

Implementasi Algoritma *Indobert* Dan *Smote* Pada Analisis Sentimen Masyarakat Terhadap Program Makan Bergizi Gratis

DOI: <http://dx.doi.org/10.35889/jutisi.v15i4.3901>

Creative Commons License 4.0 (CC BY – NC)

Syarif Hidayatullah^{1*}, Ahmad Jazuli², Estiwijayanti³

Teknik Informatika, Universitas Muria Kudus, Kudus, Indonesia.

*e-mail *Corresponding Author*: 202251017@std.umk.ac.id

Abstract

The Free Nutritious Meal (MBG) Program is one of the Indonesian government's public policies that has generated diverse public responses on social media, making sentiment analysis essential for understanding public perceptions as a basis for policy evaluation. This study aims to analyze public sentiment toward the MBG Program on the X social media platform using the IndoBERT model combined with the Synthetic Minority Over-sampling Technique (SMOTE) to address class imbalance. A total of 2,260 tweets were collected through web scraping and subsequently preprocessed through text cleaning, normalization, tokenization, sentiment labeling using the InSet Lexicon, data balancing with SMOTE, and IndoBERT fine-tuning. Model performance was evaluated using accuracy, precision, recall, and F1-score metrics. The results show that the proposed model achieved an 80% accuracy, with approximately 75% of the data classified as negative sentiment, primarily driven by concerns regarding the state budget, potential corruption, and program distribution. Meanwhile, positive sentiment mainly highlighted the program's benefits for children's nutrition and local economic empowerment. These findings demonstrate that the combination of IndoBERT and SMOTE is effective for public policy sentiment analysis and can provide valuable insights for evaluating the implementation of the MBG Program.

Keywords: Sentiment Analysis; Free Nutritious Food; IndoBERT; SMOTE; FFNN

Abstrak

Program Makan Bergizi Gratis (MBG) menjadi salah satu kebijakan pemerintah yang memunculkan beragam respons masyarakat di media sosial, sehingga diperlukan analisis sentimen untuk memahami persepsi publik sebagai bahan evaluasi kebijakan. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap Program MBG di media sosial X menggunakan model IndoBERT dengan penerapan Synthetic Minority Over-sampling Technique (SMOTE) untuk mengatasi ketidakseimbangan kelas. Data dikumpulkan melalui proses scraping dan menghasilkan 2.260 tweet, kemudian diproses melalui tahapan pembersihan teks, normalisasi, tokenisasi, pelabelan sentimen menggunakan InSet Lexicon, penyeimbangan data dengan SMOTE, serta fine-tuning model IndoBERT. Evaluasi dilakukan menggunakan metrik akurasi, presisi, recall, dan F1-score. Hasil penelitian menunjukkan bahwa model memperoleh akurasi sebesar 80%, dengan dominasi sentimen negatif sekitar 75% yang dipengaruhi oleh isu anggaran, potensi korupsi, dan distribusi program, sedangkan sentimen positif menyoroti manfaat program terhadap gizi anak dan pemberdayaan ekonomi lokal. Temuan ini menunjukkan bahwa kombinasi IndoBERT dan SMOTE efektif untuk analisis sentimen kebijakan publik serta dapat menjadi masukan bagi evaluasi pelaksanaan Program MBG.

Kata kunci: Analisis Sentimen; IndoBERT; Makan Bergizi Gratis; FFNN; SMOTE.

1. Pendahuluan

Program Makan Bergizi Gratis (MBG) merupakan salah satu kebijakan strategis pemerintah Indonesia yang bertujuan meningkatkan kualitas gizi anak sekolah, balita, dan ibu hamil sebagai upaya mendukung pembangunan sumber daya manusia yang unggul. Program ini

menjadi bagian dari agenda prioritas nasional dan diharapkan mampu menurunkan angka stunting, meningkatkan kualitas pendidikan, serta memperkuat ketahanan pangan nasional. Namun, sejak diumumkan hingga tahap implementasi, program tersebut memunculkan berbagai tanggapan dari masyarakat, mulai dari dukungan terhadap manfaat sosial yang ditawarkan hingga kritik terhadap kesiapan anggaran, mekanisme pelaksanaan, serta dampaknya terhadap kondisi fiskal negara. Di era digital, media sosial X menjadi salah satu ruang utama bagi masyarakat untuk menyampaikan opini terhadap berbagai kebijakan pemerintah secara cepat dan terbuka. Besarnya volume opini yang terus berkembang menjadikan analisis sentimen sebagai pendekatan penting untuk memahami persepsi masyarakat secara sistematis. Informasi mengenai sentimen publik tersebut dapat dimanfaatkan sebagai bahan evaluasi bagi pemerintah untuk meningkatkan efektivitas komunikasi dan implementasi kebijakan publik.

Fenomena tingginya aktivitas masyarakat di media sosial X menyebabkan opini mengenai Program Makan Bergizi Gratis tersebar dalam jumlah besar dan dinamis. Karakteristik data pada media sosial berupa penggunaan bahasa tidak baku, singkatan, slang, emoji, maupun ironi menjadikan proses identifikasi sentimen secara manual sulit dilakukan dan membutuhkan waktu yang lama. Selain itu, opini yang berkembang tidak hanya bersifat positif dan negatif, tetapi juga mengandung sentimen netral yang sering kali sulit dibedakan karena dipengaruhi oleh konteks kalimat. Kondisi tersebut menyebabkan pemerintah maupun peneliti memerlukan pendekatan berbasis Natural Language Processing (NLP) yang mampu mengolah data teks secara otomatis dengan tingkat akurasi yang tinggi. Oleh karena itu, analisis sentimen terhadap Program Makan Bergizi Gratis di media sosial X menjadi penting untuk dilakukan agar dapat memberikan gambaran yang objektif mengenai persepsi masyarakat terhadap kebijakan tersebut serta menjadi dasar pengambilan keputusan yang lebih tepat.

