

Penerapan Text Mining pada Komentar Sosmed Instagram dan Youtube Mengenai Program Makan Bergizi Gratis untuk Analisis Sentimen Lintas Isu Model Indobert dan Teknik Smote

Selvia Anggraini ^{1*}, Cindi Wulandari ², Lukman Sunardi ³, Bunga Intan ⁴

^{1, 2, 3, 4} Universitas Bina Insan, Indonesia

* v.selvianggraini@gmail.com

Abstrak

Urgensi penelitian ini adalah meningkatkan akurasi klasifikasi sentimen berbahasa Indonesia dengan mengatasi ketidakseimbangan data melalui penerapan *Stopword Removal* dan SMOTE. Program Makan Bergizi Gratis (MBG) merupakan kebijakan pemerintah yang memperoleh beragam respons masyarakat melalui media sosial. Instagram dan YouTube menyediakan ruang interaksi yang memungkinkan masyarakat menyampaikan dukungan, kritik, maupun tanggapan terhadap pelaksanaan program. Besarnya volume komentar menyebabkan identifikasi sentimen secara manual menjadi tidak efisien, sehingga diperlukan pendekatan *text mining* dan pemodelan bahasa yang mampu mengolah data teks berbahasa Indonesia dalam jumlah besar. Penelitian ini bertujuan menerapkan *text mining* pada komentar Instagram dan YouTube mengenai Program MBG menggunakan model IndoBERT dan teknik *Synthetic Minority Over-sampling Technique* (SMOTE), serta menganalisis sentimen berdasarkan isu yang berkaitan dengan program. Data dikumpulkan menggunakan teknik *web scraping* dan menghasilkan 52.326 komentar. Setelah proses *preprocessing* dan pembersihan data, sebanyak 47.861 komentar digunakan dalam analisis. Tahapan *preprocessing* meliputi *case folding*, *cleaning*, normalisasi kata tidak baku, tokenisasi, *stopword removal*, dan *stemming*. Data kemudian diberi label menjadi sentimen positif, netral, dan negatif. Distribusi awal terdiri atas 3.131 komentar positif, 31.136 netral, dan 4.025 negatif. Ketidakseimbangan kelas ditangani menggunakan SMOTE sebelum proses klasifikasi dengan IndoBERT. Evaluasi model dilakukan menggunakan *accuracy*, *precision*, *recall*, F1-score, dan *confusion matrix*. Hasil pengujian menunjukkan bahwa model IndoBERT dengan SMOTE menghasilkan *accuracy* sebesar 89,88%, *precision* 92,55%, *recall* 89,88%, dan F1-score 90,63%. Hasil *confusion matrix* menunjukkan bahwa kesalahan klasifikasi terutama terjadi antara kelas netral dan negatif. Analisis lintas isu menunjukkan bahwa kecenderungan sentimen berbeda pada isu-isu terkait pelaksanaan Program MBG. Hasil penelitian menunjukkan bahwa kombinasi IndoBERT dan SMOTE dapat digunakan untuk mengidentifikasi kecenderungan sentimen masyarakat pada data komentar media sosial serta memberikan gambaran yang lebih terarah mengenai respons publik terhadap pelaksanaan Program MBG.

Kata Kunci: Analisis Sentimen, IndoBERT, Makan Bergizi Gratis, Media Sosial, SMOTE

Pendahuluan

Program Makan Bergizi Gratis (MBG) merupakan program strategis pemerintah yang diarahkan untuk mendukung pemenuhan kebutuhan gizi masyarakat, khususnya anak-anak usia sekolah, serta mendukung pembangunan sumber daya manusia. Implementasi program dalam skala luas menjadikan aspek kualitas makanan, keamanan pangan,

distribusi, logistik, pengawasan, dan koordinasi antarinstansi sebagai bagian penting dalam keberhasilan pelaksanaannya. Penelitian terdahulu menunjukkan bahwa implementasi awal MBG memiliki potensi dalam mendukung pembangunan sumber daya manusia, tetapi masih menghadapi tantangan pada aspek keamanan pangan, logistik, koordinasi regulasi, dan mekanisme pemantauan (Suprpto et al., 2025). Penelitian lainnya juga menunjukkan bahwa MBG diarahkan untuk meningkatkan akses terhadap makanan bergizi dan mendukung pembangunan sumber daya manusia, sementara implementasinya masih menghadapi persoalan fasilitas, proses penyediaan makanan, pedoman operasional, dan koordinasi pelaksanaan (Babo et al., 2026).

Media sosial berkembang menjadi sumber data yang relevan dalam analisis opini publik karena menghasilkan *user-generated content* dalam jumlah besar (Wafda, 2024). Komentar pada Instagram dan YouTube dapat memuat dukungan, kritik, pertanyaan, keluhan, maupun tanggapan informatif terhadap suatu isu. Analisis sentimen memungkinkan data tekstual tersebut diklasifikasikan ke dalam kategori tertentu sehingga kecenderungan opini masyarakat dapat dianalisis secara kuantitatif. Penggunaan data media sosial untuk analisis sentimen juga telah berkembang pada berbagai isu di Indonesia, termasuk isu kebijakan publik dan opini masyarakat terhadap program pemerintah.

Komentar YouTube memiliki karakteristik teks yang tidak terstruktur karena mengandung variasi bahasa, *stopword*, slang, dan bentuk ekspresi informal. Penelitian terdahulu menunjukkan bahwa karakteristik tersebut menjadi salah satu tantangan dalam analisis sentimen komentar YouTube berbahasa Indonesia sehingga diperlukan pengolahan teks sebelum proses klasifikasi (Aribowo et al., 2021). Penelitian tersebut secara khusus mengembangkan model analisis sentimen lintas domain pada komentar YouTube Indonesia dan menegaskan pentingnya penanganan karakteristik bahasa informal dalam data media sosial.

Permasalahan serupa muncul pada data komentar Instagram dan YouTube mengenai Program MBG yang digunakan dalam penelitian ini. Data komentar yang diperoleh melalui proses *web scraping* memiliki variasi bentuk penulisan sehingga tidak dapat langsung digunakan sebagai masukan model. Oleh karena itu, diperlukan tahapan *text preprocessing* untuk mengurangi *noise* dan menyeragamkan bentuk data. Penelitian terdahulu menunjukkan bahwa data hasil *crawling* media sosial masih berbentuk mentah dan tidak terstruktur sehingga perlu melalui tahapan *preprocessing* sebelum digunakan dalam proses analisis (Rifaldi et al, 2023).

Tahapan *preprocessing* menjadi penting karena karakteristik bahasa pada media sosial berbeda dengan teks formal. Proses seperti *case folding*, *cleaning*, normalisasi, tokenisasi, *stopword removal*, dan *stemming* dapat digunakan untuk menghasilkan representasi teks yang lebih terstruktur. Namun, setiap tahapan *preprocessing* juga perlu diterapkan secara selektif karena penghapusan kata tertentu berpotensi menghilangkan informasi yang berkaitan dengan sentiment (Salu et al., 2026). Penelitian terdahulu menunjukkan bahwa penerapan *stopword removal* dan SMOTE memberikan pengaruh berbeda terhadap klasifikasi sentimen Bahasa Indonesia (Negara, 2025). Hasil penelitian tersebut menunjukkan bahwa SMOTE memberikan peningkatan performa pada model tertentu, sedangkan *stopword removal* justru dapat menurunkan performa karena berpotensi menghilangkan informasi sentimen.

Pemilihan IndoBERT dalam penelitian ini berkaitan dengan kebutuhan untuk memperoleh representasi kontekstual Bahasa Indonesia. Penelitian terdahulu mengembangkan IndoBERT sebagai model bahasa *pre-trained* untuk Bahasa Indonesia dan mengevaluasinya pada berbagai tugas dalam IndoLEM (Koto et al., 2020). Hasil penelitian menunjukkan bahwa IndoBERT mencapai performa *state-of-the-art* pada sebagian besar tugas yang digunakan.

