

ANALISIS SENTIMEN PROGRAM MAKAN BERGIZI GRATIS PADA MEDIA SOSIAL X DENGAN ALGORITMA *SUPPORT VECTOR MACHINE*

Gali Hirzi Farizi^{1*}, Ade Andri Hendriadi², Azhari Ali Ridha³

^{1,2,3}Universitas Singaperbangsa Karawang; Jl. HS.Ronggo Waluyo, Puseurjaya, Telukjambe Timur Karawang, Indonesia; (0267) 641177

Keywords:

Analisis Sentimen, Support Vector Machine, TF-IDF, MBG, Media Sosial X.

Correspondent Email:

galihhirzi@gmail.com

Abstrak. Program Makan Bergizi Gratis (MBG) merupakan salah satu kebijakan pemerintah yang bertujuan meningkatkan kualitas gizi masyarakat dan menurunkan angka stunting di Indonesia. Tanggapan masyarakat terhadap program selama pelaksanaannya banyak disampaikan melalui media sosial X, termasuk dukungan, kritik, dan informasi tentang bagaimana program dijalankan. Tujuan dari penelitian ini adalah untuk menganalisis persepsi masyarakat terhadap Program MBG dan juga mengevaluasi bagaimana algoritma Support Vector Machine (SVM) bekerja dalam klasifikasi persepsi. Data selection, preprocessing, transformation, data mining, evaluation, dan knowledge presentation adalah semua bagian dari metode Knowledge Discovery in Databases (KDD). Data diperoleh melalui proses crawling tweet berbahasa Indonesia dari Januari hingga November 2025 menggunakan kata kunci terkait MBG. Tahapan preprocessing termasuk case folding, cleaning, tokenize, normalize, stopword, dan stemming. Selanjutnya, TF-IDF digunakan untuk pembobotan kata. Untuk melakukan klasifikasi, algoritma SVM digunakan pada beberapa skenario pengujian SVM Normal dan Balanced. Skenario A, yang menggunakan SVM Normal dengan pembagian data 90:10, adalah hasil terbaik dengan nilai akurasi 72%, presisi 73%, recall 72%, dan f1-score 72%. Hasil penelitian menunjukkan bahwa SVM Normal mampu melakukan klasifikasi sentimen dengan cukup baik dan stabil. Selain itu, tanggapan masyarakat terhadap program MBG cenderung negatif.



Copyright © 2026 The Author(s), Published by Jurnal Informatika dan Teknik Elektro Terapan. This is an open-access article distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0).

Abstract. *The Free Nutritious Meals Program (MBG) is a government policy aimed at improving the nutritional quality of the population and reducing stunting rates in Indonesia. Public responses to the program during its implementation have been widely shared on the social media platform X, including expressions of support, criticism, and information about how the program is being implemented. The objective of this study is to analyze public perceptions of the MBG Program and to evaluate how the Support Vector Machine (SVM) algorithm performs in classifying these perceptions. Data selection, preprocessing, transformation, data mining, evaluation, and knowledge presentation are all components of the Knowledge Discovery in Databases (KDD) methodology. Data was obtained through the process of crawling Indonesian-language tweets from January to November 2025 using keywords related to MBG. Preprocessing steps included case folding, cleaning, tokenization, normalization, stopword removal, and stemming. Subsequently, TF-IDF was used for word weighting. To perform classification, the SVM algorithm was applied across several testing scenarios: Normal SVM and Balanced SVM. Scenario A, which used Normal SVM with a 90:10 data split, yielded the best results with an accuracy of 72%, precision of 73%, recall of 72%, and an F1-score of 72%. The research results*

indicate that Normal SVM is capable of performing sentiment classification quite well and consistently. Furthermore, public sentiment toward the MBG program tends to be negative.

