

Studi Komparatif IndoBERT dan SVM-TF-IDF untuk Analisis Sentimen Program Makan Bergizi Gratis Nasional

Septian Nuno Zildjian^{*1}, Dian Nurdiana², Fonda Leviany³, Medi Taruk⁴

¹²³Program Studi Sains Data, Universitas Terbuka, Jakarta, ⁴Program Studi Informatika, Fakultas Teknik, Universitas Mulawarman

^{*1}septiannuno123@gmail.com, ²dian.nurdiana@ecampus.ut.ac.id, ³fonda.leviany@ecampus.ut.ac.id, ⁴meditaruk@unmul.ac.id

Abstrak

Program Makan Bergizi Gratis (MBG) merupakan kebijakan sosial yang memicu beragam respons masyarakat, khususnya di media sosial. Penelitian ini bertujuan menganalisis opini publik terhadap Program MBG melalui komentar TikTok, mengidentifikasi distribusi sentimen, serta membandingkan kinerja model IndoBERT dan Support Vector Machine (SVM) berbasis TF-IDF dalam klasifikasi sentimen berbahasa Indonesia. Data komentar diproses melalui tahapan preprocessing, meliputi pembersihan teks, penghapusan duplikasi, normalisasi bahasa, pengolahan emoji, case folding, dan penyesuaian format teks. Hasil penelitian menunjukkan bahwa IndoBERT memberikan performa yang lebih baik dibandingkan SVM-TF-IDF, dengan nilai accuracy sebesar 89,61% dan F1-makro 83,47%, sedangkan SVM-TF-IDF memperoleh accuracy 78,45% dan F1-makro 64,67%. Distribusi sentimen menunjukkan dominasi sentimen negatif (71,7%), diikuti sentimen netral (14,3%) dan positif (14,1%). Temuan ini mengindikasikan adanya kritik dan kekhawatiran masyarakat terhadap pelaksanaan Program MBG, namun tidak dapat dimaknai sebagai penolakan secara menyeluruh. Penelitian ini menunjukkan bahwa analisis sentimen berbasis Natural Language Processing dapat menjadi pendekatan yang efektif untuk memahami dinamika opini publik dan mendukung evaluasi komunikasi kebijakan berbasis data..

Kata Kunci—Analisis Sentimen; IndoBERT; MBG; Natural Language Processing; Support Vector Machine

1. PENDAHULUAN

Media sosial menjadi ruang bagi masyarakat untuk menyampaikan opini terhadap berbagai isu, termasuk kebijakan pemerintah. Analisis sentimen sebagai salah satu pendekatan Natural Language Processing (NLP) memungkinkan klasifikasi opini ke dalam sentimen positif, netral, dan negatif [1][2]. TikTok dipilih karena komentar pengguna bersifat spontan dan banyak mengandung bahasa informal. Program Makan Bergizi Gratis (MBG) merupakan kebijakan yang memunculkan beragam respons terkait gizi, anggaran, distribusi, kualitas makanan, dan efektivitas pelaksanaannya. Oleh karena itu, analisis sentimen terhadap komentar TikTok diperlukan untuk memahami kecenderungan opini publik terhadap Program MBG secara lebih sistematis.

Analisis sentimen telah banyak diterapkan untuk mengkaji respons masyarakat terhadap berbagai kebijakan publik [3][5][15][16][18][20]. SVM berbasis TF-IDF dikenal efektif untuk klasifikasi teks, sedangkan IndoBERT memiliki keunggulan dalam memahami konteks bahasa Indonesia [4][7][10][17]. Namun, penelitian mengenai analisis sentimen komentar TikTok terhadap Program Makan Bergizi Gratis (MBG) masih terbatas. Oleh karena itu, penelitian ini menganalisis distribusi sentimen opini publik terhadap MBG, membandingkan performa IndoBERT dan SVM-TF-IDF, serta menentukan model yang lebih efektif. Hasil penelitian diharapkan mendukung pengembangan metode analisis sentimen dan evaluasi komunikasi kebijakan berbasis data..