Kemampuan IndoBERT pada analisis sentimen Bahasa Indonesia juga ditunjukkan oleh penelitian tersebut menggunakan IndoBERT dalam kombinasi dengan beberapa arsitektur *deep learning* untuk tugas analisis sentimen dan klasifikasi emosi (Ahmadian et al., 2024). Model yang dikembangkan memperoleh akurasi 93% pada tugas analisis sentimen di dataset IndoNLU. Temuan tersebut memperkuat penggunaan IndoBERT sebagai model yang relevan untuk klasifikasi teks Bahasa Indonesia. Penggunaan IndoBERT pada isu kebijakan publik di Indonesia juga telah diterapkan pada beberapa penelitian. Penelitian terdahulu menggunakan IndoBERT untuk menganalisis sentimen publik terhadap kebijakan kenaikan PPN 12% dengan data dari media social (Manoppo et al., 2025). Penelitian tersebut menunjukkan penerapan model *transformer* Bahasa Indonesia untuk mengidentifikasi kecenderungan sentimen masyarakat terhadap isu kebijakan.

Dalam konteks Program MBG, penelitian terdahulu menggunakan IndoBERT untuk menganalisis sentimen masyarakat pada platform X (Muhabbab et al., 2025). Penelitian tersebut menggunakan 7.958 unggahan berbahasa Indonesia dan memperoleh akurasi 92% serta *macro F1-score* sebesar 0,90. Hasil tersebut menunjukkan bahwa IndoBERT memiliki kemampuan yang baik dalam mengklasifikasikan respons masyarakat terhadap isu MBG. Meskipun demikian, penelitian tersebut menggunakan platform X, sedangkan penelitian ini menggunakan komentar Instagram dan YouTube sehingga memiliki karakteristik sumber data dan pola interaksi pengguna yang berbeda. Selain persoalan representasi bahasa, ketidakseimbangan kelas merupakan permasalahan penting dalam klasifikasi sentimen. Distribusi data yang tidak seimbang dapat menyebabkan model lebih dominan mempelajari kelas mayoritas dan menghasilkan kemampuan prediksi yang kurang optimal pada kelas minoritas. Permasalahan tersebut menjadi relevan dalam penelitian ini karena distribusi sentimen awal menunjukkan jumlah komentar netral yang jauh lebih besar dibandingkan komentar positif dan negatif.

SMOTE digunakan untuk mengatasi ketidakseimbangan tersebut. Penelitian terdahulu memperkenalkan *Synthetic Minority Over-sampling Technique* sebagai metode untuk menghasilkan sampel sintesis pada kelas minoritas berdasarkan hubungan dengan *nearest neighbor* (Chawla et al., 2002). Penggunaan SMOTE pada klasifikasi data Bahasa Indonesia juga telah diteliti. Penelitian terdahulu menunjukkan bahwa penerapan SMOTE memberikan pengaruh positif terhadap performa klasifikasi pada eksperimen sentiment analysis Bahasa Indonesia (Negara, 2025). Penelitian lainnya juga menerapkan SMOTE pada analisis sentimen isu resesi di Indonesia dan menunjukkan bahwa model dengan data yang telah diseimbangkan menghasilkan evaluasi yang lebih baik dibandingkan model tanpa penyeimbangan pada beberapa scenario (Kristiyanti et al., 2024).

Kombinasi IndoBERT dan SMOTE dalam penelitian ini digunakan untuk menangani dua karakteristik utama dataset, yaitu kompleksitas teks Bahasa Indonesia dan ketidakseimbangan distribusi kelas. IndoBERT digunakan untuk memperoleh representasi kontekstual teks, sedangkan SMOTE digunakan untuk memperbaiki distribusi kelas sebelum proses pembelajaran. Pendekatan tersebut berbeda fungsi sehingga

penerapannya diarahkan untuk mengatasi permasalahan yang berbeda dalam dataset penelitian.

Penelitian terdahulu terkait MBG menunjukkan bahwa hasil klasifikasi dapat berbeda berdasarkan platform, jumlah data, metode pelabelan, dan konfigurasi model (Muhabbab et al., 2025). Misalnya, memperoleh akurasi 92% pada data platform X dengan IndoBERT, sedangkan penelitian mengenai isu kebijakan publik lainnya menunjukkan variasi performa ketika karakteristik data dan skenario eksperimen berbeda. Hal tersebut menunjukkan bahwa performa model tidak hanya ditentukan oleh arsitektur yang digunakan, tetapi juga dipengaruhi oleh karakteristik dataset, proses preprocessing, distribusi kelas, dan sumber data.

Berdasarkan kondisi tersebut, penelitian ini diarahkan pada tiga aspek utama. Pertama, menggunakan komentar Instagram dan YouTube sebagai sumber data untuk memperoleh respons masyarakat dari dua platform dengan karakteristik interaksi yang berbeda. Kedua, mengombinasikan IndoBERT dengan SMOTE untuk menangani kompleksitas representasi Bahasa Indonesia dan ketidakseimbangan kelas. Ketiga, melakukan analisis sentimen lintas isu untuk mengidentifikasi variasi kecenderungan sentimen masyarakat pada aspek-aspek yang berkaitan dengan Program MBG. Pendekatan tersebut diharapkan memberikan analisis yang lebih spesifik dibandingkan hanya melihat distribusi sentimen secara keseluruhan.

Tujuan penelitian ini adalah menerapkan *text mining* pada komentar Instagram dan YouTube mengenai Program MBG menggunakan IndoBERT dan SMOTE, mengevaluasi kemampuan model dalam mengklasifikasikan sentimen positif, netral, dan negatif, serta menganalisis kecenderungan sentimen masyarakat berdasarkan isu yang berkaitan dengan Program MBG. Kebaharuan penelitian ini terletak pada integrasi analisis sentimen lintas platform Instagram dan YouTube serta lintas isu terhadap Program Makan Bergizi Gratis menggunakan model IndoBERT yang diperkuat dengan teknik SMOTE untuk menangani ketidakseimbangan data, sehingga menghasilkan analisis persepsi masyarakat yang lebih komprehensif dan representatif

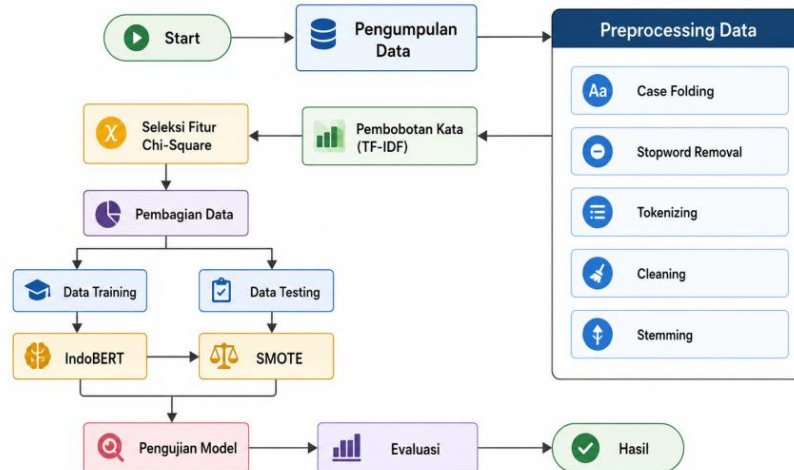
Metode

Desain Penelitian

Penelitian menggunakan pendekatan kuantitatif eksperimental dengan penerapan model *deep learning* IndoBERT dan teknik SMOTE pada data komentar media sosial. Desain eksperimental digunakan karena penelitian melibatkan serangkaian tahap pengolahan data, pembentukan dataset, penyeimbangan kelas, pelatihan model, pengujian, dan evaluasi performa. Rangkaian tahapan tersebut sejalan dengan penelitian sentiment analysis berbasis IndoBERT yang menerapkan *data preprocessing*, pembentukan dataset, pengembangan model, pelatihan, dan evaluasi menggunakan *accuracy*, *precision*, *recall*, dan F1-score pada data media sosial berbahasa Indonesia (Adhantoro et al., 2026; Medantoro & Muljono, 2026).