Berbagai penelitian sebelumnya telah mengkaji analisis sentimen terhadap kebijakan pemerintah menggunakan berbagai pendekatan machine learning maupun deep learning. Putri et al. [1] meneliti analisis sentimen masyarakat terhadap Pemilihan Presiden Indonesia 2024 menggunakan model IndoBERT berbasis zero-shot learning pada data Twitter dan memperoleh akurasi sebesar 60%, sehingga menunjukkan bahwa model masih mengalami keterbatasan dalam mengidentifikasi sentimen negatif pada data berbahasa Indonesia. Selanjutnya, Sidauruk dan Herowati [2] melakukan analisis sentimen terhadap kinerja pemerintahan Prabowo–Gibran di media sosial X menggunakan model IndoBERT yang telah melalui proses fine-tuning. Penelitian tersebut menunjukkan bahwa IndoBERT mampu menghasilkan akurasi sebesar 78% dan memiliki kemampuan yang lebih baik dibandingkan dengan Support Vector Machine (SVM) dalam memahami konteks bahasa Indonesia. Penelitian yang secara khusus membahas Program Makan Bergizi Gratis dilakukan oleh Rohman dan Agung [3] dengan membandingkan algoritma IndoBERT, Support Vector Machine (SVM), dan Random Forest menggunakan 54.105 data tweet dari media sosial X. Hasil penelitian menunjukkan bahwa IndoBERT memperoleh performa terbaik dengan akurasi sebesar 87,60%, lebih tinggi dibandingkan SVM (83,39%) dan Random Forest (70,28%), sehingga membuktikan keunggulan model berbasis Transformer dalam memahami konteks bahasa Indonesia. Selain itu, Dinanti et al. [4] melakukan analisis sentimen terhadap Program Makan Bergizi Gratis menggunakan algoritma Naïve Bayes dan Random Forest berdasarkan metodologi CRISP-DM pada 6.419 komentar di YouTube. Penelitian tersebut menunjukkan bahwa Random Forest menghasilkan akurasi sebesar 77,43%, lebih tinggi dibandingkan Naïve Bayes yang hanya mencapai 65,64%. Dominasi sentimen netral sebesar 56,40% mengindikasikan bahwa masyarakat masih bersikap wait-and-see terhadap implementasi program MBG. Meskipun berbagai penelitian tersebut telah menunjukkan keberhasilan penerapan metode analisis sentimen pada isu kebijakan pemerintah maupun Program Makan Bergizi Gratis, sebagian besar penelitian masih berfokus pada perbandingan performa algoritma atau penggunaan platform media sosial tertentu, seperti Twitter atau YouTube. Selain itu, penelitian yang secara khusus memanfaatkan IndoBERT sebagai model utama untuk menganalisis dinamika sentimen masyarakat terhadap implementasi Program Makan Bergizi Gratis di media sosial X dengan data yang lebih mutakhir masih relatif terbatas. Oleh karena itu, penelitian ini berupaya mengisi kesenjangan tersebut dengan menerapkan model IndoBERT untuk memperoleh pemetaan sentimen masyarakat yang lebih akurat terhadap Program Makan Bergizi Gratis sebagai bahan evaluasi kebijakan pemerintah.

Berdasarkan kesenjangan tersebut, penelitian ini mengusulkan penerapan model IndoBERT untuk menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis pada media sosial X. IndoBERT merupakan model pre-trained language model yang

dikembangkan khusus untuk bahasa Indonesia dan mampu menghasilkan representasi bahasa berdasarkan konteks melalui arsitektur Transformer, sehingga sesuai untuk menangani karakteristik bahasa media sosial yang tidak baku, singkatan, serta variasi kosakata[5]. Kemampuan tersebut juga didukung oleh penelitian terdahulu yang menunjukkan bahwa IndoBERT dapat memberikan performa yang kompetitif bahkan lebih baik dibandingkan metode konvensional seperti SVM dan Random Forest dalam beberapa tugas analisis sentimen berbahasa Indonesia; pada penelitian terkait Program Makan Bergizi Gratis, IndoBERT memperoleh akurasi 87,60%, sedangkan SVM 83,39% dan Random Forest 70,28% [3]. Meskipun demikian, penggunaan IndoBERT belum secara otomatis mengatasi permasalahan ketidakseimbangan kelas, yang dapat menyebabkan model lebih cenderung memprediksi kelas mayoritas dan menurunkan kemampuan dalam mengenali kelas minoritas. Oleh karena itu, penelitian ini menerapkan Synthetic Minority Over-sampling Technique (SMOTE) pada ruang embedding yang dihasilkan IndoBERT untuk membentuk sampel sintesis pada kelas minoritas sehingga distribusi data pelatihan menjadi lebih seimbang; pendekatan SMOTE pada representasi embedding juga telah digunakan dalam penelitian analisis sentimen dan terbukti dapat membantu meningkatkan kemampuan model dalam menangani data yang tidak seimbang [6]. Selanjutnya, representasi fitur dari token [CLS] yang diperoleh dari IndoBERT digunakan sebagai masukan bagi Feed-Forward Neural Network (FFNN) sebagai classifier terpisah, karena representasi [CLS] dapat digunakan untuk merangkum informasi kontekstual keseluruhan urutan teks dan selanjutnya diteruskan ke lapisan klasifikasi [7]. Dengan demikian, kombinasi IndoBERT sebagai ekstraktor fitur, SMOTE sebagai metode penyeimbangan kelas, dan FFNN sebagai classifier diharapkan mampu mempertahankan informasi kontekstual bahasa Indonesia sekaligus meningkatkan kemampuan model dalam mengenali kelas sentimen minoritas, sehingga menghasilkan klasifikasi opini masyarakat terhadap Program Makan Bergizi Gratis yang lebih representatif dan dapat memberikan informasi empiris sebagai bahan evaluasi kebijakan pemerintah. Pendekatan ini diharapkan dapat meningkatkan generalisasi model pada data opini media sosial yang dinamis dan sangat tidak seimbang, sekaligus memberikan kontribusi bagi studi analisis sentimen kebijakan publik di Indonesia. Dengan mengeliminasi ambiguitas dan memfokuskan analisis pada polaritas positif dan negatif, pendekatan ini diyakini mampu menghasilkan evaluasi yang lebih tajam, objektif, dan komprehensif mengenai tingkat penerimaan publik terhadap program Makan Bergizi Gratis.

2. Metodologi

Penelitian ini menggunakan pendekatan eksperimental komputasional dalam bidang Natural Language Processing (NLP). Fokus utama adalah mengumpulkan opini masyarakat di platform X mengenai Program Makan Bergizi Gratis (MBG), kemudian mengklasifikasikannya ke dalam dua kelas sentimen biner (positif dan negatif) menggunakan model IndoBERT [8].

2.1 Tahapan Penelitian

Penelitian dilakukan melalui tahapan sistematis sebagai berikut: pengumpulan data, prapemrosesan data, pelabelan sentimen, pembagian dataset, pemodelan dengan IndoBERT, serta evaluasi performa model.

2.2 Pengumpulan Data

Data primer berupa teks cuitan diambil dari platform X menggunakan teknik crawling dengan kata kunci "makan bergizi gratis" dan "MBG" [9]. Pengumpulan data dilakukan pada periode tertentu dengan pembatasan bahasa Indonesia. Total data mentah yang berhasil dikumpulkan sebanyak 2.260 cuitan, kemudian dibersihkan dari duplikat menjadi 2.218 cuitan.

