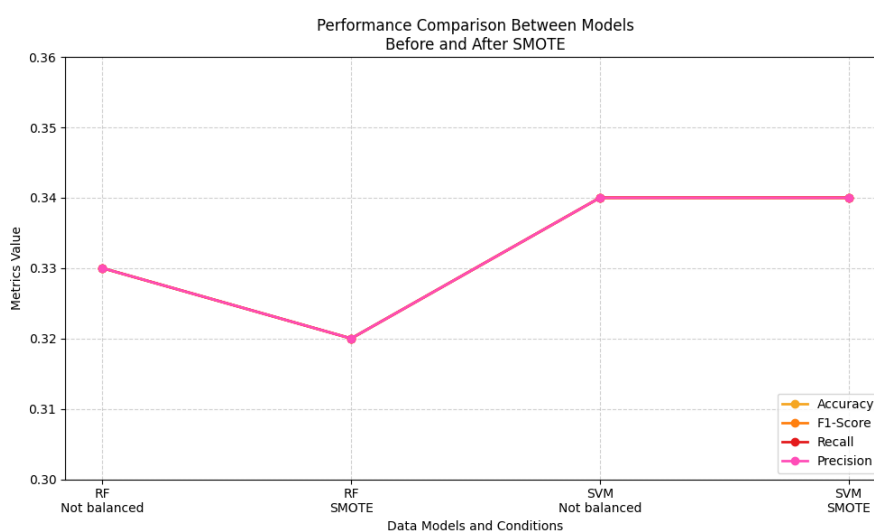
Hasil evaluasi menunjukkan bahwa SVM lebih stabil dibandingkan RF dalam mengklasifikasikan sentimen publik terhadap Program MBG. Pada kondisi sebelum SMOTE, kedua model menunjukkan performa rendah (akurasi 0,33–0,34; *macro-average F1-score* $\leq$ 0,34), menandakan ketidakmampuan dalam membedakan sentimen MBG secara seimbang akibat distribusi data yang timpang [6] [4]. Penerapan SMOTE memperbaiki distribusi prediksi antar kelas, khususnya pada kelas Netral, namun peningkatan kinerja masih terbatas. RF mengalami penurunan akurasi (0,32) dengan *macro-average F1-score* tetap rendah, sedangkan SVM mampu mempertahankan akurasi (0,34) dan menunjukkan prediksi yang lebih konsisten antar kelas. Temuan ini mengindikasikan bahwa untuk analisis sentimen kebijakan publik seperti MBG, SVM lebih adaptif terhadap data

hasil penyeimbangan, sementara RF memerlukan optimasi lanjutan agar mampu menangkap kompleksitas opini publik secara lebih efektif [4] [15].

Hasil penelitian ini menunjukkan bahwa analisis sentimen kebijakan publik merupakan tantangan utama. Ini bukan hanya masalah distribusi data yang tidak seimbang, tetapi juga bahasa yang digunakan masyarakat dalam media sosial. Komentar publik sering mengandung ironi, sindiran, atau konteks sosial tertentu yang sulit diwakili hanya melalui pendekatan berbasis frekuensi kata seperti TF-IDF. Oleh karena itu, dapat dilakukan untuk meningkatkan kinerja model di masa mendatang dengan memasukkan metode pembelajaran representasi bahasa yang lebih semantik, seperti penerapan kata atau model berbasis transformasi yang dapat mengidentifikasi lebih akurat hubungan antar kata [21].

Table 7. Perbandingan Kinerja Algoritma

Algoritma	Kondisi Data	Accuracy	Precision	Recall	F1-Score
RF	Tidak Seimbang	0.33	0.33	0.33	0.33
RF	Seimbang (SMOTE)	0.32	0.32	0.32	0.32
SVM	Tidak Seimbang	0.34	0.34	0.34	0.34
SVM	Seimbang (SMOTE)	0.34	0.34	0.34	0.34



Gambar 8. Perbandingan Kinerja Antara Model Sebelum dan Sesudah SMOTE

3.5 Analisis Wordcloud

Analisis *wordcloud* memberikan gambaran kualitatif terhadap wacana publik terkait MBG. *Wordcloud* sentimen *positif* didominasi oleh kata-kata yang mencerminkan optimisme terhadap tujuan sosial program, seperti “gizi”, “makan”, “gratis”, dan “sekolah”. Sebaliknya, *wordcloud* sentimen *negatif* menampilkan kata-kata yang berkaitan dengan kekhawatiran implementasi dan efektivitas kebijakan, seperti “mbg”, “program”, “gizi”, dan “gratis”. Temuan ini menegaskan bahwa perdebatan publik mengenai MBG tidak hanya bersifat kuantitatif, tetapi juga menunjukkan perbedaan fokus wacana antara harapan sosial dan kekhawatiran operasional [4] [6] [11].



