

# Analisis Sentimen Masyarakat Terhadap Program Makan Bergizi Gratis Menggunakan SVM

<sup>1</sup>Hendra Subastian, <sup>2</sup>Budi Triandi, <sup>3</sup>Roslina

<sup>1,2,3</sup>Ilmu Komputer, Universitas Potensi Utama, Medan, Indonesia

\*Korespondensi: [hendra.stokbinaguna@gmail.com](mailto:hendra.stokbinaguna@gmail.com)

Submit : 26 Mei 2026 | Diterima : 21 Juni 2026 | Terbit : 15 Juli 2026

## ABSTRACT

*The Free Nutritious Milk and Meals program is a strategic national policy that has sparked diverse public opinions in the digital space. This study aims to analyze public sentiment toward the implementation of the program using data from the Twitter/X social media platform. The applied approach is based on text mining utilizing Term Frequency-Inverse Document Frequency (TF-IDF) weighting and classification via the Support Vector Machine (SVM), Logistic Regression dan Naïve Bayes algorithm. Through web scraping techniques, an initial dataset of 6,210 tweets was collected, which was subsequently reduced to 4,121 data points after the deduplication process. Manual labeling revealed an imbalanced dataset, dominated by the positive class (82.76%), followed by negative (9.73%), and neutral (7.58%) sentiments. Prior to modeling, the data underwent preprocessing stages including cleaning, case folding, tokenization, stopword removal, and stemming. The SVM model evaluation demonstrated a solid performance with an accuracy of 0.849, an F1-score of 0.839, and an MCC of 0.444, with the highest sensitivity achieved in the positive category (recall of 0.946). Conversely, the model's performance on minority classes (negative and neutral) remained low due to data imbalance bias. A comparative experiment with Logistic Regression yielded slightly higher accuracy (0.857) and AUC (0.932); however, SVM proved to be more consistent in maintaining classification stability across classes. Overall, this approach effectively maps public response, and future research is recommended to implement imbalanced data handling methods to optimize accuracy in minority classes.*

**Keywords:** Data Imbalance, Free Nutritious Milk and Meals Program, Sentiment Analysis, Support Vector Machine, Text Mining, Twitter/X.

## ABSTRAK

Program Makan Bergizi Gratis merupakan kebijakan strategis nasional yang memicu beragam opini publik di ruang digital. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap implementasi program tersebut menggunakan data dari media sosial Twitter/X. Pendekatan yang diterapkan berbasis *text mining* melalui pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) dan klasifikasi menggunakan algoritma *Support Vector Machine* (SVM), *Logistic Regression* dan *Naïve Bayes*. Melalui teknik *web scraping*, diperoleh dataset awal sebanyak 6.210 twit, yang kemudian menyusut menjadi 4.121 data setelah tahapan deduplikasi. Proses pelabelan manual menunjukkan sebaran sentimen yang tidak seimbang (*imbalanced dataset*), didominasi oleh kelas positif (82,76%), diikuti negatif (9,73%), dan netral (7,58%). Sebelum pemodelan data melalui tahapan *preprocessing* yang meliputi *cleaning*, *case folding*, *tokenization*, *stopword removal*, dan *stemming*. Hasil pengujian model SVM menunjukkan performa yang cukup baik dengan nilai akurasi sebesar 0,849, F1-score 0,839, dan MCC 0,444, di mana sensitivitas tertinggi berada pada kategori positif dengan *recall* mencapai 0,946. Sebaliknya, performa pada kelas minoritas (negatif dan netral) masih tergolong rendah akibat bias ketidakseimbangan data. Eksperimen komparatif dengan *Logistic Regression* menghasilkan akurasi (0,857) dan AUC (0,932) yang sedikit lebih tinggi, namun SVM terbukti lebih konsisten dalam menjaga stabilitas klasifikasi antar-kelas. Secara keseluruhan, pendekatan ini efektif memetakan respons publik, dan disarankan bagi penelitian selanjutnya untuk menerapkan metode penanganan *imbalanced data* guna mengoptimalkan akurasi pada kelas minoritas.

**Kata Kunci:** Analisis Sentimen, Makan Bergizi Gratis, Twitter/X, *Support Vector Machine*, *Text Mining*, Ketidakseimbangan Data.

## PENDAHULUAN

Program Makan Bergizi Gratis yang diinisiasi oleh pemerintah Indonesia merupakan salah satu kebijakan strategis dalam meningkatkan kesejahteraan masyarakat, khususnya bagi pelajar dan kelompok rentan. Kebijakan ini tidak hanya berfokus pada penanggulangan masalah kekurangan gizi, tetapi juga berperan dalam peningkatan kualitas sumber daya manusia (SDM) serta mendukung pembangunan nasional yang berkelanjutan.