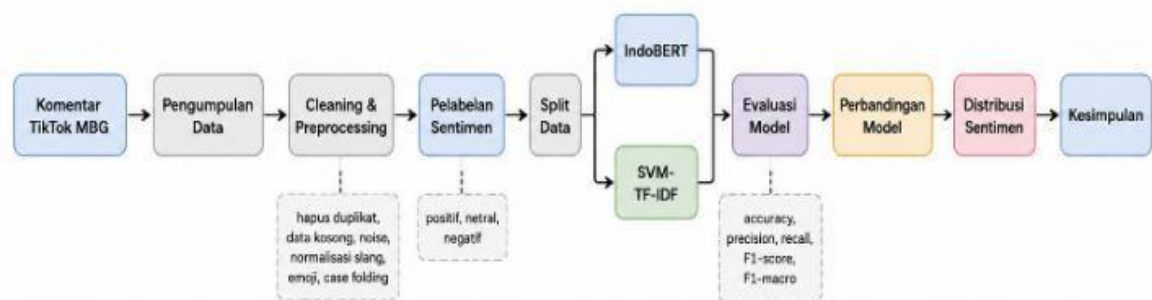
2. METODE PENELITIAN

2.1. Desain Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode analisis sentimen berbasis *Natural Language Processing* (NLP). Pendekatan ini digunakan karena analisis sentimen bertujuan untuk mengidentifikasi dan mengklasifikasikan opini dalam teks ke dalam kategori positif, netral, dan negatif [2]. Data yang dianalisis berupa komentar TikTok terkait Program Makan Bergizi Gratis (MBG), sehingga hasil penelitian disajikan dalam bentuk distribusi sentimen dan evaluasi performa model secara numerik. Pemilihan metode ini didukung oleh penelitian terdahulu yang menunjukkan bahwa analisis sentimen

banyak digunakan untuk mengkaji opini publik terhadap isu sosial dan kebijakan pemerintah di Indonesia. Beberapa studi telah menerapkan SVM pada isu kenaikan harga BBM, kenaikan pajak hiburan, kebijakan lockdown, UU TNI, dan RUU Perampasan Aset [3], [5], [16], [18], [20]. Selain itu, beberapa penelitian juga membandingkan SVM dengan Naive Bayes, Decision Tree, Random Forest, dan IndoBERT dalam klasifikasi sentimen [13], [14], [17], [19].

Desain penelitian ini bersifat komparatif karena membandingkan IndoBERT sebagai model berbasis transformer bahasa Indonesia dan SVM-TF-IDF sebagai model machine learning klasik. IndoBERT digunakan karena mampu memahami konteks bahasa Indonesia melalui model pra-latih berbasis transformer [7], [8], sedangkan SVM-TF-IDF digunakan sebagai pembanding karena SVM efektif untuk klasifikasi teks dan TF-IDF dapat merepresentasikan teks ke dalam fitur numerik [9], [10]. Secara umum, alur penelitian meliputi pengumpulan data, preprocessing, pelabelan dan pembagian data, pemodelan, evaluasi, serta analisis hasil sentimen. Alur penelitian ditampilkan pada Gambar 1.

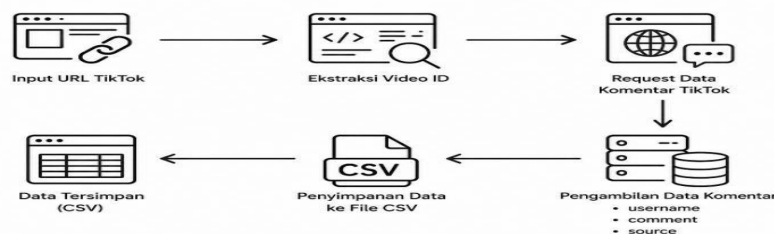


Gambar 1. Kerangka Berpikir Penelitian

2.2. Pengumpulan Data

Data penelitian berupa komentar pengguna TikTok yang membahas Program Makan Bergizi Gratis (MBG). Komentar media sosial digunakan karena dapat merepresentasikan opini digital masyarakat terhadap isu sosial dan kebijakan publik [1], [2]. Dataset awal terdiri dari **35.162** baris data dengan tiga atribut utama, yaitu **username**, **comment**, dan **source**. Pengumpulan data dilakukan menggunakan tool scraping sederhana yang dibuat secara mandiri dengan memanfaatkan URL postingan TikTok sebagai input utama. Tool bekerja dengan mengekstraksi video ID dari URL, mengirimkan request ke endpoint komentar TikTok, lalu menyimpan komentar yang diperoleh ke dalam file CSV. Data yang dikumpulkan terdiri dari atribut username, comment, dan source, dengan atribut comment sebagai fokus utama analisis sentimen.

