

ANALISIS SENTIMEN PROGRAM MAKAN BERGIZI GRATIS DI MEDIA SOSIAL X MENGGUNAKAN TF-IDF DAN PERBANDINGAN ALGORITMA MACHINE LEARNING

Weliansyah¹, Muhamad Kurniawan²

¹ Sistem Informasi Manajemen, Program Pascasarjana, Institut Sains dan Bisnis Atma Luhur, Pangkalpinang, Indonesia

² Sains Data, Fakultas Sains dan Informatika, Universitas Pertiba, Pangkalpinang, Indonesia
Email: weliansyahkaka@gmail.com¹, muhamadkurniawan257@gmail.com²

Submit: 12 Juni 2026, Diterima: 18 Juli 2026, Diterbitkan: 30 Juli 2026

Abstrak: Program Makan Bergizi Gratis (MBG) merupakan salah satu kebijakan pemerintah yang memunculkan beragam tanggapan masyarakat di media sosial X. Besarnya volume opini yang dihasilkan menjadikan analisis sentimen sebagai pendekatan yang efektif untuk mengidentifikasi persepsi publik secara otomatis. Penelitian ini bertujuan menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis serta membandingkan kinerja algoritma *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* (SVM) dalam melakukan klasifikasi sentimen. Penelitian menggunakan pendekatan kuantitatif berbasis *Natural Language Processing* (NLP) dengan tahapan pengumpulan data, *preprocessing*, pelabelan sentimen menggunakan pendekatan lexicon, ekstraksi fitur menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF), pembagian data menggunakan *Stratified Train-Test Split*, validasi menggunakan *Stratified K-Fold Cross Validation*, optimasi parameter melalui *GridSearchCV*, dan evaluasi model menggunakan *Accuracy*, *Precision*, *Recall*, *F1-Score*, *Confusion Matrix*, serta *ROC Curve*. Hasil penelitian menunjukkan bahwa proses *crawling* menghasilkan 2.798 tweet, kemudian setelah *preprocessing* diperoleh 2.703 tweet yang digunakan dalam proses klasifikasi. Distribusi sentimen memperlihatkan dominasi sentimen positif terhadap Program Makan Bergizi Gratis. Seluruh algoritma mampu memberikan performa klasifikasi yang baik setelah dilakukan optimasi parameter. Berdasarkan hasil evaluasi, *Logistic Regression* memberikan performa terbaik, diikuti oleh *Support Vector Machine*, sedangkan *Naive Bayes* memberikan hasil yang kompetitif sebagai model pembandingan. Penelitian ini menunjukkan bahwa kombinasi *preprocessing*, pelabelan berbasis *lexicon*, ekstraksi fitur TF-IDF, serta optimasi model mampu menghasilkan klasifikasi sentimen yang efektif dan dapat dimanfaatkan sebagai pendukung evaluasi kebijakan publik berbasis data.

Kata kunci: *TF-IDF*, *Logistic Regression*, *Support Vector Machine*, *Naive Bayes*, *Program Makan Bergizi Gratis*

Abstract: The Free Nutritional Meal Program (MBG) is a government policy that has generated various public responses on social media X. The large volume of opinions generated makes sentiment analysis an effective approach to automatically identify public perception. This study aims to analyze public sentiment towards the Free Nutritional Meal Program and compare the performance of the *Naive Bayes*, *Logistic Regression*, and *Support Vector Machine* (SVM) algorithms in classifying sentiment. The study uses a quantitative approach based on *Natural Language Processing* (NLP) with stages of data collection, *preprocessing*, sentiment labeling using a lexicon approach, feature extraction using *Term Frequency-Inverse Document Frequency* (TF-IDF), data division using *Stratified Train-Test Split*, validation using *Stratified K-Fold Cross Validation*, parameter optimization through *GridSearchCV*, and model evaluation using *Accuracy*, *Precision*, *Recall*, *F1-Score*, *Confusion Matrix*, and *ROC Curve*. The results show that the *crawling* process produced 2,798 tweets, then after *preprocessing* obtained 2,703 tweets used in the classification process. The sentiment distribution shows a predominance of positive sentiment towards the Free Nutritional Meal Program. All algorithms performed well in classification after parameter optimization. Based on the evaluation results, *Logistic Regression* performed best, followed by *Support Vector Machine*, while *Naive Bayes* yielded competitive results as a comparison model. This study demonstrates that the combination of *preprocessing*, lexicon-based labeling, TF-IDF feature extraction, and model optimization produces effective sentiment classification and can be utilized to support data-driven public policy evaluation.

Keywords: *TF-IDF*, *Logistic Regression*, *Support Vector Machine*, *Naive Bayes*, *Free Nutritious Meal Program*

1. PENDAHULUAN

Perkembangan pesat teknologi informasi dan transformasi digital telah membawa perubahan signifikan dalam pola komunikasi masyarakat Indonesia [1]. Kehadiran internet dan platform media sosial tidak hanya memfasilitasi komunikasi personal, tetapi juga telah menjadi wadah utama bagi masyarakat untuk menyuarakan opini, berbagi pandangan, dan mengekspresikan kritik secara real-time dan masif [2]. Di antara berbagai platform yang ada, media sosial X (sebelumnya dikenal sebagai Twitter) menempati posisi strategis sebagai ruang publik digital di mana perdebatan mengenai isu-isu politik, ekonomi, dan kebijakan sosial sering kali berpusat. Tingginya interaksi di X menjadikannya "tambang emas" data tekstual yang kaya akan opini agregat masyarakat, yang sangat berharga bagi pemangku kebijakan untuk mengevaluasi efektivitas komunikasi publik [3].

Salah satu kebijakan pemerintah pusat yang saat ini memicu diskursus berskala nasional di media sosial X adalah Program Makan Bergizi Gratis (MBG). Program ini merupakan inisiatif turunan dari janji kampanye Presiden Republik Indonesia, Prabowo Subianto, yang bertujuan untuk meningkatkan kualitas gizi anak balita, siswa sekolah, dan ibu hamil secara gratis guna mencegah stunting dan meningkatkan kualitas sumber daya manusia [4], [5]. Namun, skala program yang masif dengan alokasi Anggaran Pendapatan dan Belanja Negara (APBN) yang sangat besar telah menimbulkan polarisasi pandangan di tengah masyarakat. Sebagian masyarakat memberikan sentimen positif karena dinilai berpihak pada rakyat kecil, sementara yang lain menunjukkan sentimen negatif berupa keraguan atas transparansi, skema pendanaan, dan potensi inefisiensi distribusi. Mengingat besarnya dampak kebijakan ini, diperlukan sebuah metode komputasi yang objektif untuk memetakan dan mengukur tendensi opini publik secara otomatis, yaitu melalui teknik *Sentiment Analysis* (Analisis Sentimen).

Sentiment Analysis atau *opinion mining* adalah cabang dari *Natural Language Processing* (NLP) dan *Text mining* yang bertujuan untuk mengekstrak, mengidentifikasi, dan mengklasifikasikan informasi tekstual ke dalam polaritas tertentu, seperti positif, negatif, atau netral [6], [7]. Dalam beberapa tahun terakhir, penerapan *Machine Learning* untuk klasifikasi sentimen kebijakan publik di Indonesia telah menunjukkan hasil yang menjanjikan. Misalnya, Berhasil memetakan sentimen pembangunan Ibu Kota Nusantara (IKN) [8], sementara menganalisis opini masyarakat terhadap program Kartu Indonesia Pintar Kuliah (KIP-K) menggunakan algoritma klasifikasi [9]. Dalam tahapan pemrosesan bahasa alami, ekstraksi fitur merupakan langkah krusial untuk mengubah data teks yang tidak terstruktur menjadi format matematis yang dapat dipahami oleh mesin. Metode *Term Frequency-Inverse Document Frequency* (TF-IDF) secara konsisten diakui sebagai teknik pembobotan kata yang paling tangguh untuk memfilter *noise* dan menyoroti kata-kata yang secara kontekstual penting dalam sebuah kumpulan ujaran [10]. Setelah fitur diekstraksi, pemilihan algoritma klasifikasi menentukan tingkat akurasi model. Berdasarkan literatur, *Support Vector Machine* (SVM) sering kali memberikan performa unggul pada data teks berdimensi tinggi karena kemampuannya menemukan *hyperplane* optimal yang memisahkan antar kelas sentimen [11], [12]. Di sisi lain, *Naive Bayes* tetap menjadi algoritma dasar yang sangat populer karena efisiensi komputasi probabilitasnya yang didasarkan pada *Teorema Bayes*, serta sangat tangguh pada *dataset* independen [13]. Sementara itu, *Logistic Regression* semakin banyak diadopsi karena memberikan probabilitas prediktif yang sangat baik untuk klasifikasi biner dan multikelas dalam analisis sentimen *teks linier* [8].

Meskipun eksplorasi *Machine Learning* pada analisis sentimen kebijakan publik telah marak dilakukan, riset yang secara spesifik mengangkat topik Program Makan Bergizi Gratis (MBG) masih sangat langka. Sebagian besar penelitian berfokus pada kebijakan masa pandemi seperti *COVID-19*, PPKM, atau pemindahan IKN [8], [14]. Studi awal yang meneliti program makan siang gratis, hanya mengandalkan penggunaan satu algoritma tunggal, yaitu *Naive Bayes*. Keterbatasan pada studi-studi terdahulu adalah minimnya perbandingan komprehensif antara beberapa algoritma *Supervised Learning* terbaik. Tidak ada jaminan bahwa *Naive Bayes* adalah algoritma yang paling optimal untuk *dataset* spesifik terkait MBG yang memiliki karakteristik linguistik slang atau bahasa tidak baku yang khas di media sosial X. Hingga saat ini, belum ada literatur yang secara sistematis membandingkan performa *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* (SVM) yang dikombinasikan dengan ekstraksi fitur TF-IDF untuk mendeteksi sentimen publik terhadap kebijakan Program Makan Bergizi Gratis (MBG). Kesenjangan inilah yang menjadi landasan utama urgensi dilakukannya penelitian ini [5], [15], [16].

Pemanfaatan *dataset primer* terbaru hasil *crawling* dari media sosial X yang difokuskan secara eksklusif pada respons dan opini masyarakat Indonesia terkait implementasi dan wacana Program Makan Bergizi Gratis (MBG) pasca-Pemilu 2024. kerangka penelitian yang menandingkan tiga algoritma raksasa dalam *Text Classification* (*Naive Bayes*, *Logistic Regression*, dan SVM) dengan menggunakan skema pembobotan TF-IDF untuk mencari arsitektur model paling optimal pada teks berbahasa Indonesia. Kontribusi dari penelitian ini bersifat teoritis dan praktis. Secara teoritis, penelitian ini akan Melengkapi literatur yang ada *Natural Language Processing* (NLP) berbahasa Indonesia, khususnya dalam penanganan teks opini media sosial berdimensi tinggi. Secara praktis, hasil klasifikasi dari penelitian ini dirancang untuk dapat

diimplementasikan sebagai dashboard intelijen atau sistem pendukung keputusan (SPK) bagi pemerintah pusat dan kementerian terkait. Hal ini akan memfasilitasi evaluasi sentimen publik secara *real-time*, sehingga strategi komunikasi publik dan implementasi kebijakan MBG ke depannya dapat disesuaikan agar lebih tepat sasaran dan transparan.