Penerapan teknik SMOTE dalam penelitian ini digunakan sebagai bagian dari proses penyeimbangan kelas sebelum pelatihan model. Pendekatan tersebut sejalan dengan penelitian terdahulu yang mengintegrasikan SMOTE dengan representasi berbasis IndoBERT pada klasifikasi sentimen Bahasa Indonesia dan mengevaluasi model menggunakan *accuracy*, *precision*, *recall*, serta macro F1-score (Roihan et al., 2026).

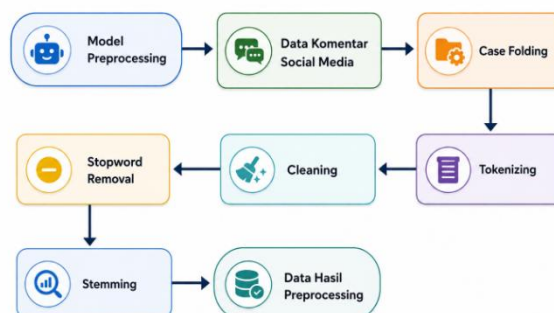
Gambar 1 tersebut menunjukkan alur penelitian yang dimulai dari pengumpulan data, kemudian dilakukan preprocessing data yang meliputi *case folding*, *stopword removal*, *tokenizing*, *cleaning*, dan *stemming*. Selanjutnya, data melalui tahap pembobotan kata menggunakan TF-IDF dan seleksi fitur menggunakan Chi-Square. Data yang telah diproses kemudian dibagi menjadi data training dan data testing. Pada tahap pemodelan, IndoBERT digunakan untuk proses klasifikasi, sedangkan SMOTE diterapkan untuk menangani ketidakseimbangan kelas pada data. Selanjutnya dilakukan pengujian model dan evaluasi menggunakan metrik evaluasi yang telah ditentukan. Tahap akhir dari penelitian adalah memperoleh hasil klasifikasi dan kinerja model.



Gambar 1. Alur Penelitian

Sumber dan Pengumpulan Data

Data primer berupa komentar berbahasa Indonesia dari Instagram dan YouTube yang dikumpulkan melalui *web scraping* menggunakan Python dan Google Colab. Data berasal dari konten berita, akun resmi pemerintah, dan opini publik terkait MBG, sebagaimana pendekatan pengumpulan komentar YouTube yang digunakan (Yulita et al., 2023). Komentar kemudian digunakan sebagai unit analisis dan diproses melalui *preprocessing*, mengingat komentar media sosial memiliki variasi bahasa dan *slang* (Aribowo et al., 2021).



Gambar 2. Tahapan Preprocessing

Berdasarkan tahapan preprocessing pada gambar 2 tersebut, *preprocessing* berfungsi mengubah komentar media sosial yang masih bersifat tidak terstruktur menjadi data teks yang lebih seragam dan siap digunakan dalam proses analisis. *Case folding* dan *cleaning* digunakan untuk menyeragamkan serta mengurangi *noise*, *tokenizing* mengubah teks menjadi unit-unit token, *stopword removal* mengurangi kata yang kurang informatif, sedangkan *stemming* menyeragamkan bentuk kata ke bentuk dasarnya. Data yang telah

melalui seluruh tahapan tersebut kemudian digunakan sebagai masukan pada proses pelabelan sentimen dan klasifikasi menggunakan IndoBERT.

Pelabelan Sentimen

Pelabelan dilakukan menggunakan pendekatan *rule-based* dengan mengelompokkan komentar ke dalam tiga kategori, yaitu positif, netral, dan negatif. Pendekatan berbasis leksikon dapat digunakan untuk menghasilkan label sentimen dalam jumlah besar dengan memanfaatkan kata-kata bermuatan sentimen, termasuk variasi *slang* dan ekspresi informal (Anaissi et al., 2021). Pendekatan ini efisien untuk pelabelan awal, tetapi masih memiliki keterbatasan dalam memahami konteks, ironi, dan sarkasme. Karena itu, hasil pelabelan digunakan sebagai dasar pembentukan dataset sebelum proses pelatihan IndoBERT.

Penyeimbangan Data Menggunakan SMOTE

SMOTE diterapkan untuk menangani ketidakseimbangan kelas dengan menghasilkan sampel sintesis pada kelas minoritas. Penerapan SMOTE bersama IndoBERT pada data media sosial berbahasa Indonesia menunjukkan relevansinya dalam menangani distribusi kelas yang tidak seimbang (Ihalauw et al., 2025). Pendekatan serupa juga diterapkan pada sentiment analysis untuk meningkatkan representasi kelas minoritas dalam proses klasifikasi (Badriyah et al., 2025).

Model IndoBERT

IndoBERT digunakan sebagai model utama untuk klasifikasi sentimen karena merupakan model *pre-trained* yang dikembangkan untuk Bahasa Indonesia dan mampu menghasilkan representasi teks berbasis konteks. Penelitian terdahulu mengembangkan IndoBERT sebagai model bahasa Bahasa Indonesia (Koto et al., 2020). Sedangkan penelitian lainnya menunjukkan penerapan *fine-tuning* IndoBERT pada *user-generated content* berbahasa Indonesia untuk klasifikasi sentiment (Perwira et al., 2025). Dalam penelitian ini, IndoBERT digunakan untuk mengklasifikasikan komentar ke dalam tiga kategori sentimen melalui proses *fine-tuning* pada dataset penelitian. Pendekatan *fine-tuning* memungkinkan model *pre-trained* disesuaikan dengan karakteristik dataset dan tugas klasifikasi yang digunakan.

Pembagian Data dan Pelatihan

Data yang telah diproses dan diseimbangkan dibagi menjadi data *training* dan *testing*. Data *training* digunakan untuk melatih model, sedangkan data *testing* digunakan untuk mengevaluasi kemampuan generalisasi pada data yang tidak digunakan selama pelatihan. Pembagian data dan proses *fine-tuning* IndoBERT merupakan pendekatan yang digunakan dalam klasifikasi sentimen Bahasa Indonesia (Aprilah et al., 2025; Perwira et al., 2025). Pelatihan dilakukan menggunakan jumlah *epoch* yang telah ditentukan, sedangkan nilai *loss* diamati selama proses pembelajaran untuk memantau perkembangan model.

Evaluasi Model

Evaluasi dilakukan menggunakan *accuracy*, *precision*, *recall*, F1-score, dan *confusion matrix*. *Accuracy* mengukur proporsi prediksi yang benar, *precision* menunjukkan ketepatan prediksi suatu kelas, *recall* mengukur kemampuan model mengenali anggota kelas, sedangkan F1-score merupakan gabungan *precision* dan *recall*. *Confusion matrix* digunakan untuk melihat pola kesalahan klasifikasi antar kelas. Penggunaan kombinasi

metrik tersebut juga diterapkan pada evaluasi model IndoBERT untuk klasifikasi sentimen Bahasa Indonesia (Singgalen, 2025).

Analisis Sentimen Lintas Isu

Analisis lintas isu dilakukan setelah model menghasilkan prediksi sentimen. Komentar kemudian dikelompokkan berdasarkan isu terkait Program MBG, lalu distribusi sentimen positif, netral, dan negatif dihitung pada setiap isu. Pendekatan berbasis aspek memungkinkan sentimen dianalisis secara lebih spesifik berdasarkan aspek yang dibicarakan dalam teks (Perwira et al., 2025). Analisis ini digunakan untuk membandingkan variasi respons masyarakat antarisu, sehingga informasi yang diperoleh tidak hanya menggambarkan distribusi sentimen secara keseluruhan, tetapi juga menunjukkan isu yang memiliki kecenderungan sentimen berbeda.