Di sisi lain, implementasi kebijakan publik sering kali diikuti oleh beragam respons dari masyarakat. Perkembangan media sosial, khususnya platform Twitter/X, telah menjadi media utama bagi masyarakat untuk menyampaikan opini secara terbuka, baik berupa dukungan, kritik, maupun tanggapan netral. Dinamika opini tersebut mencerminkan persepsi publik yang dapat memengaruhi tingkat kepercayaan terhadap pemerintah, terutama dalam konteks penyebaran informasi yang cepat dan masif (Kiratli, 2023; Parr, 2025). Oleh karena itu, analisis terhadap opini masyarakat di media sosial menjadi penting untuk memahami bagaimana kebijakan publik diterima oleh masyarakat secara luas.

Dalam konteks tersebut, analisis sentimen berbasis Natural Language Processing (NLP) dan Machine Learning (ML) menjadi pendekatan yang efektif untuk mengolah data teks dalam jumlah besar. Penelitian terbaru menunjukkan bahwa teknik NLP mampu mengidentifikasi pola sentimen secara otomatis dan memberikan hasil yang akurat pada data media sosial (T Joseph, 2024; Joyce Y. M. Nip and Benoit Berthelie, 2024). Salah satu algoritma yang banyak digunakan dalam klasifikasi teks adalah Support Vector Machine (SVM), yang dikenal memiliki performa stabil serta efektif dalam menangani data berdimensi tinggi dan berbagai permasalahan klasifikasi (Zhou & Shen, 2022; Pambudi et al., 2024).

Sejumlah penelitian sebelumnya menunjukkan bahwa kinerja algoritma Machine Learning pada analisis sentimen sangat bergantung pada karakteristik data dan teknik pra-proses. (Zahra et al. 2026) menemukan bahwa Linear SVM mampu mengungguli IndoBERT pada dataset ulasan e-commerce Tokopedia dengan akurasi 97,60%. Sementara itu, (Hazrati dan Rahmatulloh, 2026) menunjukkan bahwa Logistic Regression menghasilkan kinerja terbaik dengan akurasi 88,53% pada analisis ulasan LinkedIn. Di sisi lain, (Vazli dan Suryono, 2026) bahwa SVM memberikan hasil terbaik sebesar 91% setelah penerapan SMOTE pada data tweet tentang Pertamina, mengungguli Naïve Bayes dan Random Forest.

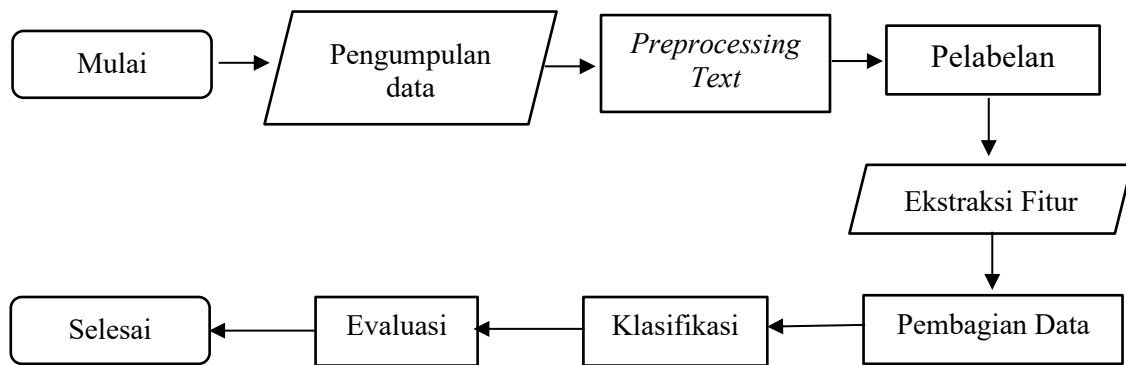
Berdasarkan hasil tersebut menunjukkan bahwa tidak terdapat algoritma tunggal yang selalu unggul dalam seluruh kondisi, karena performa model sangat dipengaruhi oleh jenis data, keseimbangan kelas, serta teknik ekstraksi fitur seperti TF-IDF.

Berdasarkan tinjauan literatur, penelitian-penelitian sebelumnya masih didominasi oleh studi pada domain e-commerce dan media sosial umum, serta belum secara khusus mengkaji penerapan ketiga algoritma tersebut pada konteks kebijakan publik di Indonesia, khususnya program *Makan Bergizi Gratis*. Selain itu, penelitian terdahulu juga belum secara konsisten mengevaluasi perbandingan SVM, Logistic Regression, dan Naïve Bayes pada data Twitter/X berbahasa Indonesia yang bersifat informal, dinamis, dan berpotensi tidak seimbang.