Gambar 9. Wordcloud Positif



Gambar 10. Wordcloud Negatif

3.6 Ketidakseimbangan Data dan Keterbatasan Evaluasi Berbasis Akurasi

Hasil penelitian yang didukung oleh [6] [41] menegaskan bahwa permasalahan utama dalam klasifikasi sentimen Program MBG terletak pada distribusi kelas yang tidak seimbang dan keterbatasan evaluasi berbasis akurasi. Meskipun RF dan SVM menghasilkan nilai akurasi yang relatif stabil pada kisaran 0,32–0,34, nilai *macro-average F1-score* yang rendah menunjukkan bahwa kedua model belum mampu membedakan kelas sentimen secara merata [11]. Temuan ini mengonfirmasi bahwa akurasi bukan indikator yang memadai dalam analisis sentimen kebijakan publik, karena dapat menyamarkan kegagalan model dalam mengenali sentimen minoritas [11] [12]. Dalam konteks MBG, kondisi ini berpotensi mengaburkan masukan publik yang relevan terhadap aspek implementasi dan akuntabilitas program [12].

3.7 Efektivitas SMOTE dan Implikasi Pemodelan Sentimen MBG

Penerapan SMOTE terbukti efektif dalam memperbaiki keseimbangan distribusi data dan mengurangi bias prediksi antar kelas. Namun, peningkatan performa klasifikasi yang dihasilkan masih bersifat marginal, khususnya pada RF, yang menunjukkan keterbatasan dalam membangun batas keputusan yang jelas pada data teks berbasis TF-IDF [9] [41]. SVM menunjukkan kinerja yang lebih stabil sebelum dan sesudah SMOTE, mengindikasikan bahwa pendekatan margin-based lebih adaptif terhadap data berdimensi tinggi dan sparse. Meskipun demikian, hasil ini juga menegaskan bahwa penyeimbangan data saja belum cukup untuk meningkatkan kemampuan diskriminatif model secara signifikan. Oleh karena itu, analisis sentimen kebijakan MBG memerlukan kombinasi penyeimbangan data dan pendekatan pemodelan yang lebih kontekstual agar mampu merepresentasikan kompleksitas bahasa dan wacana publik secara lebih akurat [4] [11].

Keterbatasan SMOTE juga menunjukkan bahwa pendekatan statistik semata belum cukup untuk memahami dinamika opini masyarakat yang kompleks dan berbagai macam dalam evaluasi kebijakan publik [27]. Faktor-faktor sosial, ekonomi, dan politik, termasuk cara orang menggunakan media sosial, memengaruhi pandangan publik tentang kebijakan seperti Program MBG [44]. Oleh karena itu, dalam penelitian yang akan datang, pengembangan model analisis sentimen harus mempertimbangkan penggunaan metode pemrosesan bahasa alami yang lebih kontekstual. Teknik-teknik seperti model pembelajaran mendalam berbasis transformasi atau representasi semantik berbasis embedding kata harus dipertimbangkan untuk mendapatkan pemahaman yang lebih mendalam tentang makna implisit dan hubungan semantik antar kata dalam diskursus kebijakan publik.

4. KESIMPULAN

Studi ini menganalisis sentimen publik terhadap Program MBG dengan menggunakan metode *Random Forest* (RF) dan *Support Vector Machine* (SVM) pada data teks dari platform X. Hasil penelitian menunjukkan bahwa kedua model memiliki kinerja yang terbatas, dengan nilai akurasi dan *macro-average F1-score* berada pada kisaran 0,32 hingga 0,34, yang mengindikasikan adanya tantangan dalam klasifikasi sentimen multi-kelas. Ketidakseimbangan kelas teridentifikasi sebagai faktor krusial yang memengaruhi efektivitas model. Meskipun penerapan SMOTE mampu memperbaiki distribusi kelas dan mengurangi bias prediksi, dampaknya terhadap peningkatan kinerja secara keseluruhan relatif marginal. SVM menunjukkan perilaku yang lebih stabil dibandingkan RF, khususnya dalam menangani fitur TF-IDF yang berdimensi tinggi dan bersifat jarang (*sparse*), namun kedua model belum mampu sepenuhnya menangkap kompleksitas kontekstual dalam wacana terkait MBG.