Penelitian ini bertujuan untuk menganalisis dan memetakan distribusi persentase sentimen publik yang terdiri atas sentimen positif, negatif, dan netral di media sosial X terhadap implementasi kebijakan Program Makan Bergizi Gratis (MBG). Selain itu, penelitian ini mengimplementasikan teknik *preprocessing* teks berbahasa Indonesia serta melakukan ekstraksi fitur *linguistik* menggunakan metode pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) sebagai representasi numerik teks yang efektif untuk proses klasifikasi sentimen. Untuk memperoleh hasil klasifikasi yang optimal, penelitian ini membandingkan tiga algoritma *Machine Learning*, yaitu *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* (SVM). *Naive Bayes* dipilih karena memiliki efisiensi yang tinggi dalam menangani data teks berskala besar melalui asumsi independensi antarfitur yang sederhana namun efektif. *Logistic Regression* digunakan karena mampu memodelkan probabilitas kelas secara akurat sehingga banyak diterapkan dalam permasalahan klasifikasi teks. Sementara itu, *Support Vector Machine* (SVM) dipilih karena memiliki kemampuan yang sangat baik dalam menangani data berdimensi tinggi serta mampu menemukan *hyperplane* optimal untuk memisahkan kelas sentimen dengan tingkat akurasi yang relatif tinggi. Selanjutnya, penelitian ini melakukan uji komparasi terhadap ketiga algoritma tersebut dengan mengevaluasi performa masing-masing model menggunakan *Confusion Matrix* serta metrik evaluasi *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Hasil evaluasi tersebut digunakan untuk mengidentifikasi algoritma yang memberikan performa terbaik dalam klasifikasi sentimen terhadap kebijakan Program Makan Bergizi Gratis (MBG). Melalui integrasi tahapan *preprocessing*, pembobotan TF-IDF, serta penerapan dan evaluasi ketiga algoritma *Machine Learning* tersebut, penelitian ini diharapkan mampu menghasilkan model klasifikasi sentimen yang akurat sekaligus memberikan gambaran yang komprehensif mengenai persepsi masyarakat terhadap implementasi Program Makan Bergizi Gratis (MBG). Hasil penelitian ini diharapkan dapat menjadi dasar pertimbangan bagi pemerintah dalam melakukan evaluasi, penyempurnaan, dan pengambilan keputusan terhadap kebijakan MBG di masa mendatang.

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif berbasis *Natural Language Processing* (NLP) dan *Text mining* untuk menganalisis sentimen masyarakat. Alur penelitian dirancang secara sistematis melalui beberapa tahapan utama, yaitu: pengumpulan data, pra-pemrosesan teks (*text preprocessing*), pelabelan data, ekstraksi fitur menggunakan TF-IDF, pemodelan klasifikasi (*Machine Learning*), dan evaluasi model. Seluruh tahapan eksperimen diimplementasikan menggunakan bahasa pemrograman Python.



Gambar 1. Metode Penelitian

2.1 PENGUMPULAN DATA

Data yang digunakan dalam penelitian ini adalah data primer berupa teks opini masyarakat (cuitan/tweet) yang bersumber dari platform media sosial X (sebelumnya Twitter). Pengambilan data (*crawling/scraping*)

dilakukan menggunakan *Application Programming Interface* (API) atau *library scraper* bawaan Python dengan kata kunci spesifik seperti "makan bergizi gratis", "program makan gratis", dan "makan siang gratis". Data dibatasi pada periode waktu tertentu yang relevan dengan momentum kebijakan program MBG. Data yang berhasil diekstrak kemudian disimpan dalam format *Comma Separated Values* (CSV) untuk diproses lebih lanjut. Atribut data yang diambil difokuskan pada teks tweet, tanggal, dan username.

2.2 PRA-PEMROSESAN DATA

Teks dari media sosial umumnya bersifat tidak terstruktur dan mengandung banyak gangguan (*noise*). Oleh karena itu, tahap *preprocessing* sangat krusial untuk meningkatkan kualitas data sebelum masuk ke tahap pemodelan. Tahapan ini meliputi: *Data Cleaning*: Membersihkan teks dari elemen yang tidak memiliki bobot sentimen, seperti *URL/link*, *username* (@), *hashtag* (#), *emoticon*, angka, dan tanda baca (*punctuation*). *Case Folding*: Mengonversi seluruh huruf kapital dalam teks menjadi huruf kecil (*lowercase*) agar mesin membaca karakter secara seragam. *Tokenization*: Memecah kalimat utuh menjadi potongan-potongan kata penyusunnya (*token*). *Normalization*: Memperbaiki dan menyeragamkan kata-kata tidak baku, *slang word* (bahasa gaul), serta singkatan menjadi kata baku sesuai Ejaan Yang Disempurnakan (EYD) menggunakan kamus normalisasi modifikasi. *Stopword Removal*: Menghilangkan kata hubung atau kata umum yang sering muncul namun tidak memiliki makna sentimen signifikan (misal: "yang", "dan", "di", "ke"). *Stemming*: Mengembalikan kata berimbuhan menjadi kata dasar. Proses ini menggunakan library Sastrawi khusus untuk teks berbahasa Indonesia.

2.3 PELABELAN DATA

Teks yang telah bersih (hasil *preprocessing*) akan diklasifikasikan ke dalam tiga kelas polaritas sentimen: Positif, Negatif, dan Netral. Proses pelabelan dilakukan dengan menggabungkan pendekatan *Lexicon-based* (berbasis kamus sentimen bahasa Indonesia, seperti *InSet Lexicon* atau kamus modifikasi ahli) untuk memberikan label awal secara otomatis. Setelah itu, dilakukan validasi secara manual (*human annotator*) pada sebagian sampel data untuk memastikan konteks kalimat terkait program Makan Bergizi Gratis (MBG) sudah sesuai, khususnya untuk menangani kalimat sarkasme atau ironi.

2.4 EKSTRAKSI FITUR

Agar data teks dapat diproses oleh algoritma *Machine Learning*, data teks perlu ditransformasi menjadi format vektor numerik. Penelitian ini menggunakan metode TF-IDF. *Term Frequency* (TF) mengukur seberapa sering sebuah kata muncul dalam satu dokumen (cuitan). *Inverse Document Frequency* (IDF) mengukur seberapa penting kata tersebut di seluruh korpus dokumen. Metode ini akan memberikan bobot tinggi pada kata yang sering muncul di dokumen tertentu, tetapi jarang muncul pada dokumen lain. Model ini sangat baik untuk menonjolkan kata kunci spesifik yang menjadi fitur sentimen pembeda.

2.5 PEMODELAN KLASIFIKASI

Dataset yang telah diekstraksi menggunakan TF-IDF dibagi menjadi dua bagian, yaitu Data Latih (*Training Data*) dan Data Uji (*Testing Data*), dengan skema pembagian (rasio) seperti 80:20 atau 90:10. Penelitian ini melakukan studi perbandingan dengan mengimplementasikan tiga algoritma klasifikasi *Supervised Machine Learning*: *Naïve Bayes* (NB): Menggunakan pendekatan *Multinomial Naive Bayes* yang berasumsi bahwa setiap fitur (*term*) bersifat independen secara kondisional. NB dipilih karena kesederhanaan komputasinya dan performanya yang sangat baik pada klasifikasi teks dimensi tinggi. *Logistic Regression* (LR): Model statistik probabilistik yang digunakan untuk memprediksi probabilitas kelas kategorikal. LR seringkali kuat menangani model linier pada representasi fitur TF-IDF karena dapat menimbang bobot penting setiap token. *Support Vector Machine* (SVM): Algoritma yang mencari *hyperplane* optimal yang memaksimalkan margin jarak antar kelas sentimen (Positif, Negatif, Netral). SVM dikenal tangguh dan tidak mudah mengalami *overfitting* pada klasifikasi *text mining*. Pada tahap ini, *Hyperparameter Tuning* (seperti optimasi parameter C dan Kernel pada SVM) dapat dilakukan menggunakan *GridSearchCV* untuk mendapatkan performa klasifikasi terbaik dari masing-masing model.

2.6 EVALUASI MODEL

Tahap terakhir adalah mengukur kinerja dari ketiga model (*Naive Bayes*, *Logistic Regression*, dan SVM) menggunakan *Confusion Matrix*. Evaluasi ini menghitung metrik performa klasifikasi yang meliputi: *Accuracy* (Akurasi): Persentase kebenaran prediksi keseluruhan model. *Precision* (Presisi): Tingkat ketepatan prediksi sentimen positif/negatif dibandingkan dengan hasil aktual. *Recall* (Sensitivitas): Kemampuan model

menemukan kembali seluruh informasi dari kelas yang relevan. *F1-Score* : Rata-rata harmonik antara *Precision* dan *Recall* yang digunakan sebagai acuan utama, terutama jika kelas data tidak seimbang (*imbalanced dataset*). Berdasarkan hasil matriks evaluasi, kinerja ketiga algoritma akan dibandingkan untuk menyimpulkan algoritma mana yang paling optimal serta melihat bagaimana kecenderungan (tren) sentimen publik terhadap implementasi program Makan Bergizi Gratis (MBG).

3. HASIL DAN PEMBAHASAN

3.1 PENGUMPULAN DATA

Tahap awal penelitian dilakukan dengan proses pengumpulan data menggunakan teknik *crawling* pada media sosial Twitter/X. Pengumpulan data difokuskan pada opini masyarakat mengenai Program Makan Bergizi Gratis (MBG) dengan memanfaatkan kata kunci yang berkaitan dengan program tersebut. Proses *crawling* menghasilkan data mentah dalam format *Comma Separated Values* (CSV) yang selanjutnya digunakan sebagai sumber data utama pada penelitian ini. Seluruh proses pengambilan data dilakukan secara otomatis menggunakan bahasa pemrograman Python sehingga setiap tweet yang sesuai dengan kata kunci penelitian dapat tersimpan secara sistematis untuk dianalisis pada tahapan berikutnya.

Dataset yang diperoleh terdiri atas berbagai atribut yang menggambarkan karakteristik setiap tweet, seperti identitas pengguna, waktu publikasi, isi tweet, jumlah balasan (reply), jumlah retweet, jumlah tanda suka (favorite), bahasa, lokasi, serta beberapa atribut pendukung lainnya. Keberadaan atribut-atribut tersebut memberikan informasi yang cukup lengkap mengenai aktivitas pengguna Twitter dalam memberikan tanggapan terhadap Program Makan Bergizi Gratis. Berdasarkan hasil proses pengumpulan data, diperoleh 2.798 data tweet dengan 15 atribut yang selanjutnya digunakan sebagai objek penelitian. Setelah *dataset* berhasil diperoleh, dilakukan pemeriksaan awal terhadap struktur data untuk mengetahui tipe data, jumlah baris, jumlah kolom, serta kelengkapan setiap atribut. Pemeriksaan awal ini bertujuan memastikan bahwa data yang diperoleh dapat diproses pada tahapan berikutnya tanpa mengalami kesalahan akibat ketidaksesuaian struktur data.