2.3 Prapemrosesan Data

Tahap pra-pemrosesan ini dapat dilakukan secara langsung melalui proses *crawling* atau dengan mengunggah data .csv hasil *crawling*. Tahap pra-pemrosesan ini mencakup pembersihan data, konversi ke huruf kecil, normalisasi, tokenisasi, dan penghapusan kata *stop word* agar data yang dihasilkan lebih mudah dipahami oleh mesin [10]. Oleh karena itu, dilakukan beberapa tahapan prapemrosesan sebagai berikut.

- 1) Data cleansing: proses ini mencakup penghapusan karakter atau simbol yang tidak relevan atau mengganggu, seperti tanda baca berlebihan, emoji, dan karakter khusus, untuk menyederhanakan data dan meningkatkan konsistensinya. Langkah ini bertujuan

untuk mengefisienkan data agar lebih mudah diproses serta meningkatkan konsistensi dalam representasi teks, seperti persamaan (1).

$$[T_{\{clean\}} = f_{\{clean\}}(T)] \quad (1)$$

- 2) Case Folding: mengubah semua huruf dalam teks menjadi huruf kecil untuk memastikan konsistensi dan mengurangi redundansi dalam analisis teks, seperti persamaan (2).

$$[T_{\{lower\}} = Lowercase(T_{\{clean\}})] \quad (2)$$

- 3) Normalization: Normalization bertujuan mengubah kata tidak baku, singkatan, bahasa gaul (slang), maupun kesalahan penulisan (typo) menjadi bentuk baku berdasarkan kamus normalisasi yang telah disusun.

- 4) Tokenization: Tahap untuk mengubah setiap kata menjadi bagian dari urutan kalimat, di mana setiap kata akan dipisahkan, seperti persamaan (3).

$$Tokens(T) = \{w_1, w_2, \dots, w_n\} \quad (3)$$

- 5) Stopword Removal: Stopword removal merupakan proses menghapus kata-kata yang memiliki frekuensi tinggi tetapi tidak memberikan informasi penting terkait polaritas sentimen. Daftar stopwords yang digunakan berasal dari NLTK Bahasa Indonesia dan ditambahkan dengan daftar stopwords kustom sesuai dengan karakteristik data media sosial, seperti persamaan (4).

$$Tokens' = \{w_i \mid w_i \notin S\} \quad (4)$$

2.4 Pelabelan Data

Pelabelan dilakukan secara otomatis menggunakan InSet Lexicon (Indonesia Sentiment Lexicon). Setiap cuitan diberi skor sentimen; skor > 0 dikategorikan positif, skor < 0 dikategorikan negatif. Data dengan skor netral dieliminasi. Hasil pelabelan: 556 data positif dan 1.662 data negatif.

2.5 Pembagian Dataset

Dataset dibagi dengan rasio 90:10 menggunakan stratified split, yaitu 90% untuk data latih dan 10% untuk data uji.

2.6 Pemodelan

Pemodelan pada penelitian ini menerapkan ekstraksi fitur tekstual dengan pendekatan transfer ilmu (transfer learning) serta pemrosesan multivariat untuk menanggulangi data yang imbalanced. Tahapan utama pemodelan mencakup:

1. Ekstraksi Fitur dengan IndoBERT: Teks yang telah dibersihkan dimanifestasikan menjadi representasi token, kemudian diinput ke dalam arsitektur dasar BertModel untuk memperoleh word embeddings berdimensi tinggi. Ekstraksi spesifik difokuskan pada token representasi kalimat utuh, yakni token [CLS] (dimensi 768) yang mewakili fitur semantik tweet tersebut [11], seperti persamaan (5).

$$x_{baru} = x_i + \lambda(x_{nn} - x_i) \quad (5)$$

2. Penyeimbangan Data dengan SMOTE: Rasio matriks fitur embedding dievaluasi menggunakan teknik pemisahan data (Train-Test split). Untuk mengatasi kelemahan jumlah kelas sentimen minoritas pada data latih, diterapkan algoritma SMOTE (Synthetic Minority Over-sampling Technique). SMOTE menggeneralisasi contoh sintesis di antara titik-titik sampel fitur minoritas pada ruang vektor IndoBERT sehingga jumlah data antara sentimen positif dan negatif merata secara matematis [12], seperti persamaan (6).

$$x_{baru} = x_i + \lambda(x_{nn} - x_i) \quad (6)$$

3. Klasifikasi Feed Forward Neural Network (FFNN): Vektor token [CLS] berdimensi 768 yang telah stabil kemudian disalurkan ke model jaringan saraf tiruan sederhana (classifier). Arsitektur ini terdiri atas serangkaian Fully Connected Layer tersembunyi yang diregulasi dengan fungsi aktivasi (ReLU) serta fungsi penahan laju overfitting (Dropout). Pada ujung lapisan klasifikasi, arsitektur memprediksi nilai keluaran berupa persentase probabilitas untuk dua kelas klasifikasi untuk evaluasi fungsi kerugian (Cross-Entropy Loss), seperti persamaan (7).

$$f = h_{CLS} \in R^{768} \quad (7)$$

2.7 Evaluasi Model

Untuk mengukur seberapa baik model IndoBERT dalam mengklasifikasikan sentimen opini masyarakat, evaluasi dilakukan menggunakan Confusion Matrix sebagai alat ukur. Evaluasi ini akan menghitung nilai True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN) [13].

1) Accuracy(Akurasi)

Menggambarkan seberapa akurat model dalam mengklasifikasikan data secara keseluruhan (baik positif maupun negatif) dengan benar, seperti persamaan (8).

$$Accuracy = \frac{TP + TN}{TP + FP + FN + TN} = x 100\% \quad (8)$$

2) Precision (Presisi)

Menggambarkan tingkat ketepatan antara informasi yang diminta dan jawaban prediksi yang diberikan oleh model (seberapa banyak data yang diprediksi positif memang benar-benar positif), seperti persamaan (9).

$$Precision = \frac{TP}{TP + FP} = x 100\% \quad (9)$$

3) Recall (Sensitivitas)

Menggambarkan tingkat keberhasilan model dalam menemukan kembali seluruh informasi dari data yang sebenarnya bersentimen positif, seperti persamaan(10).

$$Recall = \frac{TP}{TP + FN} = x 100\% \quad (10)$$

4) F-1 Score

Merupakan perbandingan rata-rata harmonik dari nilai presisi dan recall, sangat berguna untuk mengevaluasi model apabila terdapat sedikit ketidakseimbangan kelas dalam dataset, seperti persamaan (11).

$$F1 - score = \frac{2 \times recall \times precision}{recall + precision} \quad (11)$$

3. Hasil dan Pembahasan

Bagian ini menyajikan hasil implementasi komputasi dan evaluasi performa model IndoBERT-FFNN yang telah dilatih.