Secara metodologis, studi ini memberikan kontribusi empiris dengan menunjukkan bahwa penyeimbangan data semata tidak cukup untuk meningkatkan kinerja analisis sentimen pada teks kebijakan publik, serta menegaskan pentingnya penggunaan metrik evaluasi yang berorientasi pada keseimbangan kelas, seperti *macro-average F1-score*. Dari perspektif kebijakan, sentimen minoritas merupakan sinyal penting terkait implementasi program dan akuntabilitas, sehingga perlu diintegrasikan dalam proses evaluasi MBG. Penelitian selanjutnya disarankan untuk mengombinasikan teknik penyeimbangan data dengan model semantik kontekstual guna meningkatkan robustitas analisis.

Temuan penelitian ini memiliki konsekuensi praktis untuk pengembangan sistem yang memanfaatkan analisis opini publik dalam proses pengambilan keputusan kebijakan. Pemangku kebijakan dapat menggunakan analisis sentimen media sosial sebagai sumber awal untuk mengevaluasi persepsi masyarakat terhadap pelaksanaan program MBG. Dengan memahami pola sentimen publik dengan lebih baik, pemerintah atau lembaga terkait dapat membuat rencana komunikasi kebijakan yang lebih sesuai dengan kebutuhan masyarakat.

Hasil penelitian ini dapat digunakan sebagai referensi dalam pengajaran analisis data kebijakan berbasis pembelajaran mesin. Penelitian ini khususnya menekankan betapa pentingnya pengolahan data tidak seimbang, memilih model klasifikasi yang tepat, dan menggunakan teknik pemrosesan bahasa alami untuk memperoleh

pemahaman yang lebih baik tentang pendapat publik. Hasil-hasil ini dapat digunakan dalam konteks pendidikan dan pengembangan kapasitas analitik data. Metode ini dapat membantu membangun sistem pemantauan kebijakan publik yang lebih adaptif dan berbasis bukti (*data-driven policy evaluation*).