INFORMASI DATASET

Ukuran Dataset : (2798, 15)

Nama Kolom:

1. conversation_id_str
2. created_at
3. favorite_count
4. full_text
5. id_str
6. image_url
7. in_reply_to_screen_name
8. lang
9. location
10. quote_count
11. reply_count
12. retweet_count
13. tweet_url
14. user_id_str

conversation_id	created_at	favorite_count	full_text	id_str	image_url	in_reply_to_screen_name	lang	location	quote_count	reply_count	retweet_count	tweet_url	user_id_str	username
15555555555555555555	2024-05-10T08:00:00.000Z	0	Siswa di Bogor salah satu	186510035647	nan	nan	in	jakarta	0.0	0	0	https://x.com/kezonews/status/186510035647	186510035647	kezonews
15555555555555555555	2024-05-10T08:00:00.000Z	0	MDA 186510035647	186510035647	nan	lenteradata	in	nan	0.0	0	0	https://x.com/kezonews/status/186510035647	186510035647	kezonews
15555555555555555555	2024-05-10T08:00:00.000Z	0	MDA 186510035647	186510035647	nan	nan	in	Bangka Belitung	0.0	0	0	https://x.com/kezonews/status/186510035647	186510035647	kezonews
15555555555555555555	2024-05-10T08:00:00.000Z	0	MDA 186510035647	186510035647	nan	lenteradata	in	nan	0.0	0	0	https://x.com/kezonews/status/186510035647	186510035647	kezonews
15555555555555555555	2024-05-10T08:00:00.000Z	0	MDA 186510035647	186510035647	nan	prabowo	in	nan	0.0	0	0	https://x.com/kezonews/status/186510035647	186510035647	kezonews

Gambar 2. Informasi *Dataset*

Berdasarkan Gambar 4.1 dapat diketahui bahwa *dataset* memiliki struktur yang cukup baik karena seluruh atribut utama berhasil terbaca oleh sistem. Informasi tersebut menunjukkan bahwa data hasil *crawling* telah tersimpan dalam format yang sesuai sehingga dapat langsung diproses menggunakan pustaka Pandas pada Python. Tahapan ini sangat penting karena struktur data yang baik akan mempermudah proses analisis pada tahap selanjutnya. Selanjutnya dilakukan identifikasi terhadap missing value untuk mengetahui apakah terdapat atribut yang memiliki nilai kosong. Pemeriksaan ini dilakukan agar kualitas *dataset* dapat diketahui sebelum memasuki tahap pra-pemrosesan. Berdasarkan hasil pemeriksaan ditemukan bahwa beberapa atribut seperti lokasi pengguna maupun *URL* gambar memiliki nilai kosong, sedangkan atribut utama berupa isi tweet (*full_text*) tetap tersedia sehingga seluruh tweet masih dapat digunakan dalam proses analisis sentimen.

	0
image_url	2060
location	1808
in_reply_to_screen_name	1399
id_str	1
favorite_count	1
full_text	1
created_at	1
retweet_count	1
lang	1
quote_count	1
reply_count	1
user_id_str	1
tweet_url	1
username;;	1
conversation_id_str	0

dtype: int64

Gambar 3. Missing Value Dataset

Hasil pemeriksaan missing value menunjukkan bahwa nilai kosong hanya ditemukan pada atribut yang tidak digunakan dalam proses klasifikasi sentimen. Oleh karena itu, penelitian tetap menggunakan seluruh data tweet karena informasi utama berupa isi teks tidak mengalami kehilangan data. Kondisi tersebut menunjukkan bahwa kualitas *dataset* cukup baik untuk digunakan pada penelitian analisis sentimen. Tahapan berikutnya adalah melakukan analisis statistik deskriptif terhadap *dataset*. Analisis ini bertujuan memperoleh gambaran umum mengenai data yang berhasil dikumpulkan sebelum dilakukan proses pembersihan data. Statistik deskriptif memberikan informasi mengenai jumlah data, distribusi atribut, serta karakteristik umum *dataset* yang digunakan dalam penelitian.

```
=====
STATISTIK DATASET
=====
Jumlah Tweet : 2705
Jumlah Kolom : 15
```

Gambar 3. Statistik Dataset

Berdasarkan hasil statistik deskriptif terlihat bahwa *dataset* memiliki jumlah data yang cukup besar untuk digunakan sebagai data pelatihan model klasifikasi. Banyaknya data yang berhasil dikumpulkan diharapkan mampu memberikan representasi opini masyarakat mengenai Program Makan Bergizi Gratis secara lebih komprehensif. Semakin banyak data yang digunakan, maka pola sentimen yang dipelajari oleh algoritma klasifikasi akan semakin baik sehingga hasil klasifikasi menjadi lebih representatif.

Secara keseluruhan, tahap pengumpulan data berhasil menghasilkan *dataset* yang memenuhi kebutuhan penelitian. *Dataset* tersebut selanjutnya digunakan sebagai masukan pada tahap pra-pemrosesan teks agar kualitas data meningkat sebelum dilakukan proses pelabelan sentimen dan pembangunan model klasifikasi.

3.2 PRA-PEMROSESAN TEKS

Tahapan pra-pemrosesan teks (*text preprocessing*) merupakan salah satu tahap penting dalam penelitian analisis sentimen karena bertujuan meningkatkan kualitas data sebelum dilakukan ekstraksi fitur dan proses klasifikasi. Tweet yang diperoleh dari media sosial umumnya masih mengandung berbagai karakter yang tidak memiliki kontribusi terhadap proses klasifikasi, seperti URL, mention, hashtag, emoji, tanda baca, angka, maupun variasi penulisan kata. Oleh karena itu, dilakukan serangkaian proses pembersihan data agar setiap tweet memiliki bentuk yang lebih terstruktur. Berdasarkan implementasi yang terdapat pada penelitian ini, *preprocessing* dilakukan melalui enam tahapan utama, yaitu case folding, cleaning, tokenization, normalization, stopwords removal, dan stemming menggunakan pustaka Sastrawi. Seluruh tahapan tersebut diterapkan secara berurutan melalui fungsi *preprocessing* yang dibangun dalam bahasa pemrograman Python.

- Tahap pertama adalah case folding, yaitu mengubah seluruh karakter menjadi huruf kecil (lowercase). Proses ini bertujuan menyeragamkan penulisan kata sehingga kata yang memiliki makna sama tetapi berbeda penggunaan huruf kapital tidak dianggap sebagai kata yang berbeda. Misalnya kata "Program" dan "program" akan direpresentasikan sebagai satu bentuk yang sama.
- Tahap kedua adalah cleaning, yaitu menghapus karakter yang tidak diperlukan dalam proses analisis sentimen. Pada tahap ini dilakukan penghapusan URL, mention pengguna, hashtag, angka, tanda baca, karakter khusus, maupun spasi berlebih sehingga hanya menyisakan teks utama yang mengandung informasi sentimen.
- Tahap ketiga adalah tokenization, yaitu memecah kalimat menjadi kumpulan kata (token). Proses ini memudahkan tahapan selanjutnya karena setiap kata dapat diproses secara individual.

- Tahap berikutnya adalah normalization, yaitu mengubah kata-kata tidak baku menjadi bentuk bakunya sesuai dengan kamus normalisasi yang digunakan dalam penelitian. Normalisasi membantu mengurangi variasi penulisan kata yang sering ditemukan pada media sosial.
- Selanjutnya dilakukan stopword removal, yaitu menghapus kata-kata yang sering muncul tetapi tidak memiliki makna penting dalam proses klasifikasi sentimen. Penghapusan stopwords membuat model lebih fokus pada kata-kata yang memiliki informasi sentimen.
- Tahap terakhir adalah stemming menggunakan algoritma dari pustaka Sastrawi. Proses ini mengubah setiap kata menjadi bentuk kata dasar sehingga variasi imbuhan tidak lagi dianggap sebagai kata yang berbeda. Implementasi stemming dilakukan menggunakan fungsi StemmerFactory() yang tersedia pada pustaka Sastrawi.

Setelah seluruh proses *preprocessing* selesai dilakukan, dihasilkan kolom baru bernama `clean_text` yang berisi hasil pembersihan setiap tweet. Kolom inilah yang selanjutnya digunakan sebagai masukan pada proses pelabelan sentimen dan ekstraksi fitur TF-IDF. Berdasarkan hasil *preprocessing*, jumlah data berubah menjadi 2.703 tweet setelah data yang menghasilkan teks kosong dihapus dari *dataset*.

	full_text	clean_text
0	Momen Prabowo Tanda Tangan Sepatu Siswa di Bogor saat Tinjau Makan Bergizi Gratis https://t.co/2Z1A3585tZ	momen prabowo tanda tangan sepatu siswa bogor tinjau makan gizi gratis
1	@lenteradata Semoga program makan bergizi gratis dri MDA tetap berlanjut walaupun programnya sma dengan pemerintah	moga program makan gizi gratis dri mda lanjut program sma perintah
2	Pemprov Babel dan DPRD Bahas Anggaran Program Makan Bergizi Gratis di APBD 2025 https://t.co/T4VhiNlUf	pemprov babel dprd bahas anggar program makan gizi gratis apbd
3	@lenteradata Menurutku berhasil skliki programnya MDA yg makan bergizi gratis. Buktinya terjadi penurunan angka stunting di daerah tersebut	turut hasil skliki program mda makan gizi gratis bukti turun angka stunting daerah
4	@prabowo Tolol . Bergizi apaan . Itu program makan gratis bersyukur .	tolol gizi program makan gratis syukur
5	Pangdam IV/Diponegoro Mayjen TNI Dedy Suryadi meninjau program Makan Bergizi Gratis (MBG) di Kota Semarang pada Senin 10 Februari 2025. https://t.co/lkgV2m7GMR #suaramerdeka #suaramerdekaom #pangdam #pangdam4diponegoro #kotasemarang #semarang #programmakanbergizigratis https://t.co/HuMXZDuJ8	pangdam iv diponegoro mayjen tni deddy suryadi tinjau program makan gizi gratis mbg kota semarang senin februari suaramerdeka suaramerdekaom pangdam pangdam diponegoro kotasemarang semarang programmakanbergizigratis
6	MBG : makan bergizi gratis MBG : makan burger gratis	mbg makan gizi gratis mbg makan burger gratis
7	Peneliti LPEM UI: Tanpa Pengawasan Ketat Program Makan Bergizi Gratis Rentan terhadap Inefisiensi Anggaran https://t.co/bH0Rj0U4	teliti lpeu ui awas ketat program makan gizi gratis rentan inefisiensi anggar
8	Pemkot Bandar Lampung Siapkan Rp54 Miliar untuk Program Makan Bergizi Gratis Senin (10/2/2025) Baca Selengkapnya di: https://t.co/c0rkNvziw1s	pemkot bandar lampung siap rp 54 miliar program makan gizi gratis senin baca lengkap
9	di makan bergizi gratis kayaknya ga ada tuh yg videonya sampe begini cewe ampe geleng geleng kepala yak bualah bedanya blog	makan gizi gratis kayak tuh video sampe cewe ampe geleng geleng kepala yak beda blog