1. PENDAHULUAN

Program Makan Bergizi Gratis (MBG) adalah kebijakan nasional yang bertujuan untuk meningkatkan kualitas gizi masyarakat dan mengurangi stunting dan malnutrisi pada kelompok rentan seperti balita, anak sekolah, dan ibu hamil [1]. Program ini telah menjangkau 41.297.872 penerima manfaat hingga November 2025 melalui 14.483 Satuan Pelayanan Pemenuhan Gizi (SPPG) yang terletak di 38 provinsi. Untuk memastikan keberlanjutan dan efektivitas program, evaluasi implementasi harus dilakukan. Media sosial X menjadi sarana masyarakat untuk berbagi kritik, dukungan, dan pengalaman terkait kebijakan pemerintah secara real-time saat teknologi digital berkembang. Platform ini memiliki 23,76 juta pengguna di Indonesia dan menyediakan sumber data yang relevan untuk mengamati persepsi publik terhadap Program MBG [2]. Kualitas makanan, distribusi program, transparansi anggaran, dan kasus keracunan makanan memengaruhi persepsi masyarakat terhadap program [3].

Analisis sentimen dari data pengguna media sosial diperlukan untuk mendapatkan pemahaman yang objektif tentang kecenderungan opini publik. Dalam analisis sentimen, algoritma Support Vector Machine (SVM) banyak digunakan karena dapat generalisasi dengan baik, tahan terhadap noise, dan dapat menghasilkan klasifikasi yang akurat pada data teks berukuran besar. Tahapan preprocessing seperti case folding, cleaning, tokenize, normalize, stopword, dan stemming mendukung kinerja SVM [4]. SVM telah terbukti efektif dalam klasifikasi sentimen pada masalah kebijakan dan layanan publik. Namun, penelitian sebelumnya lebih banyak berfokus pada aspek kebijakan, tata kelola, fiskal, dan implementasi program [3].

Oleh karena itu, tujuan penelitian ini adalah untuk melihat bagaimana masyarakat melihat Program MBG. Ini dilakukan dengan melihat unggahan di media sosial X dan melihat kemampuan SVM untuk mengklasifikasikan sentimen publik. Hasil penelitian diharapkan dapat meningkatkan penelitian tentang analisis

sentimen dalam bahasa Indonesia dan juga menjadi bahan pertimbangan bagi pemerintah saat mereka membuat kebijakan berbasis data.

2. TINJAUAN PUSTAKA

2.1 Program Makan Bergizi Gratis (MBG)

Dalam upaya mencegah kekurangan gizi dan stunting, Program Makan Bergizi Gratis (MBG) didasarkan pada konsep pemenuhan gizi seimbang. Konsep ini dijelaskan dalam Buku Pemantauan Status Gizi mengenai pentingnya kecukupan zat gizi bagi pertumbuhan anak [5]. Program MBG juga bertujuan untuk membantu kelompok rentan mendapatkan lebih banyak akses ke pangan bergizi, memenuhi kebutuhan gizi harian mereka, dan mendukung optimalisasi pertumbuhan selama fase 1.000 Hari Pertama Kehidupan (HPK). Kelompok yang paling rentan terhadap masalah gizi adalah anak sekolah, balita, dan ibu hamil, yang merupakan sasaran program. Karena prevalensi balita pendek yang terus meningkat di seluruh negeri, buku tersebut menyebutkan stunting sebagai ancaman jangka panjang bagi kualitas SDM [6].

2.2 Media Sosial X

Media sosial X merupakan platform digital yang memungkinkan pengguna untuk menyampaikan opini, berbagi informasi, dan merespons isu publik secara cepat dan terbuka. Dalam konteks penelitian sosial, X sering digunakan sebagai sumber data untuk text mining dan sentiment analysis, karena karakternya yang berbasis teks singkat, publik, dan bersifat real-time.

Media sosial X efektif sebagai sumber data untuk menganalisis persepsi masyarakat terhadap kebijakan pemerintah. Di Indonesia, X juga menjadi wadah utama pembentukan opini publik terhadap kebijakan sosial, seperti MBG [7]. Platform ini memungkinkan pemerintah dan akademisi memantau dinamika persepsi publik yang mencerminkan tingkat kepercayaan, kritik, dan dukungan terhadap program pemerintah.