DAFTAR PUSTAKA

- [1] P. T. Ogan, O. I. District, and A. H. Fahlevi, "Implementation of Good Corporate Governance (GCG) Principles in PDAM Tirta Ogan, Ogan Ilir District," in *IAPA International Conference 2024 Towards World Class Bureaucracy*, 2024, pp. 187–405. doi: 10.23920/jphp.v1i2.292.1.
- [2] R. Hidayat and D. J. Ratnaningsih, "Analisis Sentimen Program Makanan Bergizi Gratis Menggunakan Algoritma Random Forest dan Naive Bayes," *J. Comput. Informatics Res.*, vol. 5, no. 1, pp. 395–400, 2025, doi: 10.47065/comforch.v5i1.2355.
- [3] P. Arsi and R. Waluyo, "Analisis Sentimen Wacana Pemindahan Ibu Kota Indonesia Menggunakan Algoritma Support Vector Machine (SVM)," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 8, no. 1, p. 147, 2021, doi: 10.25126/jtiik.0813944.
- [4] M. Azhari and P. Parjito, "Analisis Sentimen Opini Publik Program Makan Siang Gratis dengan Random Forest Pada Media," *Build. Informatics, Technol. Sci.*, vol. 6, no. 3, pp. 1932–1942, 2024, doi: 10.47065/bits.v6i3.6423.
- [5] A. R. Isnain, N. S. Marga, and D. Alita, "Sentiment Analysis Of Government Policy On Corona Case Using Naive Bayes Algorithm," *IJCCS (Indonesian J. Comput. Cybern. Syst.)*, vol. 15, no. 1, p. 55, 2021, doi: 10.22146/ijccs.60718.
- [6] Fatkhurrohman, B. I. Nugroho, and N. Fadillah, "Analisis Sentimen Program Makan Bergizi Gratis Pemerintah RI Melalui Twitter Menggunakan Metode SVM," *J. Artif. Intell. Digit. Bus.*, vol. 4, no. 3, pp. 3906–3917, 2025, doi: 10.31004/riggs.v4i3.2533.
- [7] P. Zakiyah, K. Umam, and A. A. Mahfudh, "Public Opinion on The MBG Program : Comparative Evaluation of InSet and VADER Lexicon Labeling Using SVM on Platform X," vol. 9, no. 6, pp. 3937–3944, 2025, doi: <https://doi.org/10.30871/jaic.v9i6.9978>.
- [8] A. Ma'ruf, A. E. Pajri, and P. Liana, "Sentiment Analysis of President Prabowo's Performance on Twitter (X) with a Comparative Study of SVM, XGBoost, and AdaBoost," *J. Appl. Informatics Comput.*, vol. 10, no. 1, pp. 684–697, 2026, doi: <https://doi.org/10.30871/jaic.v10i1.12138>.
- [9] A. Sitanggang, Y. Umaidah, Y. Umaidah, R. I. Adam, and R. I. Adam, "Analisis Sentimen Masyarakat Terhadap Program Makan Siang Gratis Pada Media Sosial X Menggunakan Algoritma Naive Bayes," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, 2024, doi: 10.23960/jitet.v12i3.4902.
- [10] S. Shafwah and H. Al Azies, "Komparasi SVM dan IndoBERT dalam Klasifikasi Sentimen Program Makanan Bergizi Gratis," vol. 10, no. 2, pp. 105–112, 2025, doi: 10.31544/jtera.v10.i2.2025.105-112.
- [11] S. C. Amaria and N. Hidayati, "Analisis Program Makan Bergizi Gratis Dengan Support Vector Machine (SVM) Pada Aplikasi X," *J. Sist. Inf.*, vol. 12, no. 2, pp. 16–24, 2025, doi: <https://doi.org/10.30656/jsii.v12i2.10708>.
- [12] R. Saputra and F. N. Hasan, "Analisis Sentimen Terhadap Program Makan Siang & Susu Gratis Menggunakan Algoritma Naive Bayes," *J. Teknol. Dan Sist. Inf. Bisnis*, vol. 6, no. 3, pp. 411–419, Jul. 2024, doi: 10.47233/jteksis.v6i3.1378.
- [13] M. M. Rahman, N. I. Khan, I. H. Sarker, M. Ahmed, and M. N. Islam, "Leveraging Machine Learning to Analyze Sentiment from COVID-19 Tweets: A Global Perspective," *Eng. Reports*, vol. 5, no. 3, pp. 1–23, 2023, doi: 10.1002/eng2.12572.
- [14] Y. Muhamad, I. Mahendra, and A. Faqih, "SMOTE untuk Meningkatkan Performa Naive Bayes dan Random Forest dalam Analisis Sentimen aplikasi Digitalent," *J. Din. Inform. Vol.*, vol. 14, no. 2, pp. 13–25, 2025, doi: <https://doi.org/10.31316/jdi.v14i2.347>.
- [15] V. Agresia and R. R. Suryono, "Comparison of SVM, Naive Bayes, and Logistic Regression Algorithms for Sentiment Analysis of Fraud and Bots in Purchasing Concert Ticket," *J. Inovtek Polbeng*, vol. 10, no. 2, pp. 591–602, 2025, doi: <https://doi.org/10.35314/npfydh47>.
- [16] A. N. Syafia, M. F. Hidayattullah, and W. Suteddy, "Studi Komparasi Algoritma SVM Dan Random Forest Pada Analisis Sentimen Komentar Youtube BTS," *J. Pengemb. IT*, vol. 8, no. 3, pp. 207–212, 2023, doi: <https://doi.org/10.30591/jpit.v8i3.5064>.
- [17] H. Ali, N. Hendrastuty, C. Science, and U. T. Indonesia, "Comparison of Naive Bayes Classifier,