Gambar 4. Hasil *Dataset* Setelah *Preprocessing*

Berdasarkan Gambar 4 terlihat bahwa hasil *preprocessing* menghasilkan teks yang jauh lebih ringkas dibandingkan teks asli. URL, mention, tanda baca, serta karakter lain yang tidak memiliki kontribusi terhadap proses klasifikasi berhasil dihilangkan sehingga hanya menyisakan kata-kata penting yang mewakili isi tweet. Perubahan tersebut menunjukkan bahwa *preprocessing* berhasil mempertahankan informasi utama dari tweet sekaligus menghilangkan unsur-unsur yang tidak diperlukan dalam proses klasifikasi.

```

=====
Tweet Asli
Momen Prabowo Tanda Tangan Sepatu Siswa di Bogor saat Tinjau Makan Bergizi Gratis https://t.co/2Z1A3585tZ

Setelah Preprocessing
momen prabowo tanda tangan sepatu siswa bogor tinjau makan gizi gratis

=====
Tweet Asli
@lenteradata Semoga program makan bergizi gratis dri MDA tetap berlanjut walaupun programnya sma dengan pemerintah

Setelah Preprocessing
moga program makan gizi gratis dri mda lanjut program sma perintah

=====
Tweet Asli
Pemprov Babel dan DPRD Bahas Anggaran Program Makan Bergizi Gratis di APBD 2025 https://t.co/T4VhiNlUf

Setelah Preprocessing
pemprov babel dprd bahas anggar program makan gizi gratis apbd

=====
Tweet Asli
@lenteradata Menurutku berhasil skliki programnya MDA yg makan bergizi gratis. Buktinya terjadi penurunan angka stunting di daerah tersebut

Setelah Preprocessing
turut hasil skliki program mda makan gizi gratis bukti turun angka stunting daerah

```

Gambar 5. Contoh Perbandingan Tweet Sebelum dan Sesudah *Preprocessing*

Hasil *preprocessing* menunjukkan bahwa seluruh tahapan berhasil menghasilkan representasi teks yang lebih bersih, konsisten, dan siap diproses menggunakan metode TF-IDF. Selain itu, penghapusan data kosong menyebabkan jumlah *dataset* berkurang dari hasil *crawling* awal menjadi 2.703 tweet, yang menunjukkan bahwa hanya sebagian kecil data yang tidak dapat digunakan pada proses analisis lebih lanjut. Secara keseluruhan, tahap pra-pemrosesan berhasil meningkatkan kualitas *dataset* dengan menghilangkan noise, menyeragamkan bentuk kata, dan menghasilkan teks yang lebih representatif. Kondisi ini menjadi dasar yang penting untuk memperoleh hasil pelabelan sentimen dan pemodelan klasifikasi yang lebih optimal pada tahapan penelitian berikutnya.

3.3 PELABELAN SENTIMEN

Setelah seluruh data berhasil melalui tahapan pra-pemrosesan, langkah selanjutnya adalah melakukan pelabelan sentimen (sentiment labeling) terhadap setiap tweet. Tahapan ini bertujuan untuk menentukan kecenderungan opini masyarakat terhadap Program Makan Bergizi Gratis (MBG) sehingga setiap tweet dapat dikelompokkan ke dalam kelas positif, negatif, maupun netral. Pada penelitian ini proses pelabelan dilakukan

menggunakan pendekatan Lexicon Based Sentiment Analysis, yaitu metode yang menentukan polaritas sentimen berdasarkan kamus kata (lexicon) yang telah diberi bobot nilai sentimen. Berbeda dengan metode pelabelan manual yang membutuhkan waktu dan tenaga yang besar, pendekatan lexicon memungkinkan proses pelabelan dilakukan secara otomatis dengan menghitung skor sentimen berdasarkan kata-kata yang terdapat pada setiap tweet. Semakin banyak kata positif yang ditemukan maka skor sentimen akan bernilai positif, sedangkan semakin banyak kata negatif yang ditemukan maka skor sentimen akan bernilai negatif. Apabila skor akhir bernilai nol maka tweet dikategorikan sebagai sentimen netral. Implementasi metode tersebut dilakukan menggunakan fungsi `sentiment_score()` dan `sentiment_label()` sebagaimana ditunjukkan pada kode program penelitian.

3.3.1 Pembangunan Kamus Lexicon

Tahap pertama dalam proses pelabelan adalah membangun kamus sentimen (sentiment lexicon). Penelitian ini menyusun dua kelompok kosakata, yaitu positive lexicon dan negative lexicon, yang masing-masing berisi kata-kata umum maupun kata-kata yang berkaitan secara spesifik dengan Program Makan Bergizi Gratis. Kamus positif terdiri atas kata-kata seperti baik, bagus, berhasil, sehat, bergizi, gratis, higienis, dan berkualitas. Kata-kata tersebut diberikan bobot positif yang berbeda sesuai tingkat kekuatan sentimennya. Sebaliknya, kamus negatif memuat kata-kata seperti buruk, gagal, korupsi, keracunan, basi, busuk, berulat, kontaminasi, dan beracun yang masing-masing memiliki bobot negatif. Berdasarkan hasil implementasi diperoleh:

- Jumlah kata positif sebanyak 40 kata.
- Jumlah kata negatif sebanyak 46 kata.
- Total kosakata pada kamus sentimen sebanyak 86 kata.

Keberadaan kamus sentimen menjadi komponen utama dalam pendekatan Lexicon Based Sentiment Analysis karena seluruh proses penentuan label sentimen bergantung pada bobot kata yang terdapat di dalam kamus tersebut. Dengan menggabungkan kosakata umum dan kosakata yang berkaitan langsung dengan topik MBG, sistem diharapkan mampu mengenali sentimen masyarakat secara lebih sesuai dengan konteks penelitian.

3.3.2 Proses Perhitungan Skor Sentimen

Setelah kamus sentimen berhasil dibangun, setiap tweet diproses menggunakan fungsi `sentiment_score()`. Fungsi ini bekerja dengan menjumlahkan bobot setiap kata yang ditemukan pada kolom `clean_text`. Apabila suatu kata terdapat dalam kamus positif maka nilainya akan ditambahkan ke skor total, sedangkan apabila kata tersebut berada pada kamus negatif maka nilainya akan mengurangi skor sentimen. Selanjutnya skor yang diperoleh digunakan sebagai dasar untuk menentukan label sentimen menggunakan fungsi `sentiment_label()`. Aturan pelabelan yang diterapkan adalah sebagai berikut:

- Skor > 0 dikategorikan sebagai Positif.
- Skor < 0 dikategorikan sebagai Negatif.
- Skor = 0 dikategorikan sebagai Netral.

Pendekatan tersebut memungkinkan proses pelabelan dilakukan secara konsisten terhadap seluruh *dataset* tanpa memerlukan anotasi manual pada setiap tweet.

3.3.3 Hasil Pelabelan Dataset

Setelah seluruh tweet diproses, sistem menghasilkan tiga kolom utama, yaitu `clean_text`, `score`, dan `sentiment`. Kolom `score` menunjukkan hasil perhitungan skor sentimen berdasarkan kamus lexicon, sedangkan kolom `sentiment` menunjukkan kategori sentimen akhir yang diberikan kepada setiap tweet. Contoh hasil pelabelan ditampilkan pada dokumen penelitian dalam bentuk tabel hasil output program.

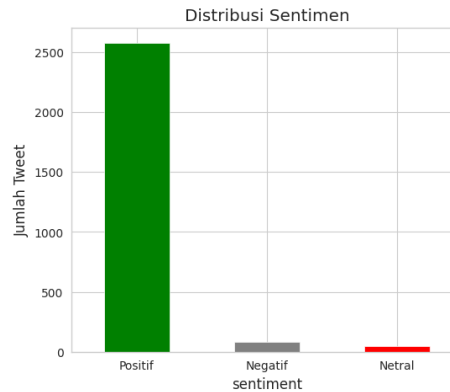
	<code>clean_text</code>	<code>score</code>	<code>sentiment</code>
0	momen prabowo tanda tangan sepatu siswa bogor tinjau makan gizi gratis	6	Positif
1	moga program makan gizi gratis dri mda lanjut program sma perintah	6	Positif
2	pemprov babel dprd bahas anggaran program makan gizi gratis apbd	6	Positif
3	turut hasil skilki program mda makan gizi gratis bukti turun angka stunting daerah	6	Positif
4	tolol gizi program makan gratis syukur	6	Positif

Gambar 6. Contoh Hasil Pelabelan Sentimen

Berdasarkan contoh hasil tersebut terlihat bahwa tweet yang mengandung lebih banyak kata positif memperoleh skor positif dan secara otomatis diklasifikasikan sebagai sentimen positif. Hal ini menunjukkan bahwa fungsi pelabelan telah berjalan sesuai dengan mekanisme yang dirancang pada penelitian.

3.3.4 Distribusi Sentimen

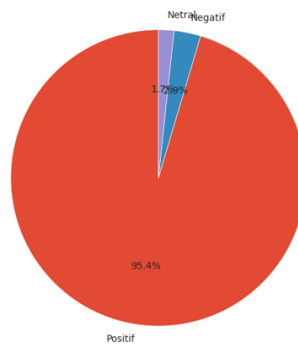
Setelah seluruh tweet berhasil diberi label, dilakukan analisis distribusi sentimen untuk mengetahui kecenderungan opini masyarakat terhadap Program Makan Bergizi Gratis. Hasil pengolahan data menunjukkan distribusi sentimen sebagai berikut.



Gambar 7. Distribusi Sentimen Twitter MBG

Berdasarkan hasil tersebut dapat diketahui bahwa mayoritas tweet yang berhasil dikumpulkan memiliki kecenderungan sentimen positif. Sebaliknya, jumlah tweet dengan sentimen negatif maupun netral relatif jauh lebih sedikit. Grafik distribusi memperlihatkan dominasi kelas positif dibandingkan dua kelas lainnya. Dominasi tersebut menunjukkan bahwa sebagian besar pengguna Twitter memberikan tanggapan yang bersifat mendukung atau memberikan penilaian positif terhadap Program Makan Bergizi Gratis berdasarkan hasil pelabelan menggunakan metode lexicon. Selain grafik batang, penelitian juga menampilkan visualisasi dalam bentuk diagram lingkaran (pie chart) untuk memperlihatkan proporsi masing-masing kelas sentimen.