Berdasarkan hal tersebut, penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap program *Makan Bergizi Gratis* menggunakan pendekatan TF-IDF dan algoritma Support Vector Machine (SVM), serta membandingkannya dengan Logistic Regression dan Naïve Bayes. Penelitian ini diharapkan dapat memberikan kontribusi dalam pengembangan analisis sentimen berbasis *machine learning* pada konteks kebijakan publik di Indonesia serta memberikan gambaran empiris mengenai persepsi masyarakat terhadap program pemerintah berdasarkan data media sosial.

## METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif berbasis analisis teks (*text mining*) untuk mengidentifikasi sentimen masyarakat terhadap program makan bergizi gratis menggunakan aplikasi *Orange Data Mining*. Data yang digunakan berupa teks tidak terstruktur yang diperoleh dari media sosial Twitter/X melalui teknik *web scraping*. Adapun tahap penelitian dapat dilihat pada *flowchart* 1 berikut :



Gambar 1. Tahapan penelitian

### Tahapan Pengumpulan Data

Pengumpulan data dilakukan menggunakan teknik *web scraping* melalui *Application Programming Interface* (API) Twitter/X. Proses ini dilakukan dengan bantuan Google Colab menggunakan kata kunci yang relevan, yaitu "makan bergizi gratis" dan "makan bergizi gratis prabowo".

### Tahapan *Preprocessing Text*

Tahap *preprocessing text* bertujuan untuk meningkatkan kualitas data teks sebelum dilakukan analisis. Proses ini meliputi beberapa tahapan, yaitu:

- Deduplikasi*, untuk menghapus data yang identik.
- Cleaning Text*, untuk menghilangkan elemen tidak relevan seperti URL, *mention*, *hashtag*, simbol, dan angka.
- Case Folding*, untuk menyeragamkan teks menjadi huruf kecil.
- Tokenisasi*, untuk memecah teks menjadi unit kata.
- Stopword Removal*, digunakan untuk menghapus kata-kata tidak bermakna sentimen. Daftar stop word yang di gunakan mengacu pada standar NLP klasik diperkenalkan oleh Fadillah Z. Tala (2003) dan di kombinasi dengan kata-kata umum yang sering muncul pada twitter.
- Stemming*, untuk mengubah kata menjadi bentuk dasar. Proses stemming dilakukan menggunakan pustaka Sastrawi yang mengimplementasikan algoritma untuk Bahasa Indonesia (Mirna adriani et al.,2007).

### Tahapan Pelabelan Data

Pelabelan dilakukan secara manual oleh peneliti dengan mengelompokkan setiap data ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral. Proses ini mempertimbangkan konteks kalimat secara keseluruhan untuk meminimalkan subjektivitas. Data yang telah diberi label digunakan sebagai *ground truth* dalam proses pelatihan dan pengujian model.

Sentimen positif diberikan pada teks yang mengandung ungkapan dukungan, apresiasi, kepuasan, atau pernyataan yang bersifat membangun. Sentimen negatif diberikan pada teks yang memuat kritik, keluhan, ketidakpuasan, atau ungkapan bernada penolakan. Sementara itu, sentimen netral diberikan pada teks yang bersifat informatif, tidak menunjukkan emosi yang kuat, serta tidak memihak pada salah satu sentimen tertentu.

### Tahapan Ekstraksi Fitur

Data teks yang telah diproses kemudian direpresentasikan ke dalam bentuk numerik menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*). Metode ini telah banyak dimanfaatkan dalam berbagai tugas pengolahan teks, seperti ekstraksi kata kunci, pengukuran kemiripan dokumen, serta klasifikasi topik (HAO LIU et al., 2022). Adapun rumus yang digunakan pada TF-IDF sebagai berikut :

Term Frequency (TF) digunakan untuk menghitung jumlah kemunculan suatu term dalam dokumen, dapat dilihat pada persamaan (1).

$$TF_{(t\ d)} = F_{(t\ d)} \quad (1)$$

Inverse Document Frequency (IDF) digunakan untuk mengukur tingkat keunikan suatu term dalam seluruh dokumen, dapat dilihat pada persamaan (2).

$$IDF = \log_{10} \left( \frac{N}{df} \right) \quad (2)$$

Nilai TF-IDF diperoleh dari hasil perkalian antara TF dan IDF, dapat dilihat pada persamaan (3)

$$TFIDF = TF \times IDF \quad (3)$$

Untuk mengurangi bias akibat perbedaan panjang dokumen, dilakukan normalisasi menggunakan L2 norm, dapat dilihat pada persamaan (4).

$$L2 = \sqrt{\sum (TFIDF^2)} \quad (4)$$

Selanjutnya, nilai TF-IDF dinormalisasi, dapat dilihat pada persamaan (5).

$$TFIDF_{norm} = \frac{TFIDF}{L2} \quad (5)$$

Dimana  $t$  merupakan term (kata),  $d$  adalah dokumen,  $N$  merupakan jumlah total dokumen dan  $df$  adalah jumlah dokumen yang mengandung term  $tf$ . Hasil akhir dari proses ini berupa matriks TF-IDF ternormalisasi yang digunakan sebagai input pada algoritma klasifikasi Support Vector Machine (SVM).