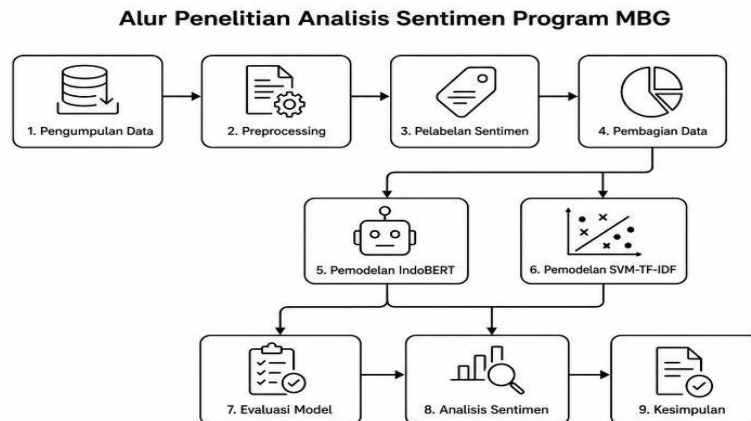
Alur proses pengumpulan data terdiri dari input URL TikTok, ekstraksi video ID, request data komentar, pengambilan atribut komentar, dan penyimpanan hasil ke dalam file CSV. Alur pengumpulan data tersebut diilustrasikan secara rinci pada Gambar 2.



Gambar 2. Alur Scaraping Data Komentar TikTok

2.3. Alur Penelitian

Alur penelitian dimulai dari pengumpulan komentar TikTok terkait Program Makan Bergizi Gratis (MBG), dilanjutkan dengan preprocessing, pelabelan sentimen, dan pembagian data menjadi data latih, validasi, dan data uji. Selanjutnya dilakukan pemodelan menggunakan IndoBERT dan SVM-TF-IDF, kemudian dievaluasi berdasarkan accuracy, precision, recall, F1-score, F1-macro, confusion matrix, dan AUC-ROC. Hasil evaluasi digunakan untuk membandingkan performa kedua model serta menganalisis kecenderungan opini publik terhadap Program MBG. Alur penelitian ditunjukkan pada Gambar 3.



Gambar 3. Alur Penelitian Analisis Sentimen Program MBG

2.4 Data Preprocessing

Tahap preprocessing dilakukan untuk membersihkan dan menormalisasi komentar TikTok sebelum pelabelan dan pemodelan sentimen. Proses ini mencakup penghapusan komentar kosong dan duplikat, normalisasi teks, serta penyaringan noise seperti komentar yang terlalu pendek, hanya berisi angka, simbol, atau tidak informatif. Dari 35.162 komentar awal, tersisa 30.963 komentar setelah penghapusan data kosong dan duplikat, kemudian menjadi 29.199 komentar setelah proses penyaringan noise. Ringkasan ukuran dataset dan pembagian data disajikan pada Tabel 1.

Tabel 1. Ringkasan Ukuran Dataset dan Pembagian Data

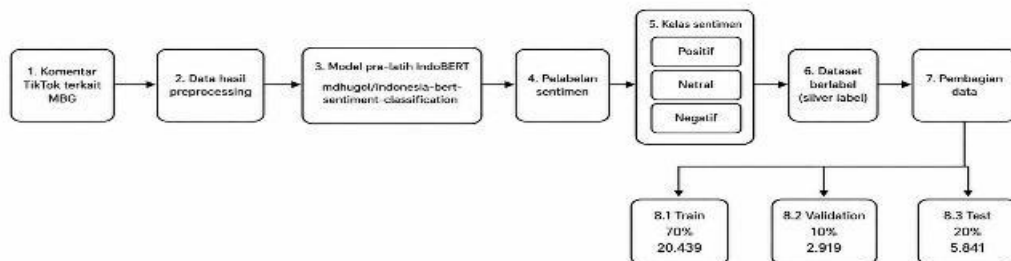
Komponen	Nilai
Data awal	35.162
Setelah buang kosong/duplikat	30.963
Data akhir modeling	29.199
Train	20.439
Validation	2.919
Test	5.841