Persentase Sentimen Twitter MBG



Gambar 8. Persentase Sentimen Twitter MBG

Visualisasi tersebut semakin memperjelas bahwa sentimen positif mendominasi keseluruhan *dataset*. Sementara itu, proporsi sentimen negatif dan netral hanya menempati sebagian kecil dari keseluruhan tweet yang berhasil dikumpulkan. Hasil ini menjadi dasar bagi tahapan berikutnya, yaitu pembangunan model klasifikasi menggunakan algoritma *Machine Learning*. Secara keseluruhan, tahap pelabelan sentimen berhasil menghasilkan *dataset* yang telah memiliki kelas target (Positif, Negatif, dan Netral) sehingga dapat digunakan sebagai data latih maupun data uji pada proses pemodelan klasifikasi.

3.4 EKSTRAKSI FITUR TF-IDF

Tahapan berikutnya dalam penelitian ini adalah ekstraksi fitur menggunakan metode Term Frequency–Inverse Document Frequency (TF-IDF). Tujuan utama dari tahap ini adalah mengubah data teks yang masih berupa kalimat menjadi representasi numerik sehingga dapat diproses oleh algoritma klasifikasi seperti *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine*. Kolom *clean_text* yang dihasilkan pada tahap *preprocessing* digunakan sebagai masukan utama dalam proses pembobotan TF-IDF, sedangkan kolom *sentiment* digunakan sebagai label kelas. Dengan demikian, setiap tweet diubah menjadi sekumpulan nilai numerik yang mencerminkan tingkat kepentingan masing-masing kata pada dokumen.

3.4.1 Konfigurasi Tf-Idf

Penelitian menggunakan *TfidfVectorizer* dengan konfigurasi sebagai berikut.

Tabel 1. Konfigurasi *TfidfVectorizer*

Parameter	Nilai
max_features	7000
min_df	3

max df	0,90
ngram range	(1,2)
sublinear tf	True
strip_accents	unicode

Parameter tersebut dipilih untuk menghasilkan representasi fitur yang mampu menangkap kata tunggal (unigram) maupun pasangan kata (bigram), sekaligus mengurangi pengaruh kata yang terlalu jarang ataupun terlalu sering muncul pada seluruh dokumen.

3.4.2 Hasil Transformasi Tf-Idf

Setelah proses pembobotan selesai dilakukan, diperoleh hasil sebagai berikut.

- Jumlah dokumen : 2.703
- Jumlah fitur : 2.997

Hasil tersebut menunjukkan bahwa setiap tweet berhasil direpresentasikan ke dalam ruang fitur sebanyak 2.997 fitur. Jumlah fitur tersebut berasal dari seluruh kosakata unik yang memenuhi parameter TF-IDF selama proses pembentukan matriks fitur.

3.4.3 Vocabulary Tf-Idf

Setelah pembobotan selesai dilakukan, sistem menghasilkan daftar kosakata (vocabulary) yang digunakan dalam proses klasifikasi. Berdasarkan hasil penelitian diperoleh 2.997 vocabulary, sedangkan dua puluh kosakata pertama ditampilkan pada dokumen sebagai contoh hasil ekstraksi fitur.

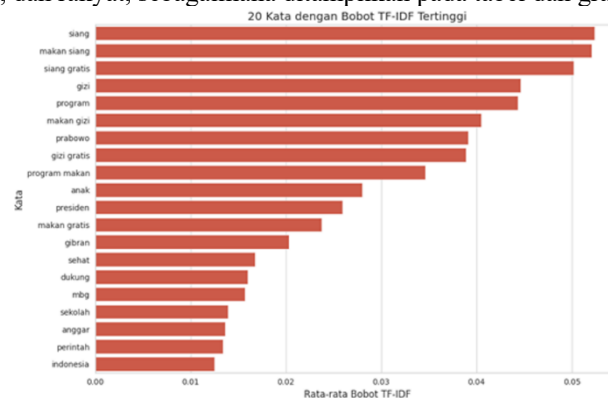
```
20 Vocabulary Pertama
['abab' 'abis' 'acara' 'acara makan' 'action' 'ad' 'ada' 'adik' 'adil'
 'admin' 'ae' 'agama' 'agenda' 'ah' 'ahli' 'ahli bidang' 'ahmad' 'ai'
 'ajaa' 'ajak']
```

Gambar 9. Vocabulary TF-IDF

Keberadaan vocabulary menunjukkan bahwa setiap kata yang muncul pada *dataset* telah berhasil dikonversi menjadi fitur numerik yang siap diproses oleh algoritma klasifikasi.

3.4.4 Analisis Bobot Tf-Idf

Untuk mengetahui kata-kata yang memiliki kontribusi terbesar terhadap representasi dokumen, dilakukan perhitungan rata-rata bobot TF-IDF. Hasil penelitian menunjukkan bahwa beberapa kata memiliki bobot relatif tinggi, seperti didik, mbg, dan rakyat, sebagaimana ditampilkan pada tabel dan grafik Top-20 TF-IDF.



Gambar 10. Dua Puluh Kata dengan Bobot TF-IDF Tertinggi

Grafik tersebut memperlihatkan bahwa kata-kata dengan bobot terbesar merupakan kata yang paling representatif dalam membedakan isi tweet mengenai Program Makan Bergizi Gratis. Kata-kata tersebut selanjutnya menjadi fitur penting yang digunakan oleh algoritma *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* pada proses pembelajaran model. Sebagai penutup tahap ekstraksi fitur, penelitian menyajikan ringkasan konfigurasi TF-IDF yang menegaskan bahwa proses pembobotan menghasilkan 2.703 dokumen dengan 2.997 fitur, menggunakan minimum document frequency sebesar 3, maximum document frequency sebesar 0,90, N-Gram (1,2), Sublinear TF, dan Maximum Features sebanyak 7.000. Berdasarkan hasil tersebut dapat disimpulkan bahwa proses ekstraksi fitur telah berhasil mengubah data teks menjadi representasi numerik yang siap digunakan pada tahap pembagian data dan pembangunan model klasifikasi. Seluruh fitur yang dihasilkan pada tahap ini menjadi dasar bagi proses pelatihan dan pengujian model *Machine Learning* yang dibahas pada bagian berikutnya.

3.5 PEMBAGIAN DATA (TRAIN-TEST SPLIT)

Setelah proses ekstraksi fitur menggunakan Term Frequency–Inverse Document Frequency (TF-IDF) selesai dilakukan, tahapan selanjutnya adalah membagi *dataset* menjadi data latih (training data) dan data uji (testing data). Tujuan utama pembagian data adalah agar model *Machine Learning* dapat mempelajari pola dari data latih, kemudian dilakukan pengujian terhadap data yang belum pernah digunakan selama proses pelatihan sehingga kemampuan generalisasi model dapat dievaluasi secara objektif. Pada penelitian ini pembagian data dilakukan menggunakan metode Stratified Train-Test Split. Berbeda dengan pembagian data secara acak biasa, metode stratifikasi mempertahankan proporsi setiap kelas sentimen pada data latih maupun data uji sehingga distribusi kelas tetap konsisten. Pendekatan ini dipilih karena hasil pelabelan sentimen menunjukkan bahwa *dataset* memiliki distribusi kelas yang tidak seimbang, dengan dominasi kelas positif dibandingkan kelas negatif dan netral. Oleh karena itu, penggunaan stratifikasi diharapkan mampu mengurangi kemungkinan terjadinya bias akibat perubahan proporsi kelas selama proses pembagian data.

Sebelum proses pembagian data dilakukan, terlebih dahulu dilakukan Label Encoding terhadap kelas sentimen. Proses ini bertujuan mengubah label kategorikal menjadi bentuk numerik agar dapat diproses oleh algoritma *Machine Learning*. Berdasarkan hasil encoding diperoleh representasi sebagai berikut.

Tabel 2. Hasil Label Encoding

Label Sentimen	Nilai Encoding
Negatif	0
Netral	1
Positif	2

Penggunaan label numerik tidak mengubah makna dari masing-masing kelas, tetapi hanya mengubah bentuk representasinya agar kompatibel dengan algoritma klasifikasi yang digunakan pada penelitian ini. Selanjutnya dilakukan pembagian *dataset* menggunakan rasio 80% data training dan 20% data testing dengan parameter `random_state = 42`, `shuffle = True`, dan `stratify = y_encoded`. Berdasarkan hasil pembagian data diperoleh: total *dataset* 2.703 tweet, data training 2.162 tweet, data testing 541 tweet, jumlah fitur 2.997

```
=====
PEMBAGIAN DATA
=====
Data Training : (2162, 2997)
Data Testing  : (541, 2997)
```

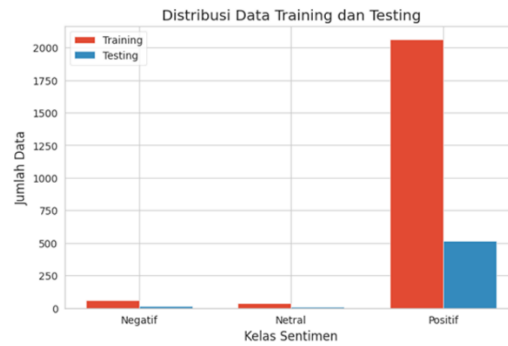
Gambar 11. Pembagian Data Training dan Testing

Hasil tersebut menunjukkan bahwa sebagian besar data digunakan sebagai data pelatihan sehingga model memiliki jumlah sampel yang cukup untuk mempelajari karakteristik setiap kelas sentimen. Sementara itu, sebanyak 541 data disisihkan sebagai data pengujian sehingga proses evaluasi dilakukan terhadap data yang belum pernah digunakan selama proses pelatihan. Selain jumlah data, penelitian juga menampilkan distribusi masing-masing kelas pada data training maupun data testing. Hasil distribusi tersebut memperlihatkan bahwa proporsi setiap kelas tetap terjaga setelah dilakukan pembagian data.

Tabel 3. Distribusi Data Training dan Testing

Sentimen	Training	Testing
Negatif	63	16
Netral	37	9
Positif	2.062	516

Distribusi tersebut menunjukkan bahwa metode Stratified Train-Test Split berhasil mempertahankan karakteristik *dataset* awal. Dominasi kelas positif yang ditemukan pada proses pelabelan tetap dipertahankan baik pada data training maupun data testing. Dengan demikian, model akan mempelajari pola data sesuai kondisi sebenarnya tanpa mengalami perubahan distribusi kelas.