- Support Vector Machine, Random Forest Algorithms for Public Sentiment Analysis of Kip-K Program on Twitter,” *J. Tek. Inform.*, vol. 5, no. 6, pp. 1701–1712, 2024, doi: <https://doi.org/10.52436/1.jutif.2024.5.6.4030>.
- [18] G. Gverdtiteli, “Authoritarian Environmentalism in Vietnam: the Construction of Climate Change as a Security Threat,” *Environ. Sci. Policy*, vol. 140, no. December 2022, pp. 163–170, 2023, doi: 10.1016/j.envsci.2022.12.004.
- [19] A. Hizqil and Y. Ruldeviani, “Sentiment analysis of online licensing service quality in the energy and mineral resources sector of the Republic of Indonesia,” *Comput. Sci. Inf. Technol.*, vol. 5, no. 1, pp. 63–71, 2024, doi: 10.11591/csit.v5i1.pp63-71.
- [20] Asro, Sudaryono, and A. Sulaiman, “Analisis Sentimen tentang Transformasi Program Makan Siang menjadi Makan Bergizi Gratis menggunakan Logistik Regression pada laman Youtube,” *J. ICT Inf. Commun. Technol.*, vol. 24, no. 1, pp. 87–94, 2024, [Online]. Available: <https://ejournal.ikmi.ac.id/index.php/jict-ikmi>
- [21] N. R. Febriyanti, K. Kusriani, and A. D. Hartanto, “Analisis Perbandingan Algoritma SVM, Random Forest dan Logistic Regression untuk Prediksi Stunting Balita,” *Edumatic J. Pendidik. Inform.*, vol. 9, no. 1, pp. 149–158, 2025, doi: 10.29408/edumatic.v9i1.29407.
- [22] I. N. Amalina and W. M. Ashari, “Analysis of the Performance Comparison Between Random Forest and SVM RBF in Detecting Cyberbullying on Imbalanced Data With the SMOTE Approach,” *Sist. J. Sist. Inf.*, vol. 14, no. 6, pp. 2768–2778, 2025, doi: <https://doi.org/10.32520/stmsi.v14i6.5574>.
- [23] D. T. Attaullah and D. Soyusiawaty, “Analisis Sentimen Program Makan Siang Gratis di Twitter/X menggunakan Metode BI-LSTM,” *Edumatic J. Pendidik. Inform.*, vol. 9, no. 1, pp. 294–303, Apr. 2025, doi: 10.29408/edumatic.v9i1.29725.
- [24] P. P. Putra, M. K. Anam, A. S. Chan, A. Hadi, and N. Hendri, “Optimizing Sentiment Analysis on Imbalanced Hotel Review Data Using SMOTE and Ensemble Machine Learning Techniques,” vol. 6, no. 2, pp. 936–951, 2025, doi: <https://doi.org/10.47738/jads.v6i2.618>.
- [25] I. Hermansyah and M. S. Hasibuan, “Sentiment Analysis of Free Nutritious Meal Programme on Social Media X using Linear Regression and Random Forest Algorithms,” *PIKSEL Penelit. Ilmu Komput. Sist. Embed. Log.*, vol. 13, no. 1, pp. 55–68, 2025, doi: 10.33558/piksel.v13i1.10633.
- [26] M. G. Hussain, B. Sultana, M. Rahman, and M. R. Hasan, “Comparison analysis of Bangla news articles classification using support vector machine and logistic regression,” *Telkomnika (Telecommunication Comput. Electron. Control.)*, vol. 21, no. 3, pp. 584–591, 2023, doi: 10.12928/TELKOMNIKA.v21i3.23416.
- [27] A. A. Rohman, A. G. Trisnapradika, and K. Kunci, “Perbandingan Algoritma NBC, SVM, Logistic Regression untuk Analisis Sentimen Terhadap Wacana KaburAjaDulu di Media Sosial X,” *Technol. Sci.*, vol. 7, no. 1, pp. 169–178, 2025, doi: 10.47065/bits.v7i1.7261.
- [28] H. Henderi and Q. Siddique, “Comparative Analysis of Sentiment Classification Techniques on Flipkart Product Reviews: A Study Using Logistic Regression, SVC, Random Forest, and Gradient Boosting,” *J. Digit. Mark. Digit. Curr.*, vol. 1, no. 1, pp. 21–42, 2024, doi: 10.47738/jdmcdc.v1i1.4.
- [29] M. R. F. Rahmatullah, P. N. Andono, and M. A. Soeleman, “Improving Random Forest Performance for Sentiment Analysis on Unbalanced Data Using SMOTE and BoW Integration : PLN Mobile Application Case Study,” vol. 12, no. 1, pp. 1–10, 2025, doi: 10.15294/sji.v12i1.19295.
- [30] T. Ahmed Khan, R. Sadiq, Z. Shahid, M. M. Alam, and M. Mohd Su’ud, “Sentiment Analysis using Support Vector Machine and Random Forest,” *J. Informatics Web Eng.*, vol. 3, no. 1, pp. 67–75, 2024, doi: 10.33093/jiwe.2024.3.1.5.
- [31] A. Putriyeki *et al.*, “Analisis Multidimensional Sentimen Masyarakat Terhadap Program Makan Bergizi Gratis Pada Media Sosial X,” *Integr. Perspect. Soc. Sci. J.*, vol. 2, no. 1, pp. 1004–1024, 2025.
- [32] I. G. Bintang, A. Budaya, and I. K. P. Suniantara, “Comparison of Sentiment Analysis Algorithms with SMOTE Oversampling and TF-IDF Implementation on Google Reviews for Public Health Centers,” *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. July, pp. 1077–1086, 2024, doi: <https://doi.org/10.57152/malcom.v4i3.1459>.
- [33] S. A. H. Bahtiar, C. K. Dewa, and A. Luthfi, “Comparison of Naïve Bayes and Logistic Regression in Sentiment Analysis on Marketplace Reviews Using Rating-Based Labeling,” *J. Inf. Syst. Informatics*, vol. 5, no. 3, pp. 915–927, Aug. 2023, doi: 10.51519/journalisi.v5i3.539.