Hasil

Hasil Pengumpulan Data

Data penelitian dikumpulkan dari Instagram dan YouTube menggunakan teknik *web scraping*. Sebanyak 52.326 komentar berbahasa Indonesia diperoleh dari konten yang berkaitan dengan Program MBG. Data mencakup beberapa atribut utama berupa platform, nama pengguna, URL sumber, dan isi komentar. Pengumpulan data dari dua platform dilakukan untuk memperoleh representasi komentar yang lebih beragam terkait respons masyarakat terhadap Program MBG.

```

=== MULAI PROSES SCRAPING YOUTUBE ===

📺 Mengambil komentar dari Video ID: EEmfTnEhVJ0
📺 Mengambil komentar dari Video ID: za8EqDmgeiM
📺 Mengambil komentar dari Video ID: KLC1Jcl9HlQ
📺 Mengambil komentar dari Video ID: D_atwzv8Zy4
✅ Total 8450 komentar YouTube berhasil diambil.

✅ Data berhasil dibuat menjadi DataFrame!
Total komentar: 8450
📁 Data berhasil disimpan ke 'youtube_comments.csv'
📁 Total komentar: 19196 disimpan di 'komentar_semua_reels.csv'
📁 Total komentar: 15471 disimpan di 'komentar_semua_reels.csv'
📁 Total komentar: 9202 disimpan di 'komentar_semua_reels.csv'

```

Gambar 3. Data Komentar Hasil Scraping

Gambar 3 memperlihatkan bentuk data mentah yang diperoleh dari kedua platform. Struktur data tersebut menjadi dasar untuk tahap *preprocessing* karena komentar masih mengandung variasi penulisan dan elemen yang tidak diperlukan untuk klasifikasi. Kondisi data mentah tersebut menunjukkan perlunya proses pembersihan dan normalisasi sebelum dilakukan analisis lebih lanjut.

Hasil Preprocessing

Setelah proses pengumpulan data, dilakukan *preprocessing* untuk mengubah komentar media sosial yang masih memiliki bentuk dan karakteristik yang beragam menjadi data teks yang lebih terstruktur. Tahapan *preprocessing* yang diterapkan meliputi *case folding*, *cleaning*, normalisasi kata tidak baku, tokenisasi, *stopword removal*, dan *stemming*. Selain proses transformasi teks, dilakukan pula penghapusan *missing value* dan data yang tidak memenuhi kebutuhan analisis. Dari proses tersebut, jumlah komentar yang digunakan dalam tahap klasifikasi menjadi 47.861 komentar.

	Teks Asli	Case Folding	Cleaning	Normalisasi	Tokenisasi	Stopword Removal	Stemming	Hasil Akhir
0	Komentar	komentar	komentar	komentar	komentar	komentar	komentar	komentar
1	tergantung yg ngelola nya	tergantung yg ngelola nya	tergantung yg ngelola nya	tergantung yang ngelola nya	tergantung yang ngelola nya	tergantung ngelola nya	gantung ngelola nya	gantung ngelola nya
2	mbg emang wajar untuk di kritik	mbg emang wajar untuk di kritik	mbg emang wajar untuk di kritik	makan bergizi gratis emang wajar untuk di kritik	makan bergizi gratis emang wajar untuk di kritik	makan bergizi gratis emang wajar kritik	makan gizi gratis emang wajar kritik	makan gizi gratis emang wajar kritik
3	"Nanti di bilang "kami kasi makanan bergizi bukan makanan kenyang" padahal bergizi juga kagak 🤔"	"nanti di bilang "kami kasi makanan bergizi bukan makanan kenyang" padahal bergizi juga kagak 🤔"	nanti di bilang kami kasi makanan bergizi bukan makanan kenyang padahal bergizi juga kagak	nanti di bilang kami kasi makanan bergizi bukan makanan kenyang padahal bergizi juga kagak	nanti di bilang kami kasi makanan bergizi bukan makanan kenyang padahal bergizi juga kagak	bilang kasi makanan bergizi makanan kenyang bergizi kagak	bilang kasi makan gizi makan kenyang gizi kagak	bilang kasi makan gizi makan kenyang gizi kagak
4	Gaenak mbgnya enakan ceker ayam bikinan mak gwehj 🤔	gaenak mbgnya enakan ceker ayam bikinan mak gwehj 🤔	gaenak mbgnya enakan ceker ayam bikinan mak gwehj	gaenak mbgnya enakan ceker ayam bikinan mak gwehj	gaenak mbgnya enakan ceker ayam bikinan mak gwehj	gaenak mbgnya enakan ceker ayam bikinan mak gwehj	gaenak mbgnya enakan ceker ayam bikin mak gwehj	gaenak mbgnya enakan ceker ayam bikin mak gwehj
5	"Yang bilang" "bersyukur" selamat anda udh kena propaganda koruptor,kalian gak butuh transparansi berarti"	"yang bilang" "bersyukur" selamat anda udh kena propaganda koruptor,kalian gak butuh transparansi berarti"	yang bilang bersyukur selamat anda udh kena propaganda koruptorkalian gak butuh transparansi berarti	yang bilang bersyukur selamat anda sudah kena propaganda koruptorkalian tidak butuh transparansi berarti	yang bilang bersyukur selamat anda sudah kena propaganda koruptorkalian tidak butuh transparansi berarti	bilang bersyukur selamat kena propaganda koruptorkalian butuh transparansi	bilang syukur selamat kena propaganda koruptorkalian butuh transparansi	bilang syukur selamat kena propaganda koruptorkalian butuh transparansi

Gambar 4. Hasil Akhir Preprocessing

Gambar 4 memperlihatkan perubahan komentar dari teks asli hingga menjadi hasil akhir *preprocessing*. Data awal masih mengandung variasi huruf kapital, kata tidak baku, tanda baca, tanda kutip, dan emoji sehingga perlu dilakukan pengolahan sebelum digunakan dalam analisis sentimen. Tahap *case folding* mengubah seluruh huruf menjadi *lowercase* untuk menyeragamkan bentuk penulisan. Selanjutnya, *cleaning* menghapus karakter atau elemen yang tidak diperlukan sehingga *noise* pada teks dapat dikurangi. Tahap normalisasi menyeragamkan kata tidak baku, singkatan, dan bentuk bahasa informal agar lebih konsisten.

Setelah itu, tokenisasi memecah komentar menjadi unit-unit kata (*token*) yang dapat diproses secara individual. Tahap *stopword removal* menghilangkan kata-kata yang dianggap kurang informatif, sedangkan *stemming* mengubah kata berimbuhan menjadi bentuk dasarnya untuk mengurangi variasi kata. Berdasarkan Gambar 10, seluruh tahapan tersebut menghasilkan teks yang lebih bersih, seragam, dan terstruktur. Hasil akhir *preprocessing* selanjutnya digunakan sebagai masukan untuk proses pelabelan sentimen dan klasifikasi menggunakan IndoBERT.

Distribusi Sentimen Awal

Setelah proses *preprocessing*, data kemudian melalui tahap pelabelan sentimen menggunakan pendekatan *rule-based*. Pelabelan menghasilkan tiga kategori sentimen, yaitu positif, netral, dan negatif. Berdasarkan hasil pelabelan, diperoleh 3.131 komentar positif, 31.136 komentar netral, dan 4.025 komentar negatif. Distribusi tersebut menunjukkan bahwa jumlah komentar pada setiap kategori belum seimbang, dengan kelas netral memiliki jumlah data yang jauh lebih besar dibandingkan kelas positif dan negatif. Hasil tersebut sesuai dengan temuan pada penelitian yang menunjukkan bahwa distribusi data awal mengalami ketidakseimbangan sebelum dilakukan proses penyeimbangan menggunakan SMOTE.