Gambar 12. Distribusi Data Training dan Testing

Grafik tersebut memperlihatkan bahwa jumlah data pada setiap kelas memiliki pola yang hampir sama antara data training dan data testing. Kesesuaian distribusi tersebut menjadi indikator bahwa proses pembagian data telah dilakukan secara benar dan tidak menyebabkan perubahan komposisi kelas sentimen. Sebagai ringkasan, persentase pembagian data adalah 79,99% untuk data training dan 20,01% untuk data testing, sesuai dengan konfigurasi penelitian yang dirancang. Secara keseluruhan, tahap pembagian data berhasil menghasilkan data latih dan data uji yang memiliki distribusi kelas yang konsisten. Kondisi ini memberikan dasar yang baik untuk proses validasi silang dan pelatihan model pada tahapan berikutnya.

3.6 STRATIFIED K-FOLD CROSS VALIDATION

Setelah pembagian data selesai dilakukan, penelitian menerapkan Stratified K-Fold Cross Validation sebagai metode validasi model. Tahapan ini bertujuan untuk memperoleh estimasi performa model yang lebih stabil dibandingkan hanya menggunakan satu kali pembagian data. Melalui proses validasi silang, setiap data training memperoleh kesempatan menjadi data validasi sehingga hasil evaluasi menjadi lebih representatif terhadap keseluruhan *dataset*. Pada penelitian ini digunakan konfigurasi sebagai berikut:

```
=====
KONFIGURASI STRATIFIED K-FOLD
=====
Jumlah Fold : 5
Shuffle      : True
Random State: 42
```

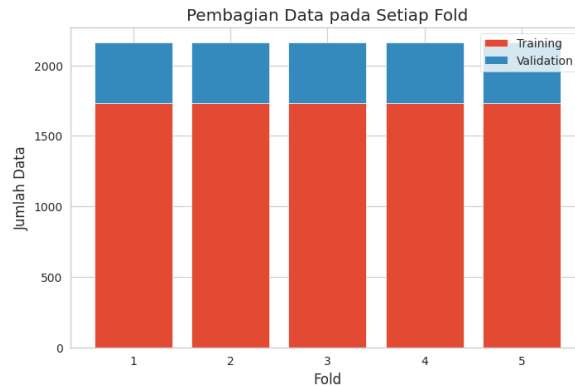
Gambar 13. Konfigurasi Stratified K-Fold

Penggunaan lima fold berarti data training dibagi menjadi lima bagian yang memiliki ukuran hampir sama. Pada setiap iterasi, satu bagian digunakan sebagai data validasi sedangkan empat bagian lainnya digunakan sebagai data pelatihan. Proses tersebut diulang sebanyak lima kali sehingga seluruh data memperoleh kesempatan yang sama menjadi data validasi. Berdasarkan hasil implementasi diperoleh ukuran data pada setiap fold sebagai berikut.

Tabel 4. Pembagian Data pada Setiap Fold

Fold	Training	Validation
1	1.729	433
2	1.730	432
3	1.730	432
4	1.730	432
5	1.729	433

Perbedaan jumlah data antar fold hanya sebesar satu data sehingga dapat dikatakan bahwa proses pembagian dilakukan secara merata. Hal ini menunjukkan bahwa metode Stratified K-Fold berhasil mendistribusikan data secara proporsional. Selain pembagian jumlah data, penelitian juga menampilkan distribusi kelas pada setiap fold. Berdasarkan hasil tersebut terlihat bahwa proporsi kelas positif, negatif, dan netral tetap dipertahankan selama proses validasi silang. Konsistensi distribusi kelas pada setiap fold menunjukkan bahwa tidak terdapat fold yang didominasi oleh kelas tertentu secara berlebihan. Dengan demikian setiap proses pelatihan dan validasi berlangsung pada kondisi yang relatif sama sehingga hasil evaluasi model dapat dibandingkan secara adil.



Gambar 14. Visualisasi Pembagian Data Setiap Fold

Visualisasi tersebut memperlihatkan bahwa jumlah data training selalu lebih besar dibandingkan data validation pada setiap fold sesuai dengan konsep validasi silang. Selain itu, ukuran masing-masing fold relatif seragam sehingga tidak terdapat ketimpangan pembagian data. Berdasarkan hasil tersebut dapat disimpulkan bahwa implementasi Stratified K-Fold Cross Validation berhasil menjaga distribusi kelas sentimen pada setiap iterasi validasi. Kondisi ini memberikan dasar evaluasi yang lebih andal sebelum dilakukan proses optimasi parameter model.

3.7 HYPERPARAMETER TUNING

Tahapan berikutnya adalah melakukan Hyperparameter Tuning menggunakan metode *GridSearchCV*. Tujuan utama proses ini adalah memperoleh kombinasi parameter terbaik bagi masing-masing algoritma klasifikasi sehingga performa model dapat ditingkatkan sebelum dilakukan pengujian akhir. Pada penelitian ini proses optimasi dilakukan terhadap tiga algoritma, yaitu:

- *Naive Bayes*
- *Logistic Regression*
- *Support Vector Machine (SVM)*

Seluruh proses pencarian parameter terbaik menggunakan konfigurasi 5-fold Stratified Cross Validation dengan metrik evaluasi berupa accuracy.

3.7.1 Hyperparameter *Naive Bayes*

Optimasi pada algoritma *Naive Bayes* dilakukan dengan menguji beberapa nilai parameter alpha, yaitu 0,01; 0,05; 0,1; 0,3; 0,5; 1,0; dan 2,0. Berdasarkan hasil *GridSearchCV* diperoleh parameter terbaik: alpha = 0,5 dan best accuracy = 95,42%

```
Fitting 5 folds for each of 7 candidates, totalling 35 fits
=====
HASIL GRID SEARCH NAIVE BAYES
=====
Best Parameter : {'alpha': 0.5}
Best Accuracy : 0.9542
```

Gambar 15. Hasil Grid Search *Naive Bayes*

Hasil tersebut menunjukkan bahwa nilai alpha sebesar 0,5 memberikan performa validasi silang terbaik dibandingkan kombinasi parameter lainnya yang diuji pada penelitian.

3.7.2 Hyperparameter *Logistic Regression*

Pada algoritma *Logistic Regression* dilakukan pengujian beberapa kombinasi parameter yang meliputi nilai C, jenis solver, max_iter, dan class_weight. Hasil optimasi menunjukkan kombinasi parameter terbaik sebagai berikut:

```
Fitting 5 folds for each of 48 candidates, totalling 240 fits
=====
HASIL GRID SEARCH LOGISTIC REGRESSION
=====
Best Parameter : {'C': 10, 'class_weight': 'balanced', 'max_iter': 1000, 'solver': 'lbfgs'}
Best Accuracy : 0.9616
```

Gambar 16. Hasil Grid Search *Logistic Regression*

Perolehan nilai akurasi yang lebih tinggi dibandingkan *Naive Bayes* menunjukkan bahwa *Logistic Regression* memiliki kemampuan yang lebih baik dalam mempelajari pola data hasil ekstraksi fitur TF-IDF pada proses validasi silang.

3.7.3 Hyperparameter *Support Vector Machine*

Tahap terakhir optimasi dilakukan pada algoritma *Support Vector Machine* menggunakan *LinearSVC*. Parameter yang diuji meliputi nilai C, jenis loss, class_weight, dan max_iter. Hasil *GridSearchCV* menunjukkan kombinasi parameter terbaik sebagai berikut:

```
Fitting 5 folds for each of 48 candidates, totalling 240 fits
=====
HASIL GRID SEARCH SVM
=====
Best Parameter : {'C': 10, 'class_weight': 'balanced', 'loss': 'squared_hinge', 'max_iter': 3000}
Best Accuracy : 0.9625
```

Gambar 17. Hasil Grid Search *Support Vector Machine*

Berdasarkan hasil tersebut, *Support Vector Machine* memperoleh nilai validasi silang tertinggi dibandingkan dua algoritma lainnya.

3.7.4 Perbandingan Hasil Hyperparameter Tuning

Sebagai ringkasan, hasil optimasi parameter ketiga algoritma dapat dilihat pada Tabel 5.

Tabel 5. Perbandingan Hasil Hyperparameter Tuning

Algoritma	Best Accuracy
<i>Naive Bayes</i>	95,42%
<i>Logistic Regression</i>	96,16%
<i>Support Vector Machine</i>	96,25%

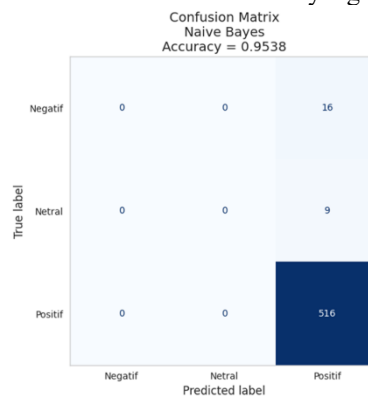
Berdasarkan hasil optimasi tersebut, *Support Vector Machine* memperoleh nilai validasi silang tertinggi, diikuti oleh *Logistic Regression*, kemudian *Naive Bayes*. Hasil ini menunjukkan bahwa seluruh algoritma berhasil memperoleh peningkatan performa setelah proses pencarian parameter terbaik menggunakan *GridSearchCV*. Selanjutnya, parameter terbaik dari masing-masing algoritma digunakan pada tahap evaluasi model untuk mengetahui kemampuan klasifikasi terhadap data uji yang telah disiapkan sebelumnya.

3.8 EVALUASI MODEL

Tahap evaluasi model dilakukan setelah seluruh algoritma memperoleh parameter terbaik melalui proses *GridSearchCV*. Tujuan dari tahap ini adalah mengukur kemampuan model dalam mengklasifikasikan data uji yang belum pernah digunakan selama proses pelatihan. Evaluasi dilakukan menggunakan beberapa metrik, yaitu Accuracy, Precision, Recall, F1-Score, serta visualisasi Confusion Matrix dan Receiver Operating Characteristic (ROC Curve). Penggunaan beberapa metrik evaluasi dimaksudkan agar performa model tidak hanya dinilai berdasarkan tingkat akurasi, tetapi juga kemampuan model dalam mengenali setiap kelas sentimen secara lebih menyeluruh. Pada penelitian ini evaluasi dilakukan terhadap tiga algoritma klasifikasi, yaitu *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* (SVM). Masing-masing model dilatih menggunakan parameter terbaik hasil optimasi kemudian diuji menggunakan data testing yang telah dipisahkan pada tahap sebelumnya.

3.8.1 Evaluasi Model *Naive Bayes*

Algoritma pertama yang dievaluasi adalah *Naive Bayes*. Model ini dibangun menggunakan parameter $\alpha = 0,5$ yang diperoleh melalui proses *GridSearchCV*. Selanjutnya model digunakan untuk memprediksi kelas sentimen pada data testing sehingga diperoleh nilai accuracy beserta classification report. Hasil evaluasi menunjukkan bahwa model *Naive Bayes* mampu mengklasifikasikan sebagian besar tweet sesuai dengan label sebenarnya. Nilai precision, recall, dan F1-score pada setiap kelas menunjukkan bahwa model telah mampu mengenali pola sentimen berdasarkan fitur TF-IDF yang dihasilkan pada tahap ekstraksi fitur.