2.3 Data Mining

Data mining dijelaskan sebagai proses analitis yang digunakan untuk menemukan pola, hubungan, dan informasi bermakna dari kumpulan data berukuran besar, termasuk dalam konteks website structure, content, dan usage mining yang menjadi cabang utama eksplorasi data pada sistem informasi [8]. Tujuan utama dari data mining adalah menghasilkan pengetahuan yang dapat digunakan untuk memprediksi kejadian atau perilaku tertentu, seperti yang ditunjukkan pada penerapan model prediksi berbasis teks dan query informasi dalam penelitian terkait web mining.

2.4 Text Mining

Text mining adalah cabang dari data mining yang berfokus pada proses mengekstraksi pola dan pengetahuan dari data berbasis teks sehingga informasi yang tersembunyi dapat diolah menjadi bentuk yang terstruktur [9].

2.5 Data Crawling

Data crawling adalah proses pengambilan data otomatis dari media sosial dengan menggunakan algoritma atau skrip tertentu untuk mengumpulkan teks, opini, dan komentar publik. Teknik ini, dalam analisis sentimen, berfungsi untuk mengumpulkan data tweet berdasarkan kata kunci atau subjek yang dibicarakan oleh pengguna. Penerapan data crawling pada media sosial X efektif dalam memperoleh dataset publik yang relevan dengan isu sosial atau politik di Indonesia [10].

2.6 KDD (Knowledge Discovery in Databases)

Knowledge Discovery in Databases (KDD) adalah proses komputasi yang menggunakan algoritma matematika untuk mengekstraksi data dan mempertimbangkan tindakan yang mungkin dilakukan di masa depan. Hasil KDD adalah pengetahuan yang sebelumnya tidak diketahui, potensial, dan bermanfaat [11].

2.7 Algoritma Support Vector Machine

Algoritma Support Vector Machine (SVM) adalah metode machine learning yang bekerja dengan membangun sebuah hyperplane optimal untuk memisahkan data ke dalam kelas tertentu. Tujuan utama SVM adalah untuk memaksimalkan jarak antar kelas sehingga model dapat melakukan generalisasi dengan

baik [12]. Dalam proses klasifikasi, SVM bergantung pada vector pendukung, yaitu titik data yang paling dekat dengan hyperplane, dan mereka bertanggung jawab untuk menentukan batas keputusan secara konsisten.

2.8 Imbalanced Data

Ketidakseimbangan data, juga dikenal sebagai Imbalanced Data, adalah salah satu cara untuk mengatasi ketidakseimbangan dalam kelas data yang lebih kecil. Jika terjadi pada kelas dengan data lebih sedikit dari kelas lainnya, teknik oversampling dapat digunakan, yang memberikan hasil yang lebih baik daripada teknik undersampling.

2.9 Class Weight

Digunakan untuk mengatasi ketidakseimbangan data, kelas berat memberikan prioritas lebih besar pada kelas minoritas daripada kelas mayoritas. Parameter ini mempengaruhi fungsi kerugian dalam model klasifikasi seperti SVM dengan memberikan penalti lebih tinggi untuk kesalahan pada kelas yang jumlahnya sedikit. Tujuan penggunaan model adalah untuk menurunkan biasnya terhadap kelas dominan dan meningkatkan kemampuan model untuk mengidentifikasi kelas secara lebih seimbang.

2.10 Python

Python adalah bahasa pemrograman interpreter yang sangat fleksibel yang mendukung berbagai paradigma pemrograman, termasuk fungsional, prosedural, dan berorientasi objek, menurut dokumentasi resminya [13].