3.1 Hasil Pengumpulan Data

Penelitian ini berhasil mengumpulkan 2.260 cuitan mentah dari platform X menggunakan teknik crawling.

Tabel 1. Data Mentah hasil Crawling

Text
@ardisatriawan Besok kalo makan mbg di sel tahanan koruptor aja Kan lebih bersih
@txtdrimedia iyaa maling...MBG, Maling Berkedok Gizi
@tempodotco gaji kecil tapi punya dapur mbg di mana-mana.
@RadioElshinta Digitalisasi MBG nasional yang transparan, tepat sasaran, dan berkelanjutan.
#makanbergizigratis #nawacita #makanbergizi https://t.co/pzt3ohJ1ll

3.2 Hasil Preprocessing Data

Pra-pemrosesan: bagaimana cara mengolah data mentah agar lebih bersih sehingga model lebih mudah dipahami. Ada 5 langkah di sini untuk tahap pra-pemrosesan yang dimulai dari cleaning, case folding, normalisasi, tokenisasi, dan penghapusan stopwords.

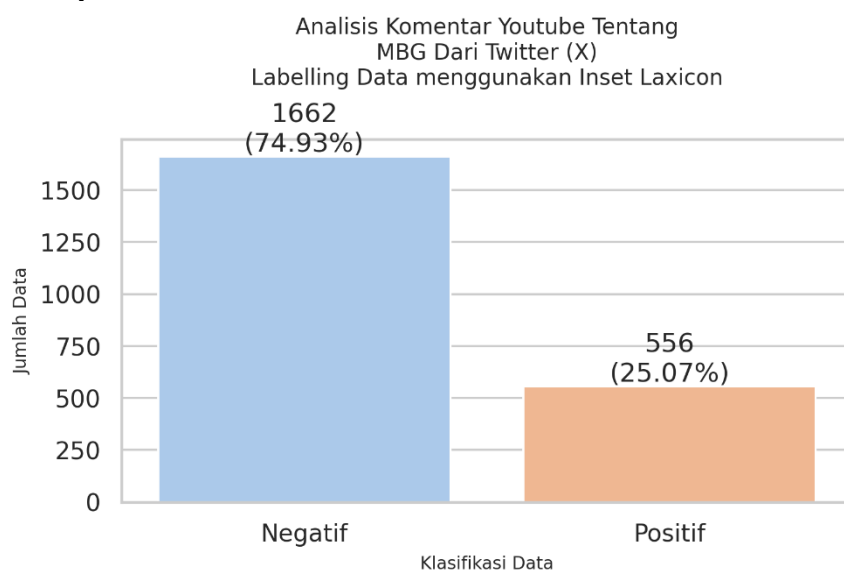
Tabel 2. Hasil Tahap Preprocessing

Proses	Sebelum	Sesudah
Cleaning	@RadioElshinta Digitalisasi MBG nasional yang transparan, tepat sasaran, dan berkelanjutan. #makanbergizigratis #nawacita #makanbergizi https://t.co/pzt3ohJ1ll	Digitalisasi MBG nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi
Case Folding	Digitalisasi MBG nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi	digitalisasi mbg nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi
Normalisasi	digitalisasi mbg nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi	digitalisasi mbg nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi
Tokenization	digitalisasi mbg nasional transparan tepat sasaran berkelanjutan makanbergizigratis nawacita makanbergizi	(digitaliasi, mbg, nasional, transparan, tepat, sasaran, berkelanjutan, makanbergizigratis, nawacita, makanbergizi)
Stopword	(digitaliasi, mbg, nasional, transparan, tepat, sasaran, berkelanjutan, makanbergizigratis, nawacita, makanbergizi)	digitalisasi mbg nasional transparan sasaran berkelanjutan makanbergizigratis nawacita makanbergizi

3.3 Hasil Pelabelan Sentimen

Setelah tahap preprocessing selesai, data kemudian diberi label sentimen untuk menentukan kategori opini pada setiap tweet. Proses pelabelan dimulai dengan menentukan apakah setiap tweet menunjukkan sentimen positif atau negatif terhadap program tersebut. Sentimen positif menunjukkan dukungan atau pandangan yang baik terhadap program makan siang gratis, sedangkan sentimen negatif mencerminkan ketidaksetujuan atau kritik.

Mengingat jumlah data yang besar, pelabelan dilakukan secara otomatis menggunakan metode berbasis leksikon. Proses ini menggunakan kamus sentimen bahasa Indonesia yang mencakup kata-kata positif dan negatif. Setiap komentar diberikan skor numerik berdasarkan kemunculan kata-kata dalam kamus tersebut, yang kemudian dikonversi menjadi dua label: Positif (skor > 0) dan Negatif (skor < 0). Hasil pelabelan menunjukkan positif sebanyak 556 data dan negatif sebanyak 1662 data.

**Gambar 1.** Hasil Labeling Inset Lexicon

Dapat dilihat bahwa data tidak seimbang, sehingga akan diterapkan SMOTE untuk menyeimbangkan data pelatihan.

Tabel 3. Pemberian Score untuk Hasil Sentimen

Stopword	Score	Sentimen
program mbg hadir dukungan nyata keluarga mental finansial lanjutkan mbg	1	Positif
namanya maling ketahuan coba ketahuan lolos kayak duit mbg ambil bajingan	-5	Negatif
anggaran mbg korupsi tutup mata anggota tutup mulut banyaknya kebocoran anggaran mbg	-5	Negatif

3.4 Hasil Pemodelan

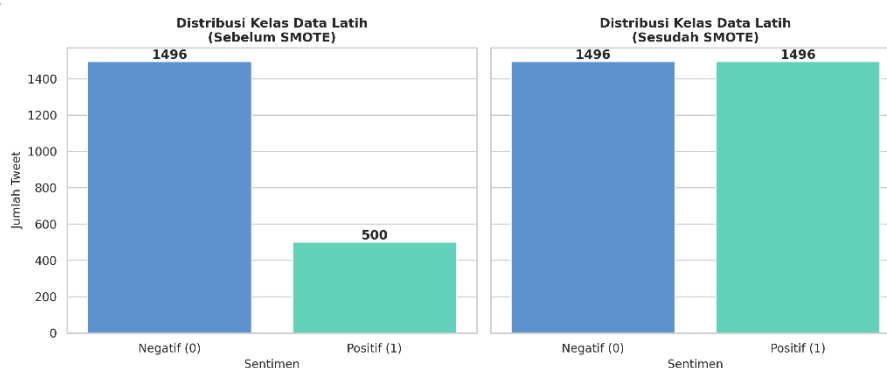
Hasil pemodelan pada penelitian ini diperoleh melalui tiga tahapan utama, yaitu ekstraksi fitur menggunakan IndoBERT, penyeimbangan data menggunakan SMOTE, dan klasifikasi menggunakan Feed Forward Neural Network (FFNN). Setiap tahapan menghasilkan keluaran (output) yang menjadi masukan bagi tahapan berikutnya hingga diperoleh hasil klasifikasi sentimen.