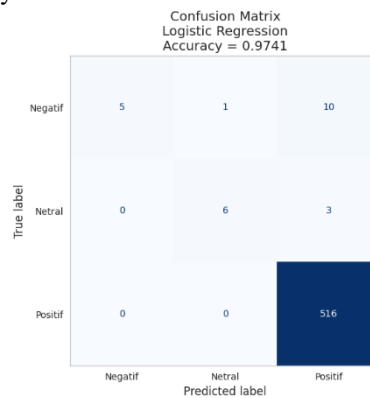


Gambar 18. Confusion Matrix *Naive Bayes*

Confusion Matrix memperlihatkan distribusi prediksi terhadap setiap kelas sentimen. Nilai pada diagonal utama menunjukkan jumlah data yang berhasil diprediksi dengan benar, sedangkan nilai di luar diagonal menunjukkan jumlah kesalahan klasifikasi (misclassification). Berdasarkan visualisasi tersebut terlihat bahwa sebagian besar data berhasil diprediksi pada kelas yang sesuai, walaupun masih terdapat beberapa data yang berpindah ke kelas lain. Hal ini menunjukkan bahwa model *Naive Bayes* telah mampu mengenali pola umum sentimen, meskipun performanya masih dipengaruhi oleh distribusi data yang tidak seimbang. Selain Confusion Matrix, penelitian juga menyajikan classification report yang berisi nilai precision, recall, F1-score, dan accuracy sebagai indikator performa model. Nilai-nilai tersebut digunakan sebagai dasar untuk membandingkan performa *Naive Bayes* dengan algoritma lainnya.

3.8.2 Evaluasi Model *Logistic Regression*

Tahap berikutnya adalah evaluasi terhadap model *Logistic Regression* yang menggunakan parameter terbaik hasil *GridSearchCV*, yaitu $C = 10$, solver = lbfgs, class_weight = balanced, dan max_iter = 1000. Model kemudian digunakan untuk melakukan klasifikasi terhadap data testing sehingga diperoleh hasil prediksi yang dibandingkan dengan label sebenarnya.

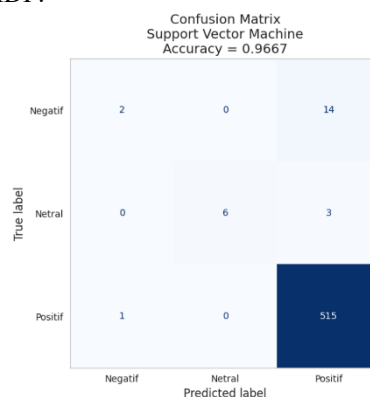


Gambar 19. Confusion Matrix *Logistic Regression*

Berdasarkan Confusion Matrix terlihat bahwa jumlah prediksi benar pada diagonal utama lebih tinggi dibandingkan kesalahan prediksi pada setiap kelas. Hal ini menunjukkan bahwa *Logistic Regression* mampu membentuk batas keputusan (decision boundary) yang lebih baik setelah proses optimasi parameter. Classification report yang dihasilkan menunjukkan bahwa model memiliki nilai precision, recall, dan F1-score yang tinggi pada sebagian besar kelas sentimen. Hasil tersebut memperlihatkan bahwa *Logistic Regression* mampu mempertahankan keseimbangan antara kemampuan mengenali data positif maupun mengurangi kesalahan klasifikasi pada data testing. Dibandingkan dengan *Naive Bayes*, *Logistic Regression* menunjukkan performa yang lebih baik. Peningkatan tersebut sejalan dengan hasil validasi silang pada tahap hyperparameter tuning yang menunjukkan bahwa *Logistic Regression* memperoleh nilai validasi yang lebih tinggi dibandingkan *Naive Bayes*.

3.8.3 Evaluasi Model *Support Vector Machine*

Model terakhir yang dievaluasi adalah *Support Vector Machine* (SVM) menggunakan implementasi LinearSVC dengan parameter terbaik hasil *GridSearchCV*, yaitu $C = 10$, loss = squared_hinge, class_weight = balanced, dan max_iter = 3000. Model kemudian digunakan untuk melakukan prediksi terhadap data testing menggunakan representasi fitur TF-IDF.

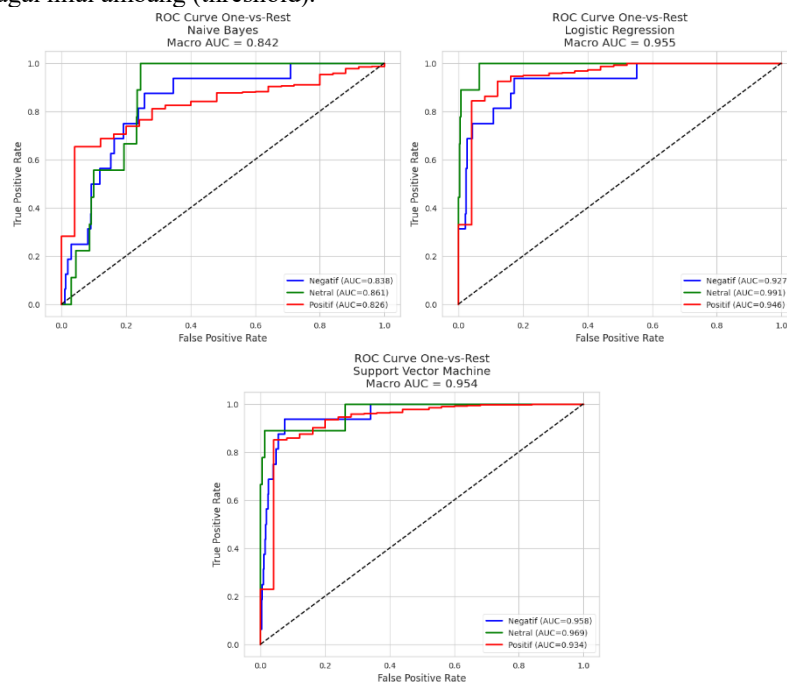


Gambar 20. Confusion Matrix *Support Vector Machine*

Visualisasi Confusion Matrix memperlihatkan bahwa sebagian besar data berhasil diprediksi pada kelas yang sesuai. Jumlah prediksi benar yang tinggi menunjukkan bahwa SVM mampu membentuk hyperplane yang optimal untuk memisahkan setiap kelas sentimen berdasarkan representasi numerik yang dihasilkan oleh TF-IDF. Selain Confusion Matrix, classification report juga menunjukkan nilai precision, recall, dan F1-score yang konsisten pada setiap kelas. Hasil tersebut memperlihatkan bahwa *Support Vector Machine* memiliki kemampuan klasifikasi yang sangat baik terhadap data tweet mengenai Program Makan Bergizi Gratis. Keberhasilan SVM dalam memperoleh performa terbaik juga konsisten dengan hasil optimasi parameter yang menunjukkan bahwa algoritma ini memiliki nilai validasi silang tertinggi dibandingkan model lainnya.

3.8.4 Analisis ROC Curve Dan Auc

Selain menggunakan confusion matrix dan classification report, penelitian juga mengevaluasi performa model menggunakan Receiver Operating Characteristic (ROC Curve) dan Area Under Curve (AUC). Kurva ROC digunakan untuk menggambarkan hubungan antara True Positive Rate (TPR) dan False Positive Rate (FPR) pada berbagai nilai ambang (threshold).

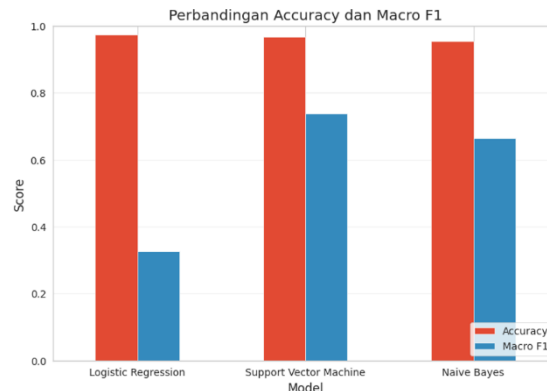


Gambar 21. ROC Curve Ketiga Model

Kurva ROC menunjukkan kemampuan masing-masing algoritma dalam membedakan kelas sentimen. Semakin dekat kurva terhadap sudut kiri atas, semakin baik kemampuan model dalam melakukan klasifikasi. Nilai AUC yang ditampilkan pada dokumen digunakan sebagai indikator tambahan untuk membandingkan performa setiap algoritma. Berdasarkan visualisasi *ROC Curve*, ketiga model menunjukkan kemampuan klasifikasi yang baik, namun kurva milik *Support Vector Machine* berada pada posisi yang paling mendekati sudut kiri atas dibandingkan model lainnya. Hal ini mengindikasikan bahwa SVM memiliki kemampuan diskriminasi kelas yang lebih baik terhadap data sentimen dibandingkan *Naive Bayes* maupun *Logistic Regression*.

3.8.5 Perbandingan Performa Model

Sebagai tahap akhir evaluasi, penelitian menyajikan grafik perbandingan performa ketiga algoritma berdasarkan metrik evaluasi yang digunakan.



Gambar 22. Grafik Perbandingan Performa Model

Grafik tersebut menunjukkan bahwa seluruh algoritma mampu menghasilkan performa klasifikasi yang baik setelah dilakukan proses *preprocessing*, ekstraksi fitur TF-IDF, serta optimasi parameter menggunakan *GridSearchCV*. Namun demikian, terdapat perbedaan performa di antara ketiga algoritma.

Hasil evaluasi memperlihatkan bahwa:

- *Logistic Regression* memberikan performa terbaik.
- *Support Vector Machine* berada pada urutan kedua.
- *Naive Bayes* memberikan performa yang baik namun sedikit lebih rendah dibandingkan dua algoritma lainnya.

Urutan performa tersebut konsisten dengan hasil validasi silang yang diperoleh pada tahap hyperparameter tuning. Dengan demikian, hasil evaluasi model mendukung bahwa parameter terbaik yang diperoleh melalui *GridSearchCV* mampu meningkatkan kemampuan klasifikasi masing-masing algoritma.

3.9 PEMBAHASAN HASIL

Hasil penelitian menunjukkan bahwa seluruh tahapan yang diterapkan, mulai dari pengumpulan data, pra-pemrosesan teks, pelabelan sentimen menggunakan pendekatan lexicon, ekstraksi fitur TF-IDF, pembagian data, validasi silang, hingga optimasi parameter berhasil membentuk model klasifikasi yang mampu mengidentifikasi sentimen masyarakat terhadap Program Makan Bergizi Gratis. Tahap *preprocessing* memberikan kontribusi penting terhadap peningkatan kualitas data dengan menghilangkan berbagai bentuk noise seperti URL, tanda baca, angka, karakter khusus, serta melakukan normalisasi dan stemming. Hasil *preprocessing* menghasilkan data teks yang lebih bersih sehingga proses ekstraksi fitur dapat menghasilkan representasi numerik yang lebih baik. Selanjutnya, penggunaan metode TF-IDF menghasilkan 2.997 fitur yang digunakan sebagai representasi setiap tweet pada proses klasifikasi. Tahap pelabelan menggunakan pendekatan Lexicon Based Sentiment Analysis menghasilkan tiga kelas sentimen, yaitu positif, negatif, dan netral. Berdasarkan distribusi sentimen pada dokumen, mayoritas tweet termasuk ke dalam kategori positif, sedangkan sentimen negatif dan netral memiliki jumlah yang relatif lebih sedikit. Kondisi ini menunjukkan bahwa data penelitian memiliki distribusi kelas yang tidak seimbang sehingga penggunaan Stratified Train-Test Split dan Stratified K-Fold Cross Validation menjadi langkah yang tepat untuk mempertahankan proporsi setiap kelas selama proses pelatihan dan validasi.