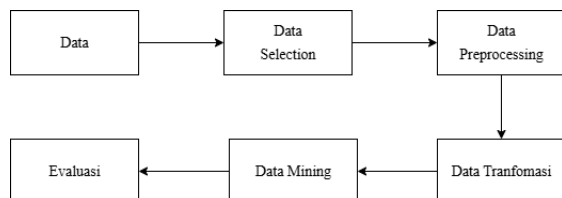
2.11 Google Colaboratory

Menurut dokumentasi resmi Google, Google Colaboratory (Colab) adalah platform komputasi berbasis cloud yang memungkinkan pengguna menjalankan kode Python tanpa memerlukan instalasi local [14].

3. METODE PENELITIAN

Knowledge Discovery in Databases (KDD) digunakan sebagai metode dalam penelitian ini karena kemampuannya untuk mengubah data tidak terstruktur mentah menjadi pengetahuan penting melalui proses sistematis seperti data selection, preprocessing, transformation, data mining, evaluation, dan knowledge

presentation. Selanjutnya, pada tahap klasifikasi sentimen, algoritma Support Vector Machine (SVM) dipilih karena kemampuannya membedakan sentimen kelas positif, negatif, dan netral berdasarkan fitur teks.



Gambar 1. Alur Metode KDD

3.1 Data Selection

Pada tahap pertama, yaitu data selection, proses diawali dengan pengumpulan data untuk mendapatkan data yang berkaitan dengan tujuan penelitian. Data diperoleh melalui crawling media sosial X dengan memanfaatkan tools Python melalui platform Google Colab.

3.2 Data Preprocessing

Proses pembersihan data mentah agar menjadi lebih terstruktur dan siap untuk digunakan dalam pemodelan. Jika data teks media sosial tidak dibersihkan, algoritma mungkin gagal menangani bahasa tidak baku, simbol, tagar, emoji, dan singkatan.

3.3 Transformation

Transformasi data adalah tahapan penting dalam proses analisis data, di mana data disiapkan untuk analisis berikutnya. Tujuan utama transformasi data adalah mengubah format dan representasi data, dan beberapa teknik yang digunakan di tahap ini termasuk transformasi diskritisasi, transformasi skala, dan penciptaan fitur dan atribut baru.

3.4 Data Mining

Proses Knowledge Discovery in Database (KDD) mencakup tahap data mining, di mana algoritma pembelajaran mesin digunakan untuk menemukan pola dan membuat model prediksi sentimen.

3.5 Evaluation

Tujuan dari tahap evaluasi adalah untuk mengevaluasi kemampuan model SVM untuk mengklasifikasikan data dengan benar. Confusion Matrix digunakan untuk melakukan evaluasi, yang membandingkan hasil prediksi

dengan label sebenarnya untuk memperoleh empat ukuran kinerja utama.

3.6 Knowledge Presentation

Tahapan knowledge presentation merupakan fase terakhir dari proses metode Knowledge Discovery in Databases (KDD) yang berfungsi untuk menyajikan hasil temuan secara informatif, mudah dipahami, dan dapat digunakan saat membuat keputusan.

4. HASIL DAN PEMBAHASAN

4.1 Data Selection

Tahapan Pertama crawling, Pada tahap ini, tool python Google Colab digunakan untuk mengumpulkan data dengan kata kunci seperti "Masalah MBG", "MBG", "Distribusi MBG", "Program MBG", dan "Keracunan MBG". Tahap ini menghasilkan data berupa berbagai kolom yang berisi teks komentar dan elemen lainnya yang relevan dengan penelitian. Data dari Januari 2025 - November 2025.