### Tahapan Pembagian Data (*Data Splitting*)

Data yang telah melalui tahap ekstraksi fitur selanjutnya dibagi menjadi data latih (*training data*) dan data uji (*testing data*). Pembagian data dilakukan untuk mengevaluasi kinerja model klasifikasi secara objektif.

Pada penelitian ini, data dibagi dengan rasio 80:20, di mana 80% digunakan sebagai data latih dan 20% sebagai data uji. Data latih digunakan untuk membangun model klasifikasi, sedangkan data uji digunakan untuk mengukur performa model terhadap data yang belum pernah dilihat sebelumnya.

### Tahapan Klasifikasi Model

- a. Support Vector Machine : Penelitian ini menerapkan SVM dengan kernel linear karena efektif dalam menangani data teks berdimensi tinggi hasil ekstraksi TF-IDF, serta mampu menghasilkan performa klasifikasi yang stabil dan akurat pada berbagai tugas klasifikasi teks (Risnawati & Wibowo, 2024). Pengaturan parameter model meliputi nilai *Cost* (C) sebesar 5 untuk mengontrol *trade-off* antara margin dan kesalahan klasifikasi, maksimum iterasi sebanyak 2.000, serta toleransi kesalahan sebesar 0,001 demi memastikan konvergensi model. Secara matematis, fungsi keputusan pada SVM linear dapat dinyatakan pada persamaan (6)

$$f(x) = w \cdot x + b \quad (6)$$

- b. Logistic Regression : Logistic Regression merupakan algoritma *supervised learning* yang digunakan untuk klasifikasi dengan memodelkan probabilitas suatu kelas melalui fungsi sigmoid, sehingga nilai prediksi berada pada rentang 0 hingga 1 dan dapat digunakan untuk menentukan kelas suatu data (Zaidi & Al Luhayb, 2023; Syahrani et al., 2023). Fungsi tersebut memetakan nilai linear ke dalam rentang interval 0 hingga 1 melalui persamaan matematika, dapat dilihat pada persamaan (7).

$$P(y = 1|x) = \frac{1}{1 + e^{-(wx+b)}} \quad (7)$$

Di mana  $x$  adalah vektor fitur,  $w$  merupakan bobot,  $b$  adalah bias, dan  $e$  ialah bilangan eksponensial. Data diklasifikasikan ke kelas positif jika probabilitasnya  $> 0.5$ , dan kelas negatif jika di bawahnya.

- c. Naïve Bayes : Naive Bayes menggunakan Teorema Bayes untuk menghitung probabilitas kelas dengan mengasumsikan setiap fitur saling independen (Ou et al., 2022). Model matematisnya dinyatakan melalui persamaan (8).

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)} \quad (8)$$

Di mana  $P(C|X)$  adalah probabilitas posterior,  $P(X|C)$  melambangkan *likelihood*,  $P(C)$  merupakan probabilitas *prior*, dan  $P(X)$  ialah probabilitas prediktor. Keputusan klasifikasi final diambil berdasarkan nilai probabilitas tertinggi.

### Tahapan Evaluasi Model

Evaluasi kinerja model dilakukan menggunakan *confusion matrix* dan metrik evaluasi. Evaluasi ini bertujuan untuk mengukur performa model dalam mengklasifikasikan sentimen secara akurat dan konsisten. Berikut rumus yang digunakan.