1) Hasil Ekstraksi Fitur Menggunakan IndoBERT

Pada tahap ini, data tweet yang telah melalui proses prapemrosesan diubah menjadi representasi numerik menggunakan model IndoBERT (indobenchmark/indobert-base-p1). Setiap tweet ditokenisasi menggunakan BertTokenizer dan diproses oleh BertModel untuk menghasilkan representasi vektor (embedding). Penelitian ini menggunakan token [CLS] sebagai representasi keseluruhan kalimat, sehingga setiap tweet direpresentasikan sebagai vektor fitur berdimensi 768. Vektor embedding tersebut mampu menangkap informasi semantik dan konteks bahasa Indonesia yang selanjutnya digunakan sebagai masukan pada tahap klasifikasi.

2) Hasil Penyeimbangan Data Menggunakan SMOTE

Setelah proses ekstraksi fitur, embedding hasil IndoBERT dipisahkan menjadi data latih dan data uji menggunakan metode Train-Test Split. Selanjutnya dilakukan penyeimbangan kelas pada data latih menggunakan algoritma Synthetic Minority Over-sampling Technique (SMOTE). Proses ini menghasilkan data sintetis pada kelas minoritas dengan membentuk sampel baru berdasarkan kedekatan antarvektor embedding. Hasil penerapan SMOTE adalah jumlah data pada setiap kelas sentimen menjadi lebih seimbang, sehingga model FFNN tidak cenderung memprediksi kelas mayoritas dan memiliki kemampuan generalisasi yang lebih baik pada kedua kelas sentimen.



Gambar 2. Penggunaan SMOTE untuk data Imbalance

3) Hasil Klasifikasi Menggunakan Feed Forward Neural Network (FFNN)

Embedding hasil IndoBERT yang telah diseimbangkan menggunakan SMOTE selanjutnya digunakan sebagai masukan ke dalam model Feed Forward Neural Network (FFNN). Arsitektur FFNN terdiri atas Fully Connected Layer dengan ukuran 768 → 256 neuron, diikuti oleh fungsi aktivasi ReLU, Dropout sebesar 0,3 untuk mengurangi overfitting, serta output layer dengan 2 neuron yang merepresentasikan kelas sentimen positif dan negatif. Nilai keluaran dari output layer kemudian diproses menggunakan

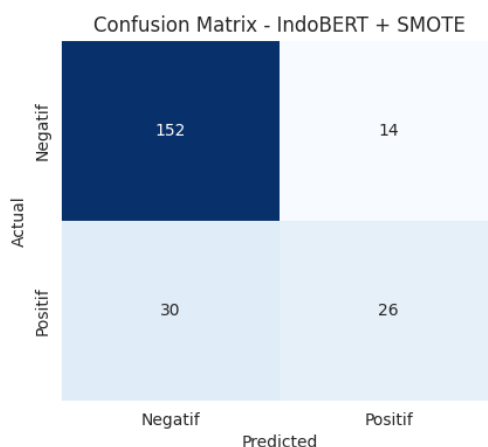
fungsi Softmax untuk menghasilkan probabilitas masing-masing kelas. Kelas dengan probabilitas tertinggi dipilih sebagai hasil prediksi sentimen untuk setiap tweet. Selama proses pelatihan, model dioptimalkan menggunakan Adam Optimizer dengan learning rate 0,0001 dan fungsi Cross-Entropy Loss sebagai fungsi objektif. Hasil akhir dari tahap ini berupa label prediksi sentimen (positif atau negatif) beserta nilai probabilitasnya yang selanjutnya digunakan pada tahap evaluasi model menggunakan metrik seperti akurasi, presisi, recall, F1-score, dan confusion matrix.

Tabel 4. Hasil Klasifikasi

Kelas	Precision	Recall	F1-Score	Support
Negatif	0,84	0,92	0,87	166
Positif	0,65	0,46	0,54	56

3.5 Hasil Evaluasi Model

Evaluasi model dilakukan untuk mengukur kemampuan model IndoBERT dalam mengklasifikasikan sentimen masyarakat terhadap Program Makan Bergizi Gratis di media sosial X. Pengujian menggunakan data uji (testing set) yang tidak digunakan selama proses pelatihan model. Hasil prediksi kemudian dibandingkan dengan label sebenarnya menggunakan Confusion Matrix sehingga diperoleh nilai True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Berdasarkan Confusion Matrix tersebut, dihitung nilai Accuracy, Precision, Recall, dan F1-Score sebagai indikator kinerja model.



Gambar 3. Confusion Matrix Indobert + SMOTE

Berdasarkan hasil Confusion Matrix, diketahui bahwa model mampu mengklasifikasikan 26 data sentimen positif dan 152 data sentimen negatif dengan benar. Namun demikian, masih terdapat 14 data positif yang diprediksi sebagai negatif (False Negative) dan 30 data negatif yang diprediksi sebagai positif (False Positive).

Evaluasi Model Accuracy, Precision, Recall, F-1 Score dapat di hitung dari nilai di atas dalam table berikut:

Tabel 5. Hasil Evaluasi Model

Metrik	Nilai
Accuracy	0,80
Precision	0,46
Recall	0,65
F-1 Score	0,54

Berdasarkan hasil Confusion Matrix, masih terdapat kesalahan klasifikasi berupa 14 data positif yang diprediksi sebagai negatif (False Negative) dan 30 data negatif yang diprediksi sebagai positif (False Positive). Kesalahan tersebut menunjukkan bahwa model IndoBERT masih mengalami kesulitan dalam memahami karakteristik bahasa pada media sosial X yang cenderung informal, singkat, menggunakan singkatan, emoji, serta mengandung makna yang bergantung pada konteks. Pada False Negative, data positif dapat salah diklasifikasikan sebagai negatif

karena sentimen positif disampaikan secara tidak langsung atau menggunakan kalimat yang ambigu. Sementara itu, False Positive dapat terjadi ketika kata-kata bernuansa positif digunakan dalam kalimat bernada kritik, sindiran, atau sarkasme, sehingga makna keseluruhan kalimat tidak sesuai dengan makna kata secara literal.