Pada langkah selanjutnya, yaitu identifikasi data, sebanyak 5000 data diperoleh dari hasil proses crawling pada media sosial X menggunakan tools Python melalui Google Colab.

conversation_id_str	created_at	favorite_count	full_text	id_str	image_url	in_reply_to_screen_name	lang
1,99E+18	Mon No 0		@law209 #####	https://pblaw20956		in	
1,98E+18	Mon No 0		@PartaiS #####	PartaiSocmed		in	
1,99E+18	Mon No 0		Babinsa (#####	https://pbs.twimg.com/		in	
1,99E+18	Mon No 0		Pengawa #####	https://pbs.twimg.com/		in	
1,99E+18	Mon No 0		@wibisoi #####	wibisonoys		in	
1,99E+18	Mon No 0		Bhabinka #####	https://pbs.twimg.com/		in	
1,99E+18	Mon No 0		@wibisoi #####	wibisonoys		in	
1,99E+18	Mon No 0		@wibisoi #####	wibisonoys		in	
1,99E+18	Mon No 1		Komitme #####	https://pbs.twimg.com/		in	
1,99E+18	Mon No 0		Di bawah #####	https://pbs.twimg.com/		in	

Gambar 2. Hasil Output Crawling

4.2 Data Preprocessing

Sebelum proses klasifikasi, tahap preprocessing dilakukan untuk membersihkan dan menyiapkan 4.639 data tweet. Sebagai bagian dari proses ini, case folding, cleaning, tokenize, normalize, stopword removal, dan stemming dilakukan. Tujuan dari proses ini adalah untuk menyeragamkan teks, menghilangkan elemen yang tidak relevan, dan menghasilkan kata dasar yang siap untuk analisis sentimen.

Tabel 1. Hasil Data Preprocessing

Preprocessing	Content
Data Mentah	Mari dukng dan sukseskan MBG https://t.co/emwzwh7fkj
Case Folding	mari dukng dan sukseskan mbg https://t.co/emwzwh7fkj
Cleaning	mari dukng dan sukseskan mbg
Tokenize	['mari', 'dukng', 'dan', 'sukseskan', 'mbg']
Normalize	['mari', 'dukung', 'dan', 'sukseskan', 'mbg']
Stopword	['mari', 'dukung', 'sukseskan', 'mbg']
Stemming	['mari', 'dukung', 'sukses', 'mbg']

4.3 Data Transformation

Pada tahap transformasi, data preprocessing diubah menjadi representasi numerik menggunakan metode Term Frequency–Inverse Document Frequency (TF-IDF). Metode ini memberikan bobot pada setiap kata berdasarkan frekuensi kemunculannya dalam suatu dokumen dan tingkat kepentingannya terhadap seluruh dokumen dalam dataset. Matriks bobot TF-IDF, hasil transformasi, menggambarkan setiap tweet sebagai vektor numerik.

Kata Rata-rata TF-IDF		
0	racun	0.038736
1	program	0.038225
2	makan	0.034401
3	distribusi	0.025855
4	anak	0.022623

Gambar 3. Output TF-IDF

4.4 Data Mining

Pada tahap ini, data yang telah diproses sebelumnya dimasukkan ke tahap data mining. Pada tahap ini, data ini akan dianalisis lebih lanjut melalui proses pemodelan, yang menggunakan algoritma yang sesuai, untuk

menemukan pola dan informasi yang terkandung dalam dataset. Dalam penelitian ini, algoritma yang digunakan adalah Support Vector Machine (SVM), yang melakukan proses klasifikasi sentimen pada data teks yang telah diproses sebelumnya.

Tabel 2. Skenario Data

Skenario	Data Uji	Keterangan
Skenario A 90:10	464	Normal
Skenario B 80:20	928	Normal
Skenario C 70:30	1392	Normal
Skenario D 60:40	1856	Normal
Skenario E 90:10	464	Balanced
Skenario F 80:20	928	Balanced
Skenario G 70:30	1392	Balanced
Skenario H 60:40	1856	Balanced

Setiap pasangan model memiliki pembagian skenario data yang sama, yaitu model A dengan model E, model B dengan model F, model C dengan model G, dan model D dengan model H. Model A hingga D menunjukkan model tanpa menggunakan teknik Class Weigh, sedangkan model E hingga H menunjukkan model dengan menggunakan teknik Class Weigh.