- a. *Accuracy* : total prediksi yang bernilai benar dihitung berdasarkan formulasi pada Persamaan (9).

$$Accuracy = \frac{\text{Jumlah prediksi benar}}{\text{Total data}} \quad (9)$$

- b. *Precision* : kebenaran dari seluruh prediksi positif yang dihasilkan diukur melalui Persamaan (10)

$$Precision = \frac{TP}{TP + FP} \quad (10)$$

- c. *Recall* : digunakan untuk mengukur proporsi data berlabel positif yang berhasil diprediksi secara akurat oleh model melalui formulasi pada Persamaan (11).

$$Recall = \frac{TP}{TP + FN} \quad (11)$$

- d. *F1-Score* : sebagai parameter evaluasi yang menyeimbangkan nilai *precision* dan *recall* secara harmonis melalui perhitungan rata-rata harmonis pada Persamaan (10).

$$F1 = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (12)$$

- e. *Matthews Correlation Coefficient* (MCC) untuk multi kelas dapat dilihat pada persamaan (13)

$$MCC = \frac{c \times s - \sum p_k t_k}{\sqrt{(s^2 - \sum p_k^2)(s^2 - \sum t_k^2)}} \quad (13)$$

Pada Persamaan (13), ccc menyatakan jumlah prediksi benar, sss total sampel, pkp\_kpk jumlah prediksi pada kelas ke-kkk, dan tkt\_ktk jumlah data aktual pada kelas ke-kkk. Semakin mendekati 1, nilai MCC menunjukkan performa model yang semakin baik

## HASIL DAN PEMBAHASAN

### Pengumpulan Data

Dataset yang berhasil dikumpulkan sebanyak 6210 tersusun dalam bentuk Tabel dengan beberapa atribut yaitu *conversation\_id\_str*, *created\_at*, *favorite\_count*, *full\_text*, *id\_str*, *image\_url*, *in\_reply\_to\_screen*, *lang*, *location*, *quote\_count*, *reply\_count*, *retweet\_count*, *tweet\_url*, *user\_id\_str*, *username*. Adapun atribut yang digunakan sebagai data analisis adalah *full\_text*. Data-data yang berhasil didapatkan kemudian disimpan dalam format .csv. berikut ini ditampilkan beberapa atribut mewakili keseluruhan attribut hasil scraping data dapat dilihat pada tabel 1.

Tabel 1. Hasil *Scraping*

| <i>created_at</i>              | <i>full_text</i>                                                                    | <i>retweet_count</i> |
|--------------------------------|-------------------------------------------------------------------------------------|----------------------|
| Sat Apr 05 08:15:36 +0000 2025 | @NoermeeBelonta Semangat warga papua dukung Makan Bergizi Gratis bikin aura positif | 0                    |
| Sat Apr 05 01:01:13 +0000 2025 | @GenZNusantara Salam sehat dari Makan Bergizi Gratis siap raih mimpi masa depan     | 0                    |

### Preprocessing Data

Pada *preprocessing data* diperoleh beberapa hasil melalui tahapan-tahapan diantaranya yaitu : Deduplikasi Data untuk menghilangkan data yang memiliki isi identik. Proses deduplikasi diperlukan agar tidak terjadi pengulangan informasi yang dapat mempengaruhi akurasi analisis sentimen pada tahap berikutnya. Dari total 6210 data awal, proses deduplikasi menghasilkan 4141 data unik, sehingga sebanyak 2069 data terdeteksi sebagai duplikat dan dihapus.

*Cleaning text* dihasilkan data bersih tanpa memiliki URL, mention, hashtag, simbol, dan angka, dapat dilihat pada tabel 2.

Tabel 2. Hasil *Cleaning Text*

| <i>full_text</i>                                                                | <i>Cleaning</i>                                                  |
|---------------------------------------------------------------------------------|------------------------------------------------------------------|
| @GenZNusantara Salam sehat dari Makan Bergizi Gratis siap raih mimpi masa depan | Salam sehat dari Makan Bergizi Gratis siap raih mimpi masa depan |

*Case folding* menghasilkan keseragaman teks menjadi huruf kecil seperti kata "Makan", "MAKAN", dan "makan" berubah menjadi bentuk yang sama, yaitu "makan", dapat dilihat pada tabel 3.

Tabel 3. Hasil *Case Folding*

| <i>full_text</i>                                                                       | <i>Case Folding</i>                                              |
|----------------------------------------------------------------------------------------|------------------------------------------------------------------|
| @GenZNusantara <i>Salam</i> sehat dari Makan Bergizi Gratis siap raih mimpi masa depan | salam sehat dari makan bergizi gratis siap raih mimpi masa depan |

Tokenisasi, dihasilkan token berupa satuan kata yang diperoleh melalui proses tokenisasi, dapat dilihat pada tabel 4

Tabel 4. Hasil *Tokenisasi*

| <i>full_text</i>                                                                       | <i>Tokenisasi</i>                                                                    |
|----------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|
| @GenZNusantara <i>Salam</i> sehat dari Makan Bergizi Gratis siap raih mimpi masa depan | salam   sehat   dari   makan   bergizi   gratis   siap   raih   mimpi   masa   depan |

Kemudian *stopword removal* menghasilkan kata-kata bermakna yang mencerminkan konteks sebenarnya dari dokumen. Kata-kata yang termasuk dalam daftar *stopword* dibuang, sehingga jumlah kata pada sebagian besar dokumen berkurang cukup signifikan. Hasil pengurangan tersebut membantu memfokuskan analisis, dapat dilihat pada tabel 5.