Tabel 1. Distribusi Sentimen Awal

Kategori Sentimen	Jumlah
Positif	3.131
Netral	31.136
Negatif	4.025

Berdasarkan Tabel 1, sentimen netral mendominasi dengan 31.136 komentar (81,31%), diikuti negatif 4.025 komentar (10,51%) dan positif 3.131 komentar (8,18%). Distribusi tersebut menunjukkan ketidakseimbangan kelas yang cukup besar, sehingga diperlukan

penyeimbangan data sebelum pelatihan model. SMOTE kemudian diterapkan untuk meningkatkan representasi kelas minoritas. Setelah proses penyeimbangan, masing-masing kelas memiliki 31.136 data, sehingga distribusi data menjadi seimbang sebelum digunakan dalam pelatihan IndoBERT.

Hasil Penyeimbangan Data

Distribusi awal menunjukkan ketidakseimbangan kelas, dengan 31.136 komentar netral, 4.025 negatif, dan 3.131 positif. Oleh karena itu, diterapkan Synthetic Minority Over-sampling Technique (SMOTE) sebelum pelatihan IndoBERT. Penerapan SMOTE bertujuan meningkatkan keterwakilan kelas minoritas dengan menghasilkan sampel sintetis berdasarkan karakteristik data yang tersedia. Proses ini diharapkan dapat mengurangi kecenderungan model untuk lebih dominan memprediksi kelas mayoritas. Dengan distribusi data yang lebih seimbang, proses pelatihan IndoBERT dapat memberikan hasil klasifikasi yang lebih representatif pada setiap kategori sentimen.

```
Sebelum SMOTE: Counter({1: 31136, 2: 4025, 0: 3131})
Sesudah SMOTE: Counter({1: 31136, 0: 31136, 2: 31136})
```

Gambar 5. Perbandingan Distribusi Data Sebelum dan Sesudah SMOTE

Gambar 5 menunjukkan bahwa sebelum SMOTE, kelas netral merupakan kelas mayoritas. Setelah SMOTE, jumlah data pada setiap kelas diseragamkan menjadi 31.136 data, sehingga distribusi ketiga kelas menjadi seimbang. Penyeimbangan dilakukan dengan menghasilkan data sintetis pada kelas minoritas hingga jumlahnya setara dengan kelas mayoritas. Hasil ini mengurangi ketimpangan distribusi data dan memberikan proporsi yang lebih seimbang bagi kelas positif, netral, dan negatif dalam proses pembelajaran IndoBERT.

Hasil Pelatihan IndoBERT

Setelah penyeimbangan data menggunakan SMOTE, model IndoBERT dilatih selama 3 epoch untuk mengklasifikasikan tiga kategori sentimen. Perkembangan proses pembelajaran diamati melalui nilai *training loss* pada setiap epoch.

```
Epoch 1/3 - Loss: 0.1256
Epoch 2/3 - Loss: 0.0656
Epoch 3/3 - Loss: 0.0493
```

Gambar 6. Nilai Training Loss Model IndoBERT pada Setiap Epoch

Berdasarkan Gambar 6, nilai *loss* mengalami penurunan secara bertahap selama proses pelatihan. Pada epoch pertama, nilai *loss* tercatat sebesar 0,1256, kemudian menurun menjadi 0,0656 pada epoch kedua, dan kembali menurun menjadi 0,0493 pada epoch ketiga. Penurunan nilai tersebut menunjukkan bahwa kesalahan model terhadap data pelatihan semakin kecil seiring bertambahnya proses pembelajaran.

Secara keseluruhan, nilai *loss* mengalami penurunan sebesar 0,0763, atau sekitar 60,75%, dari epoch pertama hingga epoch ketiga. Penurunan tersebut menunjukkan bahwa model mampu mempelajari pola yang terdapat pada data pelatihan secara bertahap. Nilai *loss* terendah diperoleh pada epoch ketiga, yaitu sebesar 0,0493, sehingga epoch tersebut menunjukkan hasil pembelajaran dengan tingkat kesalahan paling rendah selama proses training. Hasil ini menjadi dasar bahwa model telah melalui proses pembelajaran sebelum

digunakan pada tahap pengujian. Namun, penurunan *training loss* tidak digunakan sebagai satu-satunya dasar untuk menentukan performa akhir model. Kemampuan generalisasi model selanjutnya dievaluasi menggunakan data pengujian melalui metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *confusion matrix*.

Hasil Evaluasi Model

Setelah proses pelatihan, model IndoBERT diuji menggunakan data pengujian untuk mengetahui kemampuan model dalam mengklasifikasikan komentar ke dalam kategori sentimen. Evaluasi dilakukan menggunakan beberapa metrik, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*. Selain metrik agregat, *classification report* digunakan untuk melihat performa model pada masing-masing kelas sentimen secara lebih terperinci. Pengujian dilakukan terhadap 9.573 data uji.

Tabel 2. Hasil Evaluasi IndoBERT dan SMOTE

Metrik	Nilai
Accuracy	89,88%
Precision	92,55%
Recall	89,88%
F1-Score	90,63%

Berdasarkan Tabel 2, model memperoleh *accuracy* sebesar 89,88%, yang menunjukkan proporsi prediksi yang berhasil diklasifikasikan dengan benar. *Precision* sebesar 92,55% menunjukkan tingkat ketepatan prediksi model secara agregat, sedangkan *recall* sebesar 89,88% menunjukkan kemampuan model dalam mengenali data sesuai kelas aktual. Nilai *F1-score* sebesar 90,63% menunjukkan keseimbangan antara *precision* dan *recall*.

=== HASIL EVALUASI MODEL ===				
	precision	recall	f1-score	support
1	0.83	0.96	0.89	783
3	0.98	0.89	0.94	7784
5	0.57	0.89	0.69	1006
accuracy			0.90	9573
macro avg	0.79	0.91	0.84	9573
weighted avg	0.93	0.90	0.91	9573
Accuracy : 0.8987778125979317				
F1-score : 0.9063444496962423				

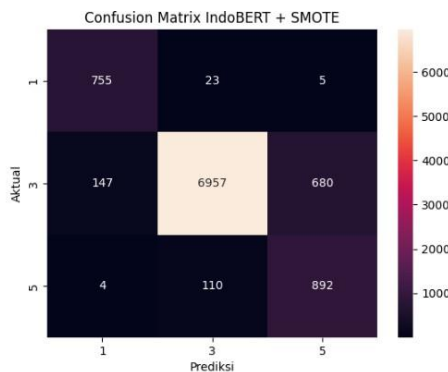
Gambar 7. Classification Report

Gambar 7 menunjukkan performa model pada masing-masing kelas. Kelas 1 memperoleh *precision* 0,83, *recall* 0,96, dan *F1-score* 0,89. Kelas 3 memperoleh *precision* tertinggi sebesar 0,98, *recall* 0,89, dan *F1-score* 0,94. Sementara itu, kelas 5 memperoleh *precision* 0,57, *recall* 0,89, dan *F1-score* 0,69. Perbedaan tersebut menunjukkan bahwa kemampuan model tidak sepenuhnya sama pada setiap kelas, terutama pada *precision* kelas 5 yang masih lebih rendah.

Secara keseluruhan, *macro average* menghasilkan *precision* 0,79, *recall* 0,91, dan *F1-score* 0,84, sedangkan *weighted average* masing-masing sebesar 0,93, 0,90, dan 0,91. Nilai *weighted average* yang lebih tinggi menunjukkan adanya pengaruh jumlah data pada masing-masing kelas terhadap performa agregat. Hasil evaluasi tersebut menunjukkan bahwa model memiliki kinerja klasifikasi yang baik, tetapi masih terdapat perbedaan performa antar kelas. Oleh karena itu, analisis *confusion matrix* digunakan untuk mengidentifikasi pola kesalahan klasifikasi secara lebih terperinci.