4.5 Evaluasi

Evaluasi yang dilakukan menggunakan Matrix Confusion membandingkan hasil prediksi dengan label sebenarnya untuk menghasilkan empat ukuran kinerja utama. Sebagai langkah terakhir, hasil evaluasi kemudian dianalisis lebih lanjut untuk menemukan pola dan informasi yang membantu pengambilan keputusan.

Tabel 3. Hasil Evaluasi

Skenario	Akurasi	Presisi	Recall	F1-Score
A 90:10 Normal	0,72	0,73	0,72	0,72
B 80:20 Normal	0,71	0,72	0,71	0,71
C 70:30 Normal	0,72	0,72	0,72	0,71
D 60:40 Normal	0,72	0,72	0,72	0,71
E 90:10 Balanced	0,70	0,73	0,70	0,71

F 80:20 Balanced	0,70	0,71	0,70	0,70
G 70:30 Balanced	0,70	0,71	0,70	0,71
H 60:40 Balanced	0,71	0,71	0,71	0,71

Skenario A dengan metode SVM Normal dan pembagian data 90:10 dianggap memiliki kinerja terbaik dalam klasifikasi sentimen. Dibandingkan dengan situasi lain, model ini menunjukkan kemampuan klasifikasi yang unggul dengan nilai akurasi 0,72, presisi 0,73, recall 0,72, dan f1-score 0,72. Senario A cenderung memiliki performa yang lebih baik terutama pada nilai presisi dan f1-score dibandingkan dengan skenario B, C, dan D yang juga menggunakan SVM Normal. Namun, skenario C dan D memiliki nilai akurasi yang sama.

4.6 Knowledge Presentation

Pada tahapan ini menyajikan visualisasi data berupa word frequency dan wordcloud dari hasil analisis sentimen terhadap data teks yang telah diklasifikasikan ke dalam tiga label, yaitu positif, netral dan negatif.



Gambar 4. Hasil Wordcloud Positif

Gambar menunjukkan visualisasi wordcloud label positif yang menunjukkan kata-kata yang paling sering muncul melalui ukuran huruf yang proporsional. Kata-kata seperti "mbg", "program", "makan", dan "bergizi" tampaknya muncul dengan jumlah yang paling besar pada visualisasi tersebut, menunjukkan bahwa kata-kata tersebut mendominasi perasaan positif masyarakat.



Gambar 5. Hasil Wordcloud Negatif

Gambar wordcloud untuk label negatif yang menunjukkan distribusi kata serupa dengan kata-kata dominan seperti "mbg", "racun", "program". Meskipun beberapa kata tampak sama dengan label positif, konteks penggunaannya dalam kalimat kemungkinan besar yang membedakan sentimennya. Secara keseluruhan, visualisasi ini memberikan gambaran umum mengenai kata-kata yang paling banyak digunakan dalam opini negatif terhadap mbg.



Gambar 6. Hasil Wordcloud Netral

Label netral biasanya digunakan dalam konteks informatif tanpa menunjukkan dukungan atau kritik yang kuat. Berdasarkan sikap netral mereka, visualisasi ini memberikan gambaran umum tentang topik diskusi masyarakat tentang program MBG.

Hasil penelitian menunjukkan bahwa, di media sosial X, persepsi masyarakat terhadap Program Makan Bergizi Gratis (MBG) didominasi oleh sentimen negatif sebesar 39,5%, diikuti oleh sentimen positif sebesar 34,8%, dan sentimen netral sebesar 25,7%. Temuan ini menunjukkan bahwa, meskipun MBG adalah program strategis pemerintah yang bertujuan untuk meningkatkan kualitas gizi masyarakat dan mengurangi angka stunting, implementasinya masih menghasilkan berbagai reaksi kritis. Sebaliknya, banyak

komentar positif menunjukkan bahwa sebagian orang percaya bahwa program ini dapat membantu meningkatkan akses terhadap makanan bergizi dan mendukung upaya untuk meningkatkan kualitas sumber daya manusia. Hasil ini menunjukkan bahwa pandangan publik tentang MBG berubah-ubah, dengan tujuan program didukung dan kritik terhadap implementasinya.