Sebagai contoh, pada data yang termasuk False Positive, terdapat kemungkinan kalimat seperti "Programnya bagus sekali, sampai banyak anak yang masih mengeluh soal makanannya". Kalimat tersebut mengandung kata "bagus", tetapi konteks keseluruhannya menunjukkan kritik sehingga seharusnya berlabel negatif. Contoh False Negative dapat berupa kalimat "Akhirnya anak-anak bisa menikmati makanan bergizi setiap hari, semoga program ini terus berjalan". Meskipun mengandung sentimen positif terhadap program, model dapat salah memprediksinya sebagai negatif apabila terdapat kata atau konteks tertentu yang menyerupai pola data negatif pada proses pelatihan. Kesalahan seperti ini menunjukkan bahwa model perlu memahami konteks kalimat secara keseluruhan, bukan hanya mengenali kata-kata tertentu. Oleh karena itu, analisis terhadap sampel data yang salah diklasifikasikan dapat digunakan sebagai bahan evaluasi untuk meningkatkan kualitas preprocessing, pelabelan data, dan performa model pada penelitian selanjutnya.

3.6 Hasil Prediksi Baru

Setelah model IndoBERT selesai dilatih dan dievaluasi, dilakukan pengujian menggunakan data teks baru yang tidak termasuk dalam data pelatihan maupun data pengujian (unseen data). Tujuan pengujian ini adalah untuk mengetahui kemampuan model dalam melakukan generalisasi terhadap opini masyarakat yang belum pernah dipelajari sebelumnya.

Contoh Prediksi 1: "Program makan bergizi gratis sangat membantu meningkatkan asupan gizi anak-anak di Indonesia."

Tabel 6. Proses Input Data Baru

Proses	Hasil
Data awal	Program makan bergizi gratis sangat membantu meningkatkan asupan gizi anak-anak di Indonesia.
Cleansing	Program makan bergizi gratis sangat membantu meningkatkan gizi anak Indonesia
Case Folding	program makan bergizi gratis sangat membantu meningkatkan gizi anak indonesia
Tokenisasi	[program, makan, bergizi, gratis, sangat, membantu, meningkatkan, gizi, anak, indonesia]
Embedding	Vektor fitur 768 dimensi (token [CLS])
Klasifikasi FFNN	Menghasilkan dua nilai logit (Negatif dan Positif)
Softmax	Negatif = 0,024; Positif = 0,976

Tabel 7. Hasil Prediksi Data Baru

Kelas	Probabilitas
Negatif	2,4%
Positif	97,6%

Validasi dilakukan dengan membandingkan hasil prediksi model dengan makna semantik setiap kalimat. Pada contoh ini, model memberikan probabilitas tertinggi pada kelas positif karena kalimat tersebut mengandung ungkapan yang mendukung program, seperti "*membantu*" dan "*meningkatkan gizi*". Selain validasi semantik, hasil prediksi juga didukung oleh evaluasi model pada data uji yang menunjukkan bahwa model memiliki nilai Accuracy, Precision, Recall, dan F1-Score yang tinggi. Hal ini mengindikasikan bahwa model memiliki kemampuan generalisasi yang baik sehingga mampu mengklasifikasikan data baru secara konsisten. Meskipun demikian, validasi dalam penelitian ini masih bersifat kualitatif melalui analisis isi kalimat.

3.7 Pembahasan

Hasil penelitian ini menunjukkan bahwa model IndoBERT yang dipadukan dengan SMOTE mampu mengidentifikasi sentimen terhadap Program Makan Bergizi Gratis dengan cukup baik, ditunjukkan oleh akurasi keseluruhan sebesar 80%. Temuan ini menjawab tujuan

penelitian, yaitu mengukur kecenderungan opini publik dan mengetahui isu-isu yang paling dominan dalam percakapan mengenai program tersebut. Dominasi sentimen negatif sekitar 75% mengindikasikan bahwa percakapan warganet lebih banyak dipengaruhi oleh kekhawatiran terhadap beban anggaran negara, potensi korupsi, dan masalah distribusi. Sementara itu, sentimen positif lebih menonjolkan manfaat program bagi gizi anak dan pemberdayaan ekonomi lokal. Dengan demikian, penelitian ini tidak hanya memetakan polaritas opini, tetapi juga memperlihatkan bahwa persepsi publik terhadap program lebih banyak bersifat kritis daripada suportif.

Secara teoritis, temuan ini sejalan dengan landasan ilmiah yang melatarbelakangi penggunaan IndoBERT dan SMOTE. BERT dirancang untuk memahami konteks kata secara dua arah, sedangkan IndoBERT merupakan adaptasi yang telah melalui pre-training pada korpus berbahasa Indonesia sehingga lebih peka terhadap nuansa bahasa lokal dan bentuk bahasa informal di media sosial. Karena itu, keberhasilan model dalam menangkap konteks semantik bahasa Indonesia informal pada data penelitian ini logis secara teoritis. Di sisi lain, SMOTE memang ditujukan untuk menyeimbangkan distribusi kelas agar kelas minoritas tidak terus didominasi oleh kelas mayoritas. Namun, pada data teks, penyeimbangan kelas tidak selalu secara otomatis meningkatkan pemahaman makna karena proses oversampling bekerja pada ruang fitur numerik, bukan pada konteks kalimat secara utuh. Hal ini menjelaskan mengapa kelas positif masih perlu ditingkatkan meskipun penyeimbangan data telah diterapkan.

Temuan penelitian ini dapat dibandingkan dengan beberapa penelitian sebelumnya yang sama-sama membahas analisis sentimen terhadap Program Makan Bergizi Gratis maupun isu kebijakan publik lain menggunakan pendekatan machine learning dan deep learning. Penelitian Siregar et al. [14] yang menerapkan Naïve Bayes untuk analisis sentimen netizen terhadap Program Makan Siang Gratis di media sosial X memperoleh akurasi 65%, dengan sentimen negatif sebagai dominan. Hasil tersebut menunjukkan bahwa pendekatan berbasis klasifikasi konvensional masih mampu memberikan gambaran awal terhadap opini publik, namun performanya belum setinggi penelitian ini yang menggunakan IndoBERT, SMOTE, dan FFNN. Dengan demikian, hasil penelitian ini mendukung temuan bahwa model berbasis representasi semantik dapat memberikan kinerja yang lebih baik dalam memahami opini masyarakat di media sosial.