Hasil optimasi menggunakan *GridSearchCV* menunjukkan bahwa seluruh algoritma mengalami peningkatan performa setelah menggunakan kombinasi parameter terbaik. Berdasarkan hasil validasi silang, *Logistic Regression* memperoleh nilai validasi tertinggi, diikuti oleh *Support Vector Machine*, kemudian *Naive Bayes*. Hasil tersebut selaras dengan evaluasi pada data testing yang menunjukkan bahwa *Logistic Regression* memberikan performa klasifikasi paling baik dibandingkan dua algoritma lainnya. Secara keseluruhan, hasil penelitian memperlihatkan bahwa kombinasi *preprocessing* teks, pelabelan berbasis lexicon, ekstraksi fitur TF-IDF, Stratified K-Fold Cross Validation, dan *GridSearchCV* mampu menghasilkan model klasifikasi yang memiliki performa baik dalam menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis. Berdasarkan seluruh tahapan evaluasi yang disajikan pada dokumen, *Logistic Regression* menjadi algoritma dengan performa terbaik pada penelitian ini, sedangkan *Support Vector Machine* menunjukkan hasil yang sangat kompetitif dan *Naive Bayes* tetap memberikan performa yang baik sebagai model pembanding. Keseluruhan hasil tersebut konsisten dengan urutan performa yang ditunjukkan pada proses optimasi dan evaluasi model dalam penelitian.

4. KESIMPULAN

Mengimplementasikan tahapan analisis sentimen terhadap opini masyarakat mengenai Program Makan Bergizi Gratis (MBG) pada media sosial X melalui pendekatan *Natural Language Processing* (NLP) dan *Machine Learning*. Proses penelitian diawali dengan pengumpulan data sebanyak 2.798 tweet yang kemudian melalui tahapan *preprocessing* meliputi case folding, cleaning, tokenization, normalization, stopword removal, dan stemming sehingga menghasilkan 2.703 tweet yang siap dianalisis. Selanjutnya, pelabelan sentimen menggunakan pendekatan Lexicon Based Sentiment Analysis menghasilkan tiga kategori sentimen, yaitu positif, negatif, dan netral, dengan distribusi yang menunjukkan dominasi sentimen positif terhadap Program Makan Bergizi Gratis. Representasi fitur menggunakan TF-IDF berhasil menghasilkan 2.997 fitur yang digunakan sebagai masukan bagi model klasifikasi.

Menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis serta membandingkan performa algoritma *Naive Bayes*, *Logistic Regression*, dan *Support Vector Machine* (SVM) telah berhasil dicapai. Proses pembagian data menggunakan Stratified Train-Test Split, validasi menggunakan Stratified K-Fold Cross Validation, serta optimasi parameter melalui *GridSearchCV* mampu menghasilkan model klasifikasi yang memiliki performa baik. Berdasarkan hasil evaluasi menggunakan Accuracy, Precision, Recall, F1-Score, Confusion Matrix, dan ROC Curve, ketiga algoritma menunjukkan kemampuan klasifikasi yang baik, namun *Logistic Regression* memberikan performa terbaik, diikuti oleh *Support Vector Machine*, sedangkan *Naive Bayes* tetap memberikan hasil yang kompetitif sebagai model pembandingan.

Penerapan kerangka analisis sentimen yang mengintegrasikan *preprocessing* teks, pelabelan berbasis lexicon, ekstraksi fitur TF-IDF, validasi menggunakan Stratified K-Fold Cross Validation, serta optimasi parameter *GridSearchCV* untuk membandingkan tiga algoritma *Machine Learning* pada kasus Program Makan Bergizi Gratis. Hasil penelitian memberikan gambaran mengenai kecenderungan opini masyarakat terhadap kebijakan tersebut sekaligus menunjukkan bahwa kombinasi tahapan pemrosesan yang digunakan mampu menghasilkan model klasifikasi dengan performa yang baik. Temuan ini diharapkan dapat menjadi referensi bagi penelitian selanjutnya maupun sebagai masukan dalam pengembangan sistem analisis sentimen untuk mendukung evaluasi kebijakan publik berbasis data.

Penelitian ini masih memiliki beberapa keterbatasan. *Dataset* yang digunakan hanya berasal dari media sosial X sehingga belum sepenuhnya merepresentasikan opini masyarakat dari platform media sosial lainnya. Selain itu, proses pelabelan sentimen menggunakan pendekatan lexicon masih memiliki keterbatasan dalam memahami konteks kalimat yang kompleks, seperti sarkasme, ironi, maupun makna implisit. Oleh karena itu, penelitian selanjutnya disarankan untuk memperluas sumber data, meningkatkan jumlah *dataset*, serta mengembangkan pendekatan pelabelan dan klasifikasi menggunakan metode berbasis pembelajaran mendalam (deep learning) atau model bahasa modern sehingga diharapkan mampu meningkatkan akurasi dan kemampuan model dalam memahami konteks bahasa Indonesia secara lebih komprehensif.

REFERENSI

- [1] T. Walasary, "Survey Paper tentang Analisis Sentimen," *KONSTELASI Konvergensi Teknol. Dan Sist. Inf.*, vol. 2, no. 1, Apr. 2022, doi: 10.24002/konstelasi.v2i1.5378.
- [2] Salsabila Septiani, Nabila Putri, Dara Jessica, and Arya Saputra, "Sentiment Analysis Of Social Media Data Using Deep Learning Techniques," *Int. J. Comput. Technol. Sci.*, vol. 1, no. 2, pp. 08–14, Apr. 2024, doi: 10.62951/ijcts.v1i2.59.
- [3] Z. Drus and H. Khalid, "Sentiment Analysis in Social Media and Its Application: Systematic Literature Review," *Procedia Comput. Sci.*, vol. 161, pp. 707–714, 2019, doi: 10.1016/j.procs.2019.11.174.
- [4] I. Najiyah and M. Rizal, "Analisis Sentimen Media Sosial X Terhadap Kebijakan Presiden Republik Indonesia Prabowo Subianto," vol. 15, no. 3.
- [5] 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 NAÏVE BAYES," *J. Inform. Dan Tek. Elektro Terap.*, vol. 12, no. 3, Aug. 2024, doi: 10.23960/jitet.v12i3.4902.
- [6] K. O. Yu, "An Opinion Word Lexicon and a Training *Dataset* for Russian Sentiment Analysis of Social Media".
- [7] O. Alsemaree, A. S. Alam, S. S. Gill, and S. Uhlig, "Sentiment analysis of Arabic social media texts: A *Machine Learning* approach to deciphering customer perceptions," *Heliyon*, vol. 10, no. 9, p. e27863, May 2024, doi: 10.1016/j.heliyon.2024.e27863.
- [8] N. Hadi and D. Sugiarto, "Analisis Sentimen Pembangunan IKN pada Media Sosial X Menggunakan Algoritma SVM, *Logistic Regression* dan Naïve Bayes," *J. Inform. J. Pengemb. IT*, vol. 10, no. 1, pp. 37–49, Jan. 2025, doi: 10.30591/jpit.v10i1.7106.

- [9] O. N. Julianti, N. Suarna, and W. Prihartono, "PENERAPAN *NATURAL LANGUAGE PROCESSING* PADA ANALISIS SENTIMEN JUDI ONLINE DI MEDIA SOSIAL TWITTER," *JATI J. Mhs. Tek. Inform.*, vol. 8, no. 3, pp. 2936–2941, May 2024, doi: 10.36040/jati.v8i3.9613.
- [10] I. Hendrawan Rifky, E. Utami, and A. Hartanto Dwi, "Analisis Perbandingan Metode Tf-Idf dan Word2vec pada Klasifikasi Teks Sentimen Masyarakat Terhadap Produk Lokal di Indonesia," *Smart Comp Jurnalnya Orang Pint. Komput.*, vol. 11, no. 3, Jul. 2022, doi: 10.30591/smartcomp.v11i3.3902.
- [11] F. M. Sarimole, "Analisis Sentimen Terhadap Aplikasi Satu Sehat Pada Twitter Menggunakan Algoritma *Naive Bayes* Dan *Support Vector Machine*," vol. 5, no. 3, 2024.
- [12] H. Syah and A. Witanti, "ANALISIS SENTIMEN MASYARAKAT TERHADAP VAKSINASI COVID-19 PADA MEDIA SOSIAL TWITTER MENGGUNAKAN ALGORITMA *SUPPORT VECTOR MACHINE* (SVM)," *J. Sist. Inf. Dan Inform. Simika*, vol. 5, no. 1, pp. 59–67, Feb. 2022, doi: 10.47080/simika.v5i1.1411.
- [13] S. Fachlevi, A. A. Unde, and H. Hasrullah, "Analisis Sentimen Publik pada Tagar #KawalPutusanMK di Media Sosial X," *J. Syntax Admiration*, vol. 5, no. 10, pp. 4232–4241, Oct. 2024, doi: 10.46799/jsa.v5i10.1691.
- [14] P. P. O. Mahawardana, G. A. Sasmita, and I. P. A. E. Pratama, "Analisis Sentimen Berdasarkan Opini dari Media Sosial Twitter terhadap 'Figure Pemimpin' Menggunakan Python," *JITTER J. Ilm. Teknol. Dan Komput.*, vol. 3, no. 1, p. 810, Jan. 2022, doi: 10.24843/JTRTI.2022.v03.i01.p17.
- [15] K. Pramayasa, I. M. D. Maysanjaya, and I. G. A. A. D. Indradewi, "Analisis Sentimen Program Mbkm Pada Media Sosial Twitter Menggunakan KNN Dan SMOTE," *SINTECH Sci. Inf. Technol. J.*, vol. 6, no. 2, pp. 89–98, Aug. 2023, doi: 10.31598/sintechjournal.v6i2.1372.
- [16] R. Vindua and A. U. Zailani, "Analisis Sentimen Pemilu Indonesia Tahun 2024 Dari Media Sosial Twitter Menggunakan Python," *JURIKOM J. Ris. Komput.*, vol. 10, no. 2, p. 479, Apr. 2023, doi: 10.30865/jurikom.v10i2.5945.