Dalam hal metodologi, metode Knowledge Discovery in Databases (KDD) berhasil mengubah data tidak terstruktur dari media sosial menjadi informasi yang dapat dianalisis. Metode ini mencakup langkah-langkah seperti data selection, preprocessing, transformasi, data mining, evaluasi, dan knowledge presentation. Sebelum pembobotan menggunakan TF-IDF, tahapan preprocessing seperti case folding, cleaning, tokenize, normalize, stopword, dan stemming mampu mengurangi noise pada data, meningkatkan kualitas representasi teks. Hasil evaluasi menunjukkan bahwa algoritma Support Vector Machine (SVM) sangat baik dalam melakukan klasifikasi sentimen terhadap data media sosial. Skenario terbaik diperoleh untuk SVM Normal dengan pembagian data latih dan data uji sebesar 90:10, dengan skor akurasi 72%, presisi 73%, recall 72%, dan F1-Score 72%. Hasil ini menunjukkan bahwa SVM mampu mengidentifikasi pola sentimen masyarakat secara konsisten meskipun karakteristiknya berbeda.

Hasil distribusi sentimen yang diperoleh diperkuat dengan analisis word frequency dan wordcloud. Dalam sentimen positif, kata-kata yang paling sering dikaitkan dengan program, gizi, dan manfaat MBG, yang menunjukkan bahwa masyarakat mendukung tujuan program untuk meningkatkan kualitas gizi dan kesejahteraan masyarakat. Di sisi lain, dalam sentimen negatif, kata-kata yang paling sering dikaitkan dengan masalah yang berkaitan dengan pelaksanaan program, yang menunjukkan bahwa kritik masyarakat lebih fokus pada masalah pelaksanaan daripada tujuan program itu sendiri. Namun, sentimen netral tidak menunjukkan kecenderungan dukungan atau penolakan yang kuat, didominasi oleh kata-kata informatif saat berbicara tentang aktivitas penyebaran informasi dan diskusi tentang MBG. Secara keseluruhan, temuan penelitian berhasil mencapai tujuan penelitian, menemukan

persepsi publik terhadap Program MBG dan mengevaluasi bagaimana algoritma SVM berfungsi untuk klasifikasi sentimen. Analisis sentimen media sosial dapat digunakan sebagai alat evaluasi kebijakan berbasis data untuk membantu pemangku kebijakan meningkatkan pelaksanaan program dan memperkuat strategi komunikasi publik.

5. KESIMPULAN

- Analisis sentimen yang dilakukan terhadap komentar pengguna media sosial X, pendapat publik terbagi menjadi tiga kategori yaitu negatif, positif, dan netral. Berdasarkan hasil visualisasi distribusi label sentimen, sebanyak 39,5% termasuk dalam kategori negatif, 34,8% termasuk dalam kategori positif, dan 25,7% termasuk dalam kategori netral, sebanyak 4639 data digunakan dalam penelitian setelah data duplikat dihapus dari data awal 5000.
- Skenario A, dengan metode SVM Normal dan pembagian data 90:10, menunjukkan hasil terbaik dalam penelitian ini dengan nilai akurasi sebesar 72%, presisi sebesar 73%, recall sebesar 72%, dan f1-score sebesar 72%. Hasilnya menunjukkan bahwa penggunaan SVM Normal dapat melakukan klasifikasi sentimen dengan lebih baik dan stabil dibandingkan dengan penerapan Class Weight Balanced. Dengan demikian, model dapat mengidentifikasi sentimen positif, negatif, dan netral secara cukup baik pada data teks media sosial X yang terkait dengan Program MBG.

UCAPAN TERIMA KASIH

Penulis menyampaikan terima kasih kepada kedua orang tua yang selalu memberikan support dan doa restu dalam setiap langkah penelitian ini, serta kepada dosen pembimbing yang telah membantu dalam penelitian yang dilakukan.