Tahapan preprocessing dalam penelitian ini meliputi *case folding*, *cleansing*, konversi emoji, normalisasi slang, *stopword removal*, dan stemming. Tahapan tersebut dilakukan karena kualitas preprocessing dapat memengaruhi hasil klasifikasi sentimen, terutama pada data media sosial yang banyak mengandung bahasa informal dan variasi penulisan [23]. Pada jalur IndoBERT, teks diproses sampai tahap teks bersih dengan tetap mempertahankan struktur dan konteks kalimat karena model berbasis transformer bekerja dengan memahami konteks antarkata dalam kalimat [7], [8], [34], [35]. Sementara itu, pada jalur SVM-TF-IDF, preprocessing dilanjutkan dengan *stopword removal* dan stemming menggunakan Sastrawi agar teks menjadi lebih ringkas dan sesuai untuk pembobotan kata berbasis TF-IDF [9], [10], [27], [28]. Tahap *case folding* dilakukan dengan mengubah seluruh huruf menjadi huruf kecil. Tahap ini bertujuan untuk menyeragamkan bentuk kata, sehingga kata yang sama tidak dianggap berbeda hanya karena perbedaan huruf kapital.

2.5 IndoBERT-Based Sentiment Labeling and Data Splitting

Pelabelan sentimen dilakukan menggunakan model pra-latih berbasis BERT, yaitu *mdhugol/indonesia-bert-sentiment-classification*. Model ini digunakan untuk menghasilkan label awal komentar TikTok terkait Program Makan Bergizi Gratis (MBG) ke dalam tiga kelas, yaitu positif, netral, dan negatif. Penggunaan model pra-latih dipilih karena dataset belum memiliki label sentimen manual. Pendekatan berbasis BERT relevan dalam pemrosesan bahasa alami karena mampu menghasilkan representasi teks secara kontekstual [35], serta didukung oleh pengembangan IndoBERT dan sumber daya NLP bahasa Indonesia [7], [8]. Label yang dihasilkan digunakan sebagai *silver label*, yaitu label hasil prediksi model. Label tersebut kemudian menjadi dasar untuk melatih dan mengevaluasi model IndoBERT

serta SVM-TF-IDF. Proses ini sejalan dengan tujuan analisis sentimen, yaitu mengklasifikasikan opini teks ke dalam kategori sentimen tertentu [2]. Setelah pelabelan selesai, data dibagi menjadi data latih, data validasi, dan data uji dengan proporsi 70%, 10%, dan 20%. Data latih digunakan untuk pelatihan model, data validasi untuk memantau performa selama pelatihan, sedangkan data uji digunakan untuk mengevaluasi performa akhir model.