Hasil penelitian ini juga dapat dibandingkan dengan penelitian Tribuana et al. [15] yang menganalisis sentimen MBG di TikTok menggunakan IndoBERT dengan penerapan SMOTE. Penelitian tersebut memperoleh akurasi 72,39% dan menunjukkan bahwa SMOTE membantu meningkatkan representasi kelas minoritas, meskipun akurasi keseluruhan justru lebih rendah dibandingkan dengan model tanpa SMOTE. Dibandingkan dengan penelitian tersebut, hasil penelitian ini lebih tinggi, yang menunjukkan bahwa kombinasi IndoBERT, SMOTE, dan FFNN pada data X mampu menghasilkan klasifikasi yang lebih baik. Perbedaan ini juga mengindikasikan bahwa efektivitas SMOTE sangat dipengaruhi oleh karakteristik platform, distribusi kelas, serta desain model klasifikasi yang digunakan.

Selain itu, penelitian Haikal dan Erfina [16] tentang sentimen MBG di X menggunakan IndoBERT menghasilkan akurasi 60,27%, dengan dominasi sentimen negatif dan netral yang paling sulit dibedakan. Hasil tersebut sejalan dengan penelitian ini karena sama-sama menegaskan bahwa IndoBERT efektif digunakan pada data media sosial berbahasa Indonesia. Namun, penelitian ini berbeda karena memfokuskan analisis pada dua kelas sentimen, yaitu positif dan negatif, setelah kelas netral dieliminasi. Perbedaan ini memberikan temuan baru bahwa ketika fokus diarahkan pada polaritas biner dan ketidakseimbangan kelas ditangani dengan SMOTE, performa model dapat meningkat secara lebih optimal.

Jika dibandingkan dengan penelitian Prianto et al. [17] yang menggunakan Random Forest dan Support Vector Machine dengan penyeimbangan data berbasis SMOTE pada MBG di X, hasil klasifikasi yang diperoleh masih relatif rendah, yakni pada kisaran 0,32–0,34. Temuan tersebut menunjukkan bahwa penyeimbangan data saja belum cukup untuk meningkatkan performa secara signifikan apabila model yang digunakan masih berbasis fitur klasik seperti TF-IDF. Dibandingkan dengan hasil tersebut, penelitian ini menunjukkan bahwa pendekatan berbasis kontekstual embedding pada IndoBERT lebih unggul dalam menangkap makna kalimat dan konteks bahasa di media sosial. Dengan demikian, penelitian ini memperkuat argumen bahwa model semantik kontekstual lebih sesuai untuk analisis sentimen kebijakan publik dalam bahasa Indonesia.

Perbandingan berikutnya dapat dilakukan dengan penelitian Fahrezi et al. [18] yang menganalisis sentimen debat Pilpres 2024 di YouTube menggunakan LSTM dan IndoBERT. Penelitian tersebut menunjukkan bahwa LSTM menghasilkan akurasi yang sangat tinggi, sedangkan IndoBERT secara mandiri belum memberikan performa yang sebaik ensemble hybrid. Temuan ini berbeda dengan penelitian ini karena objek, platform, dan karakter data yang digunakan juga berbeda. Namun, hasil tersebut tetap relevan karena menunjukkan bahwa keberhasilan model sangat bergantung pada konfigurasi arsitektur, karakteristik data, dan strategi penanganan kelas. Dalam konteks penelitian ini, penggunaan IndoBERT yang dipadukan dengan SMOTE dan FFNN menghasilkan hasil yang lebih stabil pada data opini publik di X.

Penelitian Nilapaksi et al. [19] mengenai sentimen publik terhadap Magang Berdampak 2025 di X menggunakan IndoBERT juga menunjukkan performa yang cukup baik dengan akurasi 0,7481 serta dominasi sentimen netral. Hasil tersebut mendukung temuan penelitian ini bahwa IndoBERT efektif digunakan untuk menganalisis opini publik terhadap kebijakan atau program pemerintah di media sosial. Meski demikian, penelitian ini memberikan temuan tambahan bahwa penanganan ketidakseimbangan kelas dengan SMOTE pada embedding IndoBERT dapat meningkatkan performa. Secara keseluruhan, penelitian ini tidak hanya mendukung hasil-hasil sebelumnya, tetapi juga memberikan kontribusi baru berupa pengujian kombinasi IndoBERT, SMOTE, dan FFNN pada data X bertema MBG yang menunjukkan bahwa pendekatan ini mampu meningkatkan sensitivitas model terhadap distribusi kelas yang tidak seimbang.

Adapun keterbatasan utama penelitian ini adalah ketergantungan pada pelabelan otomatis menggunakan InSet Lexicon yang kurang sensitif terhadap sarkasme, ironi, serta konteks bahasa yang tidak formal. Keterbatasan ini berpotensi membuat sebagian data latih tidak sepenuhnya merepresentasikan sentimen yang sebenarnya. Hal ini sejalan dengan penelitian terdahulu yang juga menekankan bahwa penggunaan label otomatis atau pseudo-label pada tahap awal masih perlu diperbaiki melalui pelabelan manual dan penggunaan dataset yang lebih besar. Penelitian selanjutnya disarankan menggunakan pendekatan pelabelan manual atau hybrid labeling, memperluas jumlah data, serta membandingkan IndoBERT dengan model lain yang lebih spesifik untuk bahasa media sosial, agar klasifikasi kelas positif dan netral menjadi lebih stabil dan akurat.

4. Simpulan

Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis (MBG) di media sosial X menggunakan model IndoBERT dengan penerapan SMOTE. Hasil penelitian menunjukkan bahwa model mampu mengklasifikasikan sentimen dengan akurasi sebesar 80%, sehingga efektif untuk memahami konteks bahasa Indonesia di media sosial. Penerapan SMOTE membantu mengurangi ketidakseimbangan distribusi kelas, meskipun performa pada kelas positif masih perlu ditingkatkan. Analisis sentimen juga menunjukkan bahwa opini masyarakat didominasi oleh sentimen negatif (sekitar 75%) yang dipengaruhi oleh isu anggaran, potensi korupsi, dan distribusi program, sedangkan sentimen positif lebih banyak berkaitan dengan manfaat program bagi gizi anak dan pemberdayaan ekonomi lokal.