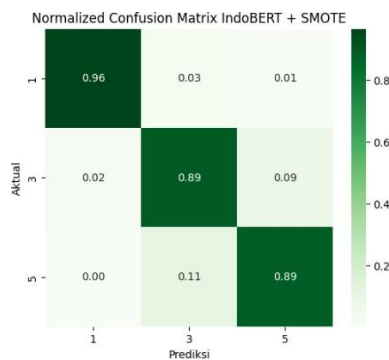
Analisis Confusion Matrix

Evaluasi menggunakan *confusion matrix* dilakukan untuk melihat ketepatan prediksi pada setiap kelas dan pola kesalahan klasifikasi. *Normalized confusion matrix* digunakan untuk menunjukkan proporsi hasil prediksi dalam bentuk persentase.



Gambar 8. Confusion Matrix Model IndoBERT + SMOTE

Berdasarkan Gambar 8, model mengklasifikasikan 755 dari 783 data kelas 1, 6.957 dari 7.784 data kelas 3, dan 892 dari 1.006 data kelas 5 secara benar. Dari total 9.573 data uji, sebanyak 8.604 data berhasil diklasifikasikan dengan benar dan 969 data mengalami kesalahan, sehingga konsisten dengan *accuracy* sebesar 89,88%.



Gambar 9. Normalized Confusion Matrix Model IndoBERT + SMOTE

Gambar 9 menunjukkan tingkat prediksi benar sebesar 96% pada kelas 1, serta masing-masing 89% pada kelas 3 dan kelas 5. Kesalahan terbesar terjadi antara kelas 3 dan kelas 5, yaitu 680 data kelas 3 diprediksi sebagai kelas 5 dan 110 data kelas 5 diprediksi sebagai kelas 3. Secara proporsional, kesalahan tersebut masing-masing mencapai sekitar 9% dan 11%.

Pola tersebut menunjukkan adanya kemiripan karakteristik teks antara kelas 3 dan kelas 5, terutama pada komentar yang menggunakan bahasa informal, kritik ringan, keluhan, atau ungkapan dengan polaritas yang tidak eksplisit. Dengan demikian, meskipun model menghasilkan tingkat klasifikasi yang baik secara keseluruhan, kelas 3 dan kelas 5 masih menjadi sumber kesalahan utama. Hasil ini menegaskan bahwa evaluasi model perlu mempertimbangkan distribusi kesalahan pada setiap kelas, tidak hanya *accuracy* secara keseluruhan.

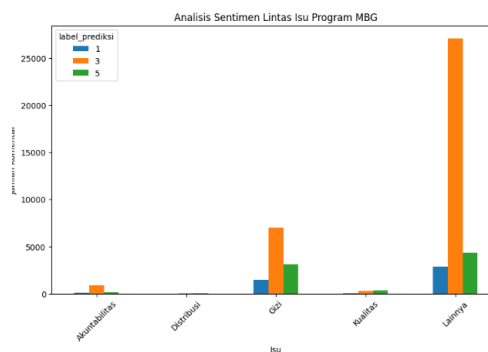
Analisis Sentimen Lintas Isu

Analisis sentimen lintas isu dilakukan untuk melihat distribusi sentimen pada setiap aspek Program Makan Bergizi Gratis (MBG). Komentar diprediksi menggunakan IndoBERT, kemudian dikelompokkan berdasarkan isu dan label sentimen.

Tabel 3. Rekap Sentimen Berdasarkan Isu

Label prediksi Isu	1	3	5
Akuntabilitas	130	892	165
Distribusi	12	71	24
Gizi	1473	6986	3134
Kualitas	53	304	340
Lainnya	2854	27095	4332

Berdasarkan Tabel 3, label 3 mendominasi pada isu Akuntabilitas (892), Distribusi (71), Gizi (6.986), dan Lainnya (27.095). Pada isu Kualitas, label 5 menjadi kategori terbesar dengan 340 komentar, diikuti label 3 sebanyak 304 dan label 1 sebanyak 53 komentar. Isu Gizi dan Lainnya memiliki volume komentar terbesar, sedangkan Distribusi memiliki jumlah paling sedikit. Dominasi kategori Lainnya menunjukkan bahwa sebagian besar komentar berada di luar empat isu utama yang ditentukan. Sementara itu, tingginya jumlah komentar pada isu Gizi menunjukkan bahwa aspek tersebut menjadi salah satu topik yang banyak dibicarakan dalam konteks MBG.

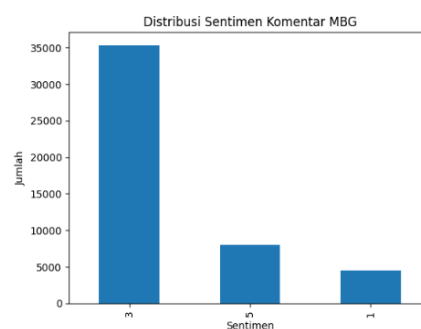


Gambar 10. Visualisasi Sentimen Lintas Isu

Gambar 10 memperlihatkan perbedaan distribusi label pada setiap isu. Hasil ini menunjukkan bahwa respons masyarakat terhadap MBG tidak seragam pada seluruh aspek, sehingga distribusi sentimen secara keseluruhan perlu dilengkapi dengan analisis berdasarkan isu. Analisis lintas isu memberikan gambaran mengenai aspek yang paling banyak memperoleh perhatian serta perbedaan kecenderungan sentimen pada masing-masing isu.

Distribusi Sentimen Keseluruhan

Setelah dilakukan prediksi terhadap seluruh komentar, distribusi sentimen secara keseluruhan dianalisis untuk mengetahui kecenderungan respons masyarakat terhadap Program Makan Bergizi Gratis (MBG). Hasil distribusi diperoleh dari keseluruhan komentar yang telah diproses dan diklasifikasikan menggunakan model IndoBERT.



Gambar 11. Distribusi Sentimen Komentar MBG

Berdasarkan Gambar 11, label 3 merupakan kategori dominan dengan 35.348 komentar (73,86%), diikuti label 5 sebanyak 7.995 komentar (16,71%) dan label 1 sebanyak 4.522 komentar (9,44%). Distribusi tersebut menunjukkan dominasi yang cukup kuat pada label 3 dibandingkan dua kategori lainnya. Hal ini menunjukkan bahwa sebagian besar komentar dalam data penelitian terkonsentrasi pada kategori yang direpresentasikan oleh label 3.

Hasil ini memberikan gambaran umum mengenai kecenderungan sentimen pada seluruh data. Namun, distribusi agregat belum menunjukkan isu yang melatarbelakangi munculnya setiap sentimen sehingga perlu diinterpretasikan bersama hasil analisis sentimen lintas isu. Dengan demikian, gambar 20 melengkapi analisis sebelumnya dengan memberikan gambaran distribusi sentimen secara keseluruhan. Analisis ini menjadi dasar untuk memahami pola sentimen sebelum dikaitkan dengan isu-isu spesifik yang berkembang dalam pembahasan Program MBG.

Analisis Wordcloud

Wordcloud digunakan sebagai analisis pendukung untuk melihat kata yang paling sering muncul pada masing-masing kategori sentimen hasil klasifikasi IndoBERT. Ukuran kata menunjukkan frekuensi kemunculannya dalam kategori tersebut.



Gambar 12. Wordcloud Sentimen Positif

Gambar 12 menunjukkan dominasi kata seperti *program*, *makan*, *gizi*, *bagus*, *anak*, *sekolah*, *enak*, *sehat*, dan *manfaat*. Kata-kata tersebut menunjukkan bahwa komentar pada kategori positif banyak berkaitan dengan manfaat, kesehatan, makanan, serta penerimaan terhadap Program MBG.



Gambar 13. Wordcloud Sentimen Netral

Gambar 13 didominasi kata seperti *gizi*, *makan*, *program*, *pemerintah*, *anak*, *sekolah*, *kritik*, *salah*, *basi*, *racun*, *kurang*, dan *korupsi*. Kemunculan kata tersebut menunjukkan adanya pembahasan mengenai kritik, permasalahan, dan evaluasi terhadap pelaksanaan Program MBG.