-
- [34] Z. S. Munmun, S. Akter, and C. R. Parvez, "Machine Learning-Based Classification of Coronary Heart Disease: A Comparative Analysis of Logistic Regression, Random Forest, and Support Vector Machine Models," *OALib*, vol. 12, no. 03, pp. 1–12, 2025, doi: 10.4236/oalib.1113054.
- [35] 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.
- [36] F. Suandi *et al.*, "Enhancing Sentiment Analysis Performance Using SMOTE and Majority Voting in Machine Learning Algorithms," in *Proceedings of the 7th International Conference on Applied Engineering (ICAE 2024)*, Atlantis Press International BV, 2024, pp. 126–138. doi: 10.2991/978-94-6463-620-8.
- [37] S. Peerbasha, A. Saleem Raja, Y. Mohammed Iqbal, and M. Surputheen, "Diabetes Prediction using Decision Tree, Random Forest, Support Vector Machine, K-Nearest Neighbors, Logistic Regression Classifiers," 2023. doi: <https://doi.org/10.46947/joaasr542023680>.
- [38] A. A. Firdaus *et al.*, "Application of Sentiment Analysis as an Innovative Approach to Policy Making: A Review," *J. Robot. Control*, vol. 5, no. 6, pp. 1784–1798, 2024, doi: 10.18196/jrc.v5i6.22573.
- [39] T. V. S. Krishna, T. S. R. Krishna, S. Kalime, C. V. M. Krishna, S. Neelima, and R. R. Pbv, "A novel ensemble approach for Twitter sentiment classification with ML and LSTM algorithms for real-time tweets analysis," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 34, no. 3, pp. 1904–1914, Jun. 2024, doi: 10.11591/ijeecs.v34.i3.pp1904-1914.
- [40] M. W. Arif, "Analisis Sentimen Kebijakan Makan Bergizi Gratis di Media Sosial Menggunakan Natural Language Processing Berbasis Python TextBlob di Indonesia Teknik Informatika , Universitas Ngudi Waluyo , Indonesia Sentiment Analysis of Free Nutritious Meal Policy on S," vol. 5, no. 9, pp. 2463–2471, 2025, doi: <https://doi.org/10.52436/1.jpti.931>.
- [41] A. Fauzi *et al.*, "Optimalisasi Random Forest untuk Sentimen Bahasa Indonesia dengan GridSearch dan SMOTE," *J. Ilmu Komput. dan Sist. Inf.*, vol. 4, no. 2, pp. 202–217, 2025, doi: <https://doi.org/10.70340/jirsi.v4i2.207>.
- [42] G. R. Ati and P. T. Prasetyaningrum, "Analysis of Community Sentiment Towards Free Nutrition Meal Programs on Twitter Using Naïve Bayes , Support Vector Machine , K-Nearest Neighbors , and Ensemble Methods," vol. 7, no. 2, pp. 1443–1460, 2025, doi: 10.51519/journalisi.v7i2.1098.
- [43] Z. Purwanti, P. Studi Sistem Informasi, S. Tinggi Ilmu Komputer Cipta Karya Informatika, K. Jakarta Timur, and D. Khusus Ibukota Jakarta, "Pemodelan Text Mining untuk Analisis Sentimen Terhadap Program Makan Siang Gratis di Media Sosial X Menggunakan Algoritma Support Vector Machine (SVM)," 2024. doi: <https://doi.org/10.35870/jimik.v5i3.1001>.
- [44] M. Hanafi and M. Furqan, "Perbandingan Analisis Sentimen Presiden 2024 Menggunakan Algoritma Support Vector Machine dan K-Nearest Neighbor," *CESS (Journal Comput. Eng. Syst. Sci.)*, vol. 10, no. 1, p. 275, 2025, doi: 10.24114/cess.v10i1.67747.