Gambar 4. Alur Pelabelan Sentimen IndoBERT

2.6 Pemodelan IndoBERT

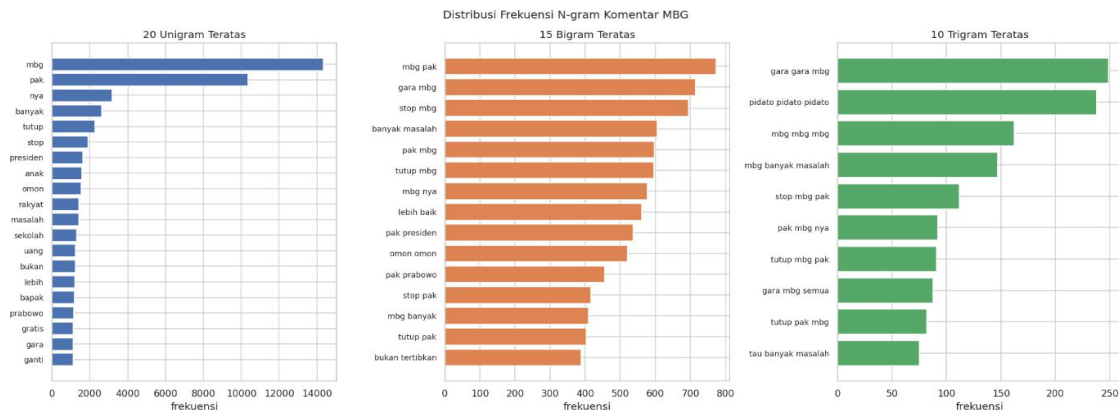
Pemodelan IndoBERT dilakukan untuk mengklasifikasikan komentar TikTok terkait Program Makan Bergizi Gratis (MBG) ke dalam tiga kelas sentimen, yaitu negatif, netral, dan positif. IndoBERT digunakan karena merupakan model berbasis transformer untuk bahasa Indonesia yang mampu memahami konteks kalimat melalui hubungan antarkata dalam teks [7], [8], [34], [35]. Input yang digunakan adalah teks_bert, yaitu teks hasil preprocessing yang masih mempertahankan struktur kalimat. Teks diproses menggunakan tokenizer IndoBERT dengan batas maksimum 64 token. Komentar yang lebih pendek diberi *padding*, sedangkan komentar yang lebih panjang dipotong melalui *truncation*. Proses *fine-tuning* dilakukan menggunakan data latih dan validasi. Karena distribusi kelas tidak seimbang, digunakan *weighted CrossEntropyLoss* agar model tidak terlalu bias terhadap kelas mayoritas. Optimizer yang digunakan adalah AdamW, dan model terbaik dipilih berdasarkan performa validasi. Evaluasi menggunakan *accuracy* dan F1-makro karena F1-makro lebih sesuai untuk dataset tidak seimbang [12].

2.7 Pemodelan SVM-TF-IDF

Pemodelan SVM-TF-IDF digunakan sebagai pembandingan terhadap IndoBERT. Model ini memakai input teks_svm, yaitu teks hasil preprocessing lanjutan yang telah melalui *stopword removal* dan stemming agar fitur kata menjadi lebih ringkas dan stabil. Komentar direpresentasikan ke dalam bentuk numerik menggunakan TF-IDF, yaitu metode pembobotan kata berdasarkan frekuensi kemunculan dan tingkat kepentingannya dalam korpus [10], [28]. Setelah itu, klasifikasi dilakukan menggunakan Support Vector Machine karena metode ini efektif untuk klasifikasi teks dan banyak digunakan dalam analisis sentimen [9], [27]. SVM digunakan sebagai baseline karena lebih ringan, cepat, dan mudah diinterpretasikan. Pengaturan parameter dilakukan menggunakan GridSearchCV untuk mencari nilai C terbaik. Hasil model kemudian dievaluasi pada data uji yang sama dengan IndoBERT agar perbandingan performa kedua model tetap adil.

2.8 Eksplorasi Awal dan Karakteristik Teks

Eksplorasi awal dilakukan untuk memahami karakteristik komentar TikTok sebelum masuk ke tahap preprocessing dan pemodelan. Tahap ini penting karena komentar media sosial umumnya bersifat pendek, informal, tidak seragam, serta banyak mengandung hashtag, mention, emoji, dan variasi penulisan. Kondisi tersebut sesuai dengan karakter media sosial sebagai ruang opini digital yang spontan [1], serta perlu diperhatikan dalam proses analisis sentimen [2]. Visualisasi distribusi panjang karakter, jumlah kata, hashtag, mention, dan emoji digunakan untuk melihat bentuk awal data. Hasil eksplorasi menunjukkan bahwa komentar memiliki panjang yang bervariasi. Komentar yang terlalu pendek perlu diperhatikan karena konteks sentimennya terbatas, sedangkan komentar yang lebih panjang menjadi pertimbangan dalam penentuan batas maksimum token pada IndoBERT. Oleh karena itu, eksplorasi ini menjadi dasar dalam proses pembersihan teks, normalisasi, dan pembatasan token [7], [8], [23], [35].



Gambar 5. Frekuensi token dominan setelah stopwords removal.

Gambar 5 menunjukkan distribusi frekuensi N-gram komentar MBG setelah *stopword removal*. Token "mbg" mendominasi unigram dengan frekuensi sekitar 14.000, diikuti "gara", "banyak", dan "makan". Pada bigram, "mbg pak" dan "gara mbg" paling sering muncul (sekitar 700–800), sedangkan trigram teratas adalah "gara gara mbg" (sekitar 250). Visualisasi ini memastikan bahwa dataset memuat konteks pembahasan Program MBG secara relevan, sekaligus mendukung penggunaan TF-IDF pada model SVM karena pembobotan kata bergantung pada kemunculan token dalam korpus [10].