Gambar 14. Wordcloud Sentimen Negatif

Gambar 14 didominasi kata *makan, gizi, program, pemerintah, anak, sekolah, masyarakat, gratis, dan masak*. Kata-kata tersebut lebih banyak menggambarkan objek dan informasi terkait MBG tanpa menunjukkan kecenderungan evaluatif yang kuat.

Secara keseluruhan, *wordcloud* menunjukkan perbedaan karakteristik leksikal antar kategori. Kategori positif cenderung memuat kata terkait manfaat dan apresiasi, kategori negatif berkaitan dengan kritik dan permasalahan, sedangkan kategori netral lebih bersifat informatif. *Wordcloud* digunakan sebagai analisis pendukung, bukan sebagai dasar penentuan label sentimen.

Pembahasan

Hasil penelitian menunjukkan bahwa kombinasi IndoBERT dan SMOTE menghasilkan *accuracy* sebesar 89,88%, *precision* 92,55%, *recall* 89,88%, dan F1-score 90,63%. Nilai tersebut menunjukkan bahwa model mampu melakukan klasifikasi sentimen pada komentar berbahasa Indonesia dengan kinerja yang baik. Penggunaan IndoBERT sesuai dengan karakteristik data media sosial karena mampu mempertimbangkan konteks kata dalam teks Bahasa Indonesia. Temuan ini sejalan dengan penelitian terdahulu mengenai penerapan IndoBERT untuk analisis sentimen Bahasa Indonesia (Ahmadian et al., 2024; Singgalen, 2025). Kemampuan IndoBERT dalam memahami konteks bahasa juga menjadi dasar pemilihannya dalam penelitian ini.

Tahapan *preprocessing* yang meliputi *case folding*, *cleaning*, normalisasi, tokenisasi, *stopword removal*, dan *stemming* digunakan untuk mengurangi variasi serta *noise* pada komentar media sosial. Proses tersebut penting karena data mengandung bahasa tidak baku, simbol, dan variasi penulisan. Pendekatan ini sejalan dengan penelitian terdahulu yang menunjukkan pentingnya penanganan variasi bahasa dan *stopword* pada data komentar berbahasa Indonesia (Aribowo et al., 2021).

Distribusi data awal menunjukkan kelas netral mendominasi, sehingga diterapkan SMOTE untuk menyeimbangkan jumlah data pada setiap kelas. Setelah balancing, masing-masing kelas memiliki jumlah data yang sama. Penerapan tersebut sesuai dengan prinsip SMOTE yang menghasilkan sampel sintetis untuk memperkuat representasi kelas minoritas (Chawla et al., 2002). Dengan demikian, SMOTE pada penelitian ini berfungsi mengurangi ketidakseimbangan distribusi data sebelum proses pelatihan model.

Berdasarkan *confusion matrix*, kelas 1 memiliki tingkat prediksi benar sekitar 96%, sedangkan kelas 3 dan kelas 5 masing-masing sekitar 89%. Kesalahan terbesar terjadi antara kelas 3 dan kelas 5, yaitu sekitar 9% data kelas 3 diprediksi sebagai kelas 5 dan 11% data kelas 5 diprediksi sebagai kelas 3. Kondisi tersebut menunjukkan adanya kemiripan karakteristik

bahasa antar kedua kelas, terutama pada komentar yang mengandung kritik ringan atau konteks yang tidak secara eksplisit menunjukkan polaritas.

Distribusi sentimen menunjukkan bahwa sentimen netral merupakan kategori dominan, diikuti negatif dan positif. Dominasi tersebut menunjukkan bahwa sebagian besar komentar lebih bersifat informatif atau belum menunjukkan dukungan maupun penolakan secara eksplisit. Sementara itu, keberadaan sentimen negatif menunjukkan adanya perhatian masyarakat terhadap aspek pelaksanaan Program MBG. Hasil ini juga sesuai dengan kesimpulan penelitian bahwa respons masyarakat didominasi kategori netral, kemudian negatif dan positif.

Analisis lintas isu menunjukkan bahwa distribusi sentimen berbeda pada setiap aspek Program MBG. Isu Gizi dan kategori Lainnya memiliki jumlah komentar paling besar, sedangkan Distribusi memiliki jumlah paling kecil. Hasil tersebut menunjukkan bahwa respons masyarakat terhadap MBG tidak bersifat homogen dan perlu dianalisis berdasarkan isu yang dibicarakan. Pendekatan ini sejalan dengan penelitian terdahulu yang menghubungkan analisis sentimen dengan aspek pada *user-generated content* Bahasa Indonesia (Perwira et al., 2025).

Hasil *wordcloud* memperkuat interpretasi tersebut. Sentimen positif lebih banyak menampilkan kata seperti *bagus*, *manfaat*, *sehat*, dan *enak*, sedangkan sentimen negatif memperlihatkan kata seperti *kritik*, *salah*, *basi*, *racun*, dan *kurang*. Pada sentimen netral, kata yang dominan antara lain *program*, *pemerintah*, *anak*, *sekolah*, *makan*, dan *gizi*. Perbedaan tersebut menunjukkan karakteristik leksikal yang berbeda pada setiap kategori sentimen, meskipun *wordcloud* digunakan sebagai analisis pendukung dan bukan sebagai dasar utama pelabelan.

Secara keseluruhan, kombinasi IndoBERT dan SMOTE mampu memberikan kinerja klasifikasi yang baik sekaligus menghasilkan gambaran mengenai persepsi masyarakat terhadap Program MBG. Namun, hasil penelitian memiliki keterbatasan karena data hanya berasal dari Instagram dan YouTube serta menggunakan *rule-based labeling* yang berpotensi menghasilkan *label noise*, khususnya pada komentar yang mengandung sarkasme, ironi, dan makna implisit. Keterbatasan tersebut perlu dipertimbangkan dalam menginterpretasikan hasil dan menjadi dasar pengembangan penelitian selanjutnya.

Kesimpulan

Penelitian ini menerapkan *text mining* untuk menganalisis sentimen komentar Instagram dan YouTube mengenai Program Makan Bergizi Gratis (MBG) menggunakan model IndoBERT dan teknik SMOTE. Dari 52.326 komentar yang dikumpulkan melalui *web scraping*, diperoleh 47.861 komentar setelah proses preprocessing dan pembersihan data. Tahapan preprocessing meliputi *case folding*, *cleaning*, normalisasi kata tidak baku, tokenisasi, *stopword removal*, dan *stemming*.

Hasil pelabelan menunjukkan bahwa sentimen netral merupakan kategori dominan dengan 31.136 komentar, diikuti sentimen negatif sebanyak 4.025 komentar dan positif sebanyak 3.131 komentar. Ketidakseimbangan tersebut kemudian ditangani menggunakan SMOTE sehingga masing-masing kelas memiliki 31.136 data. Model IndoBERT menghasilkan *accuracy* sebesar 89,88%, *precision* 92,55%, *recall* 89,88%, dan F1-score 90,63%, yang menunjukkan kinerja klasifikasi yang baik.

Hasil *confusion matrix* menunjukkan bahwa sebagian besar data berhasil diklasifikasikan dengan benar, dengan tingkat prediksi sekitar 96% pada kelas 1 dan sekitar 89% pada kelas 3 dan 5. Kesalahan klasifikasi terutama terjadi antara kelas 3 dan 5 yang memiliki kemiripan pola bahasa. Analisis lintas isu dan *wordcloud* juga menunjukkan adanya perbedaan karakteristik respons masyarakat pada setiap isu dan kategori sentimen. Secara keseluruhan, kombinasi IndoBERT dan SMOTE dapat digunakan untuk mengidentifikasi pola sentimen masyarakat terhadap Program MBG.