Secara keseluruhan, penelitian ini menunjukkan bahwa kombinasi IndoBERT dan SMOTE merupakan pendekatan yang efektif untuk analisis sentimen kebijakan publik serta dapat memberikan gambaran mengenai persepsi masyarakat terhadap Program MBG. Hasil penelitian ini diharapkan dapat menjadi masukan bagi pemerintah dalam mengevaluasi pelaksanaan program sekaligus memperkaya kajian analisis sentimen berbasis *Natural Language Processing* (NLP). Penelitian ini masih memiliki keterbatasan dalam penggunaan pelabelan otomatis berbasis InSet Lexicon yang kurang mampu mengenali sarkasme dan konteks bahasa tertentu. Oleh karena itu, penelitian selanjutnya disarankan untuk menggunakan pelabelan manual atau hybrid labeling, memperluas dataset, serta membandingkan IndoBERT dengan model bahasa lainnya guna meningkatkan akurasi klasifikasi sentimen.

Daftar Referensi

- [1] D. I. Putri, A. N. Alfian, M. Y. Putra, and P. D. Mulyo, "IndoBERT Model Analysis : Twitter Sentiments on Indonesia ' s 2024 Presidential Election," *J. Appl. Informatics Comput.*, vol. 8, no. 1, pp. 7–12, 2024.
- [2] V. E. Sidauruk and W. Herowati, "Indobert-Based Sentiment Analysis of Political Discourse

- on Platform X : The Case Of Prabowo-Gibran Administration,” *J. Appl. Informatics Comput.*, vol. 10, no. 1, pp. 673–683, 2026.
- [3] S. Rohman and I. W. P. Agung, “Komparasi algoritma indobert, svm, dan random forest dalam analisis sentimen program makan bergizi gratis,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 6, pp. 9464–9471, 2025.
 - [4] R. P. Dinanti, S. Sarah, F. Dwi, A. Hafiz, and A. Alfarizi, “Analisis Sentimen Program Makanan Bergizi Gratis Menggunakan Naïve Bayes dan Random Forest Berbasis,” *J. Inf. Syst. Comput. Sci. Inf. Technol.*, vol. 7, no. 1, pp. 153–163, 2026.
 - [5] B. Wilie *et al.*, “IndoNLU : Benchmark and Resources for Evaluating Indonesian Natural Language Understanding,” *Proc. of the 1st Conf. of the Asia-Pacific Chapter of the Assoc. Comput. Linguist. 10th Int. Jt. Conf. Nat. Lang. Process.*, vol. 7, pp. 843–857, 2020.
 - [6] F. Y. A’la, “Edumatic : Jurnal Pendidikan Informatika Optimasi Klasifikasi Sentimen Ulasan Game Berbahasa Indonesia : IndoBERT dan SMOTE untuk Menangani Ketidakseimbangan Kelas,” *Edumatic J. Pendidik. Inform.*, vol. 9, no. 1, pp. 256–265, 2025, doi: 10.29408/edumatic.v9i1.29666.
 - [7] M. C. Kenton, L. Kristina, and J. Devlin, “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding,” *Proc. 2019 Conf. North Am. Chapter Assoc. Comput. Linguist. Hum. Lang. Technol.*, vol. 1, pp. 4171–4186, 2019.
 - [8] A. Jazuli, “Optimizing Aspect-Based Sentiment Analysis Using BERT for Comprehensive Analysis of Indonesian Student Feedback,” *Appl. Sci.*, vol. 15, pp. 1–28, 2025.
 - [9] A. Jazuli, Widowati, and R. Kusumaningrum, “Auto Labeling to Increase Aspect-Based Sentiment Analysis Using K-Nearest Neighbors Method,” *E3S Web Conf.*, vol. 359, p. 05001, Oct. 2022, doi: 10.1051/e3sconf/202235905001.
 - [10] A. Jazuli and R. Kusumaningrum, “Enhancing Aspect-Based Sentiment Analysis through Data Labeling Classification on Student Reviews Using a Text Sampling Approach,” *Migr. Lett.*, vol. 20, pp. 801–810, 2023.
 - [11] A. Jazuli, A. A. Chamid, and R. Kusumaningrum, “Transformer-based semantic indexing for aspect-based sentiment analysis using an enhanced index generation algorithm with BERT,” *Int. J. Adv. Technol. Eng. Explor.*, vol. 12, no. 127, 2025.
 - [12] E. Wijayanti and C. E. Widodo, “Hybrid Machine Learning for Early Prediction of At-Risk Students with Imbalanced Data,” *J. Appl. Data Sci.*, vol. 7, no. 2, pp. 1648–1660, 2026.
 - [13] S. H. Majidah and P. Hendikawati, “Analisis Sentimen Program Makan Bergizi Gratis Menggunakan SVM DAN KNN Dengan TF-IDF DAN TF-ABS,” *E-Jurnal Mat.*, vol. 15, no. 1, pp. 1–9, 2026.
 - [14] R. Amelia, B. Siregar, A. R. Manik, A. Syahri, M. B. Akbar, and F. Ramadhani, “Penerapan Naïve Bayes Untuk Analisis Sentimen Netizen Terhadap Program Makan Siang Gratis Pada Media Social X,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 4, pp. 5621–5628, 2025.
 - [15] B. Tribuana and W. Murniati, “Analisis Sentimen Terhadap Program Makan Bergizi Gratis (MBG) pada Media Sosial Menggunakan Algoritma IndoBERT,” *METIK J. Vol.*, vol. 10, pp. 176–187, 2026.
 - [16] M. Haikal and A. Erfina, “JURNAL LOCUS : Penelitian & Pengabdian Analisis Sentimen Warganet terhadap Program Makan Bergizi Gratis (MBG) Menggunakan Model Indobert,” *J. LOCUS Penelit. Pengabdi.*, vol. 5, no. 3, pp. 1626–1636, 2026.
 - [17] S. E. Prianto, R. E. Saputro, M. I. Komputer, F. I. Komputer, and U. A. Purwokerto, “Analisis Sentimen Program Makan Bergizi Gratis Menggunakan Random Forest dan Support Vector Machine dengan Penyeimbangan Data Berbasis SMOTE,” *J. Pendidik. dan Teknol. Indones.*, vol. 6, no. 2, pp. 353–365, 2026.
 - [18] M. Fahrezi, Y. B. Pratama, and A. Pramudiyantoro, “Analisis Sentimen Debat Publik Pilpres 2024 Menggunakan Metode Algoritma LSTM dan IndoBERT Pada Platform Youtube,” *JPIM J. Penelit. Ilm. Multidisipliner*, vol. 02, no. 03, pp. 1936–1961, 2025.
 - [19] B. Mery, R. Eufra, T. N. Fatyanosa, and P. P. Adikara, “Analisis Sentimen Publik Terhadap Magang Berdampak 2025 di Platform X / Twitter Menggunakan Model Indobert,” *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 4, pp. 1–10, 2026.