Tabel 5. Hasil *Stopword Removal*

| <i>full_text</i>                                                                       | <i>Stopword Removal</i>                     |
|----------------------------------------------------------------------------------------|---------------------------------------------|
| @GenZNusantara <i>Salam</i> sehat dari Makan Bergizi Gratis siap raih mimpi masa depan | salam sehat makan bergizi gratis raih mimpi |

*Stemming*, menghasilkan kata-kata dasar yang tidak lagi memiliki awalan maupun akhiran, sehingga kata seperti "makanan", "memakan", dan "pemakan" berubah menjadi bentuk dasar "makan". Hasil stemming membuat data teks menjadi lebih seragam karena berbagai variasi bentuk kata disatukan dalam satu *representasi* yang sama, dapat dilihat pada tabel 6.

Tabel 6. Hasil *Stemming*

| <i>full_text</i>                                                                       | <i>Stemming</i>                          |
|----------------------------------------------------------------------------------------|------------------------------------------|
| @GenZNusantara <i>Salam</i> sehat dari Makan Bergizi Gratis siap raih mimpi masa depan | salam sehat makan gizi gratis raih mimpi |

### Pelabelan Data

Pelabelan data dilakukan secara manual dan dihasilkan dataset yang terdiri dari tiga kategori sentimen yaitu negatif, netral dan positif. Berdasarkan proses pelabelan manual terhadap 4.141 data tweet, diperoleh distribusi sentimen yang terdiri dari 3.427 data positif (82,76%), 403 data negatif (9,73%), dan 311 data netral (7,51%). Distribusi ini menunjukkan bahwa opini masyarakat terhadap program Makan Bergizi Gratis cenderung terpolarisasi antara dukungan, kritik dan data berbentuk informasi. Dari pelabelan ini diketahui bahwa data termasuk data *imbalance* (data tidak seimbang).

### Ekstraksi Fitur

Pada tahapan ekstraksi fitur menggunakan TF-IDF dimana kata-kata diubah menjadi numerik diperoleh nilai sparse tinggi, hal ini disebabkan normalisasi vektor (L2 *Normalization*). Pada TF-IDF dapat memperlihatkan bagaimana variasi bobot terbentuk pada dokumen yang berbeda. Meskipun beberapa kata muncul pada banyak tweet, seperti kata makan, gizi, gratis, bobotnya tidak selalu sama karena TF-IDF mempertimbangkan tingkat frekuensi lokal, dapat dilihat pada tabel 7.



Tabel 7. Hasil Ekstraksi Fitur TF-IDF

| Text                                                                   | Hasil Ekstraksi Fitur (TF-IDF)                                                                                                                                                 |
|------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| semangat warga papua<br>dukung makan gizi gratis bikin<br>aura positif | aura=0.725761, bikin=0.25538, dukung=0.182826,<br>gizi=0.0012709, gratis=0.0012709, makan=0.00099216,<br>papua=0.204825, positif=0.326816, semangat=0.319263,<br>warga=0.35204 |
| salam sehat makan gizi gratis<br>raih mimpi                            | gizi=0.00142834, gratis=0.00142834, makan=0.00111507,<br>mimpi=0.564482, raih=0.600426, salam=0.522877,<br>sehat=0.217814                                                      |

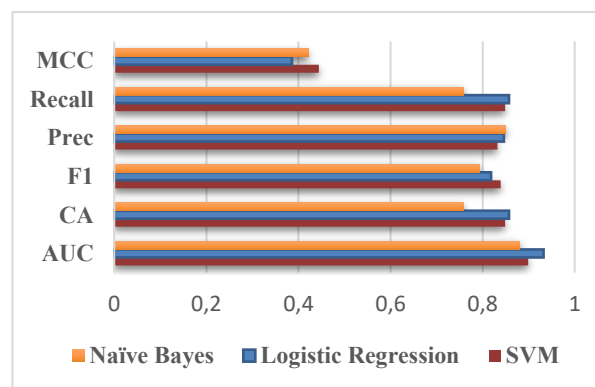
### Evaluasi Model dan Perbandingan

Perbandingan SVM terhadap model klasifikasi Logistic Regression dan Naïve Bayes menggunakan rasio perbandingan 80:20, hasil metrik evaluasi rata-rata kelas disajikan pada tabel 8.