DAFTAR PUSTAKA

- [1] Basit, M., & Ramadani, H. (2025). Analisis Implementasi Program Makan Bergizi Gratis Terhadap Perkembangan Ekonomi. *Journal of Economics Development Research*, 1(2), 49-54.
- [2] Hadi, N., & Sugiarto, D. (2025). Analisis Sentimen Pembangunan IKN pada Media

- Sosial X Menggunakan Algoritma SVM, Logistic Regression dan Naïve Bayes. *J. Inform. J. Pengemb. It*, 10(1), 37-49.
- [3] Suardi, S., & Purmadani, A. S. (2025). Makan Bergizi Gratis Dan Dampak Bagi Pertumbuhan Ekonomi. *Bisnis-Net Jurnal Ekonomi dan Bisnis*, 8(1), 323-327.
- [4] Maulana, N. A., & Darwis, D. (2025). Perbandingan Metode SVM dan Naïve Bayes untuk Analisis Sentimen pada Twitter tentang Obesitas di Kalangan Gen Z. *Jurnal Pendidikan dan Teknologi Indonesia*, 5(3), 655-666.
- [5] Taufiq, S. (2025). Indonesian Klasifikasi Status Gizi Menggunakan Algoritma Naive Bayes berdasarkan Pengukuran Antropometri. *Jurnal Ilmu Pengetahuan dan Teknologi*, 9(1), 1-5.
- [6] Komaria, A. K., Priyatna, B., Tukino, T., & Hananto, A. (2026). Segmentasi Daerah Prioritas Penanganan Stunting Di Jawa Barat Menggunakan Metode K-Means Clustering. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 10(3), 4840-4846.
- [7] Manalu, D. R., Tobing, M. C. L., & Yohanna, M. (2022). Analisis Sentimen Twitter Terhadap Wacana Penundaan Pemilu Dengan Metode Support Vector Machine. *Jurnal Manajemen Informatika & Komputerisasi Akuntansi*, 6(2), 149–156.
- [8] Jauhari, M. T., & Marleny, F. D. (2025). Penerapan Data Mining Algoritma Decision Tree Untuk Memprediksi Keterlambatan Pembayaran Piutang Usaha Pada PT XYZ. *Jurnal Informatika dan Teknik Elektro Terapan*, 13(3).
- [9] Aalst, W. Van Der, Mylopoulos, J., Ram, S., Rosemann, M., & Szyperski, C. (2022). *Process Mining Handbook*. <https://doi.org/10.1007/978-3-031-08848-3>
- [10] Dewi, A. R. P., Riyadi, S., Damarjati, C., Isa, N. A. M., & Andriyani, A. D. (2025). Sentiment Analysis Of Pro-Israel Product Boycott Action Using Indobert Method On Unbalanced Data. *Jurnal Informatika E-ISSN*, 13(2), 187–197.
- [11] Muttaqin, M. R., & Defriani, M. (2020). Algoritma K-Means Untuk Pengelompokan Topik Skripsi Mahasiswa. *ILKOM Jurnal Ilmiah*, 12(2), 121–129.
- [12] Faiza, I. M., & Andriani, W. (2022). Tinjauan Pustaka Sistematis: Penerapan Metode Machine Learning untuk Deteksi Bencana Banjir. *Jurnal Minfo Polgan*, 11(2), 59-63.
- [13] Python, S. F. (N.D.). *Python Programming Language*. Retrieved December 1, 2025, From [Python.Org](https://python.org).
- [14] Google. (N.D.). *Google Colaboratory*. Retrieved December 1, 2025, From <https://colab.research.google.com/>
- [15] Rahmawati, R., & Prihartono, W. (2025). Optimasi stok dengan clustering data transaksi penjualan menggunakan algoritma K-Means di Konter Agung Cell. *Jurnal Informatika dan Teknik Elektro Terapan*, 13(2).