Penelitian ini memiliki keterbatasan pada penggunaan *rule-based labeling* yang berpotensi menghasilkan *label noise*, terutama pada komentar yang mengandung sarkasme, ironi, dan konteks implisit. Selain itu, data hanya berasal dari Instagram dan YouTube sehingga belum sepenuhnya merepresentasikan opini masyarakat secara luas. Penelitian selanjutnya dapat menggunakan *manual labeling* atau *semi-supervised learning*, memperluas sumber data, serta membandingkan IndoBERT dengan model bahasa lainnya.

Acknowledgment

-

Daftar Pustaka

- Adhantoro, M. S., Haq, F. A., Hartanto, D., & Sudaryanto, A. P. (2026). IndoBERT-Based Sentiment Analysis of Electric Motorcycle Policy in Indonesia Using Instagram Data. *Jurnal Penelitian Sains Teknologi*, 165–181. <https://doi.org/10.23917/saintek.v2i2.17021>
- Ahmadian, H., Abidin, T. F., Riza, H., & Muchtar, K. (2024). Hybrid Models for Emotion Classification and Sentiment Analysis in Indonesian Language. *Applied Computational Intelligence and Soft Computing*, 2024(1), 2826773. <https://doi.org/10.1155/2024/2826773>
- Anaissi, A., Suleiman, B., & Zandavi, S. M. (2021). Online Tensor-Based Learning Model for Structural Damage Detection. *ACM Transactions on Knowledge Discovery from Data*, 15(6), 1–18. <https://doi.org/10.1145/3451217>
- Aprilah, T., Setiadi, D. R. I. M., & Herowati, W. (2025). Enhancing Aspect-Based Sentiment Analysis via Hugging Face Fine-Tuned IndoBERT. *Journal of Applied Informatics and Computing*, 9(6), 3821–3830. <https://doi.org/10.30871/jaic.v9i6.11409>
- Aribowo, A. S., Basiron, H., Yusof, N. F. A., & Khomsah, S. (2021). Cross-domain sentiment analysis model on Indonesian YouTube comment. *International Journal of Advances in Intelligent Informatics*, 7(1), 12. <https://doi.org/10.26555/ijain.v7i1.554>
- Babo, D. H. P., Sidi, K. R., & Imtiyaz, H. (2026). Implementation of the free nutritious meals (MBG) program: A health policy perspective from Jakarta. *Jurnal Cakrawala Promkes*, 8(1), 50–59. <https://doi.org/10.12928/jcp.v8i1.15522>
- Badriyah, Chamidy, T., & Suhartono. (2025). Application of SMOTE in Sentiment Analysis of MyXL User Reviews on Google Play Store. *JISKA (Jurnal Informatika Sunan Kalijaga)*, 10(1), 74–86. <https://doi.org/10.14421/jiska.2025.10.1.74-86>
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research*, 16, 321–357. <https://doi.org/10.1613/jair.953>

- Ihalauw, S. A., Trezandy Lapatta, N., Wiria Nugraha, D., Wirdayanti, & Ar Lamasitudju, C. (2025). Analisis Sentimen Terhadap Kinerja Awal Pemerintahan Menggunakan IndoBERT Dan SMOTE Pada Media Sosial X. *Jurnal Algoritma*, 22(2). <https://doi.org/10.33364/algoritma/v.22-2.2957>
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP. *Proceedings of the 28th International Conference on Computational Linguistics*, 757–770. <https://doi.org/10.18653/v1/2020.coling-main.66>
- Kristiyanti, D. A., Sanjaya, S. A., Tjokro, V. C., & Suhali, J. (2024). Dealing imbalance dataset problem in sentiment analysis of recession in Indonesia. *IAES International Journal of Artificial Intelligence (IJ-AI)*, 13(2), 2060. <https://doi.org/10.11591/ijai.v13.i2.pp2060-2072>
- Manoppo, M. R., Kolang, I. C., Nur Fiat, D. N., Mawara, R. M. C., Sumarno, A. D. P., Yusupa, A., & Tarigan, V. (2025). Analisis Sentimen Publik Di Media Sosial Terhadap Kenaikan PPN 12% Di Indonesia Menggunakan Indobert. *Jurnal Kecerdasan Buatan dan Teknologi Informasi*, 4(2), 152–163. <https://doi.org/10.69916/jkbti.v4i2.322>
- Medantoro, G. F. S., & Muljono, M. (2026). Comparative Analysis of IndoBERT and Classic Machine Learning Models for Sentiment Classification of Education Policy on Social Media X. *Journal of Applied Informatics and Computing*, 10(1), 548–557. <https://doi.org/10.30871/jaic.v10i1.11723>
- Muhabbab, A. Z., Bunyamin, B., & Hasmawati, H. (2025). Sentiment Analysis of Indonesia's Free Nutritious Meal Program on Platform X (Formerly Twitter) Using IndoBERT. *ZERO: Jurnal Sains, Matematika dan Terapan*, 9(3), 884. <https://doi.org/10.30829/zero.v9i3.27629>
- Negara, A. B. P. (2025). The influence of applying stopword removal and SMOTE on Indonesian sentiment classification. *Lontar Komputer: Jurnal Ilmiah Teknologi Informasi*, 14(3), 172–185. <https://doi.org/10.24843/LKJITI.2023.v14.i03.p05>
- Perwira, R. I., Permadi, V. A., Purnamasari, D. I., & Agusdin, R. P. (2025). Domain-Specific Fine-Tuning of IndoBERT for Aspect-Based Sentiment Analysis in Indonesian Travel User-Generated Content. *Journal of Information Systems Engineering and Business Intelligence*, 11(1), 30–40. <https://doi.org/10.20473/jisebi.11.1.30-40>
- Rifaldi, D., Abdul Fadlil, & Herman. (2023). Teknik Preprocessing Pada Text Mining Menggunakan Data Tweet “Mental Health.” *Decode: Jurnal Pendidikan Teknologi Informasi*, 3(2), 161–171. <https://doi.org/10.51454/decode.v3i2.131>
- Roihan, M. R., Muhammad Itqan Mazdadi, Muliadi, Irwan Budiman, & Fatma Indriani. (2026). Comparative performance of IndoBERT-based deep learning models with SMOTE for trade tariff sentiment analysis. *Journal of Soft Computing Exploration*, 7(2), 432–446. <https://doi.org/10.52465/josce.v7i2.122>
- Salu, B. M., Irwanda, D. A., Hasibuan, N. R. K., & Arifin, S. (2026). Pengembangan Sistem Informasi Persediaan Alat Tulis Kantor Berbasis Web Di PT BPRS Amanah Bangsa. *Journal Artificial: Informatika Dan Sistem Informasi*, 4(1), 37–50. <https://doi.org/10.54065/artificial.1104>

- Singgale, Y. A. (2025). Performance Analysis of IndoBERT for Sentiment Classification in Indonesian Hotel Review Data. *Journal of Information System Research (JOSH)*, 6(2), 976–986. <https://doi.org/10.47065/josh.v6i2.6505>
- Suprpto, F. A., Praditya, E., Dewi, R. M., & Adiyoso, W. (2025). A Policy Implementation Review of the Free Nutritious Meal (MBG) Program. *The Journal of Indonesia Sustainable Development Planning*, 6(2), 297–312. <https://doi.org/10.46456/jisdep.v6i2.798>
- Wafda, A. (2024). Integrasi Machine Learning dalam Ritel: Tinjauan Komprehensif tentang Prediksi Harga, Analisis Data Pelanggan, dan Pemanfaatan Media Sosial. *Journal Artificial: Informatika Dan Sistem Informasi*, 2(2), 90–106. <https://doi.org/10.54065/artificial.543>
- Yulita, I. N., Wijaya, V., Rosadi, R., Sarathan, I., Djujandi, Y., & Prabuwno, A. S. (2023). Analysis of Government Policy Sentiment Regarding Vacation during the COVID-19 Pandemic Using the Bidirectional Encoder Representation from Transformers (BERT). *Data*, 8(3), 46. <https://doi.org/10.3390/data8030046>