Tabel 8. Metrik Evaluasi Rata-Rata Kelas Model

| Model               | AUC   | CA    | F1    | Prec  | Recall | MCC   |
|---------------------|-------|-------|-------|-------|--------|-------|
| SVM                 | 0,899 | 0,849 | 0,839 | 0,832 | 0,849  | 0,444 |
| Logistic Regression | 0,932 | 0,857 | 0,818 | 0,846 | 0,857  | 0,386 |
| Naïve Bayes         | 0,881 | 0,759 | 0,793 | 0,850 | 0,759  | 0,421 |

Hasil evaluasi perbandingan model tersebut juga disajikan pada gambar 1.



Gambar 1. Grafik Perbandingan Kinerja Model Klasifikasi

Berdasarkan Tabel 8, model SVM menunjukkan kinerja klasifikasi yang baik dengan nilai AUC sebesar 0,899, *Accuracy* sebesar 0,849, *F1-Score* sebesar 0,839, dan MCC sebesar 0,444. Nilai *F1-Score* dan MCC yang tertinggi dibandingkan model lainnya menunjukkan bahwa SVM memiliki kemampuan yang lebih baik dalam menjaga keseimbangan klasifikasi antar kelas sentimen. Sementara itu, Logistic Regression memperoleh nilai AUC (0,932) dan *Accuracy* (0,857) tertinggi, yang menunjukkan kemampuan diskriminasi dan ketepatan klasifikasi yang sangat baik. Adapun Naïve Bayes menghasilkan nilai AUC sebesar 0,881 dan *Accuracy* sebesar 0,759, sehingga memiliki performa yang lebih rendah dibandingkan kedua model lainnya. Evaluasi hasil juga dapat dilihat pada confusion matrix pada gambar 2.

| Confusin Matrix SVM |          |        |         |      | Confusin Matrix Logistic Regression |          |         |         |      | Confusin Matrix Naïve Bayes |          |        |         |      |     |      |      |
|---------------------|----------|--------|---------|------|-------------------------------------|----------|---------|---------|------|-----------------------------|----------|--------|---------|------|-----|------|------|
| Aktual              | Prediksi |        |         | Σ    | Aktual                              | Prediksi |         |         | Σ    | Aktual                      | Prediksi |        |         | Σ    |     |      |      |
|                     | Negatif  | Netral | Positif |      |                                     | Negatif  | Netral  | Positif |      |                             | Negatif  | Netral | Positif |      |     |      |      |
|                     | Negatif  | 154    | 24      | 144  |                                     | 322      | Negatif | 252     | 53   |                             | 17       | 322    | Negatif | 185  | 72  | 65   | 322  |
|                     | Netral   | 35     | 67      | 147  |                                     | 249      | Netral  | 77      | 120  |                             | 52       | 249    | Netral  | 59   | 139 | 51   | 249  |
|                     | Positif  | 68     | 81      | 2593 |                                     | 2742     | Positif | 156     | 207  |                             | 2379     | 2742   | Positif | 141  | 411 | 2190 | 2742 |
| Σ                   | 257      | 172    | 2884    | 3313 | Σ                                   | 485      | 380     | 2448    | 3313 | Σ                           | 385      | 622    | 2306    | 3313 |     |      |      |

Gambar 2. Confusion Matrix SVM, Logistic Regression, Naïve Bayes

Berdasarkan confusion matrix pada Gambar 2, seluruh model menunjukkan kemampuan yang baik dalam mengenali sentimen positif. Namun, masih terdapat kesalahan klasifikasi pada kelas negatif dan netral, yang sebagian besar diprediksi sebagai sentimen positif. Kondisi ini menunjukkan adanya pengaruh ketidakseimbangan data, di mana kelas positif yang mendominasi dataset lebih mudah dipelajari oleh model dibandingkan kelas lainnya..

### KESIMPULAN

Penerapan TF-IDF dan Support Vector Machine (SVM) pada penelitian ini mampu mengklasifikasikan sentimen masyarakat terhadap program Makan Bergizi Gratis dengan hasil yang baik. Model SVM memperoleh nilai Accuracy sebesar 0,849, F1-Score 0,839, dan MCC 0,444. Meskipun Logistic Regression menghasilkan Accuracy dan AUC yang lebih tinggi, SVM menunjukkan kemampuan yang lebih baik dalam menjaga keseimbangan prediksi antar kelas. Ketidakseimbangan distribusi data masih menjadi tantangan dalam klasifikasi kelas minoritas. Oleh karena itu, penggunaan teknik penyeimbangan data dapat dipertimbangkan pada penelitian selanjutnya untuk meningkatkan hasil klasifikasi.

### REFERENSI

- Kiratli, O. S. (2023). Social media effects on public trust in the European Union. *Public Opinion Quarterly*, 87(3), 749-763.
- Parr, C. S. (2025). Does social media undermine trust? Institutional trust in civil society and governance institutions. *Journal of Public Policy*, 1-24.
- Joseph, T. (2024). Natural language processing (NLP) for sentiment analysis in social media. *International Journal of Computing and Engineering*, 6(2), 35-48.
- Nip, J. Y., & Berthelie, B. (2024). Social media sentiment analysis. *Encyclopedia*, 4(4), 1590-1598.
- Zhou, X., & Shen, H. (2022). Communication-efficient distributed learning for high-dimensional support vector machines. *Mathematics*, 10(7), 1029.
- PAMBUDI, F., & PRANOTO, W. (2024). OPTIMASI SVM DENGAN RFE DAN ROS UNTUK MENGATASI HIGH DIMENSION DAN IMBALANCED DATA BANJIR. *JURNAL TEKNOLOGI SISTEM INFORMASI DAN APLIKASI Учредители: Universitas Pamulang*, 7(3), 1194-1203.
- Zahra, N. Z., Farhanatussaidah, S., Afifah, N. N., Muthoharoh, L., Satria, A., & Manullang, M. C. (2026). Benchmarking Logistic Regression, SVM, Naive Bayes, and IndoBERT fine-tuning for sentiment analysis on Indonesian product reviews. *arXiv preprint arXiv:2605.03439*.
- Hazrati, N., & Rahmatulloh, A. (2026). Comparison of SVM, Naive Bayes, and Logistic Regression for LinkedIn Reviews Sentiment Analysis. *JASMINE: Journal of Intelligent Systems and Machine Learning*, 1(1), 57-67.
- Vazli, H., & Suryono, R. R. (2026). KOMPARASI MODEL SVM, NAÏVE BAYES, DAN RANDOM FOREST DALAM ANALISIS SENTIMEN TERHADAP PERCAKAPAN PUBLIK TENTANG PERTAMINA. *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, 11(1), 92-105.
- Hidayati, N., Hamami, F., & Fa'rifah, R. Y. (2023). Aspect-Based Sentiment Analysis On FLIP Application Reviews (Play Store) Using Support Vector Machine (SVM) Algorithm. *Journal of Informatics and Telecommunication Engineering*, 7(1), 183-197.
- Yanuari, B., Utami, E., & Parikesit, A. A. (2024). Data Clustering for Sentiment Classification with Naïve Bayes and Support Vector Machine. *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, 8(6), 819-827.
- Yunita, R., & Kamayani, M. (2023). Perbandingan algoritma SVM dan Naïve Bayes pada analisis sentimen penghapusan kewajiban skripsi. *The Indonesian Journal of Computer Science*, 12(5).
- Tala, F. (2003). A study of stemming effects on information retrieval in Bahasa Indonesia.
- Adriani, M., Asian, J., Nazief, B., Tahaghoghi, S. M., & Williams, H. E. (2007). Stemming Indonesian: A confix-stripping approach. *ACM Transactions on Asian Language Information Processing (TALIP)*, 6(4), 1-33.



- Liu, H., Chen, X., & Liu, X. (2022). A study of the application of weight distributing method combining sentiment dictionary and TF-IDF for text sentiment analysis. *Ieee Access*, 10, 32280-32289.
- Alemerien, K., Al-Ghareeb, A., & Alksasbeh, M. Z. (2024). Sentiment analysis of online reviews: a machine learning based approach with TF-IDF vectorization. *Journal of Mobile Multimedia*, 20(5), 1089-1116. <https://doi.org/10.13052/jmm1550-4646.2055>.
- risnawati, W., & Wibowo, A. (2024). *Sentiment analysis of ICT service user using Naive Bayes classifier and SVM methods with TF-IDF text weighting*. Jurnal Teknik Informatika (JUTIF), 5(3), 709–719. <https://doi.org/10.52436/1.jutif.2024.5.3.1784>
- Zaidi, A., & Al Luhayb, A. S. M. (2023). Two statistical approaches to justify the use of the logistic function in binary logistic regression. *Mathematical Problems in Engineering*, 2023(1), 5525675. <https://doi.org/10.1155/2023/5525675>
- Syahrani, R., Suhartono, S., & Zaman, S. (2023). Regresi Logistik Multinomial untuk Prediksi Kategori Kelulusan Mahasiswa. *JISKA (Jurnal Informatika Sunan Kalijaga)*, 8(2), 102–111. <https://doi.org/10.14421/jiska.2023.8.2.102-111>
- Ou, G., He, Y., Fournier-Viger, P., & Huang, J. Z. (2022). A novel mixed-attribute fusion-based naive bayesian classifier. *Applied Sciences*, 12(20), 10443.