2.9 Analisis Pola Bahasa dan Ekspresi Komentar

Analisis pola bahasa dilakukan untuk melihat kecenderungan kata, ekspresi, dan gaya penulisan yang muncul dalam komentar TikTok terkait Program Makan Bergizi Gratis. Tahap ini penting karena media sosial menjadi ruang opini digital yang memungkinkan pengguna menyampaikan respons secara spontan, singkat, dan informal [1]. Dalam analisis sentimen, karakteristik tersebut perlu diperhatikan karena bentuk teks dapat memengaruhi proses klasifikasi opini [2]. Visualisasi word cloud digunakan untuk memberikan gambaran umum mengenai kata-kata yang sering muncul dalam komentar. Kata yang tampil dengan ukuran lebih besar menunjukkan frekuensi kemunculan yang lebih tinggi dalam keseluruhan data.



Gambar 6. Word cloud keseluruhan komentar MBG.

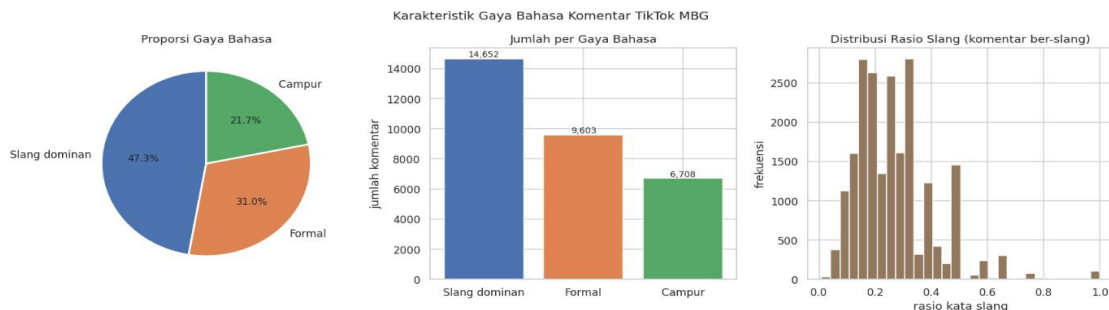
Gambar 6 menunjukkan bahwa percakapan terkait MBG tidak hanya memuat kata yang berhubungan dengan program, tetapi juga kata yang berkaitan dengan pelaksanaan, aktor kebijakan, kritik, dan respons masyarakat. Word cloud tidak digunakan sebagai bukti utama, tetapi membantu membaca konteks percakapan secara lebih cepat sebelum data digunakan dalam proses pemodelan sentimen. Selain kata, emoji juga dianalisis karena menjadi bagian dari ekspresi sentimen di media sosial. Pada dataset ini, sekitar 17,4% komentar mengandung emoji. Oleh karena itu, emoji tidak langsung dihapus karena dapat membawa makna emosi, terutama pada komentar pendek. Perlakuan ini sejalan dengan pentingnya preprocessing pada data media sosial agar informasi yang relevan tidak hilang sebelum proses klasifikasi dilakukan [22], [23].



Gambar 7. Analisis emoji dan simbol ekspresif.

Grafik tersebut menunjukkan bahwa emoji dan simbol ekspresif dapat memperkuat makna komentar. Karena itu, emoji dikonversi ke bentuk deskripsi teks agar tetap dapat diproses oleh model. Langkah ini dilakukan agar sinyal sentimen dari emoji tetap terbaca, baik pada model IndoBERT yang memanfaatkan

konteks kalimat maupun pada SVM-TF-IDF yang menggunakan representasi kata sebagai fitur [7], [8], [10], [35]. Selanjutnya, analisis slang dilakukan untuk melihat penggunaan bahasa informal dalam komentar. Komentar TikTok banyak mengandung singkatan, bahasa percakapan, dan variasi ejaan tidak baku seperti “yg”, “gak”, “aja”, atau bentuk informal lainnya.

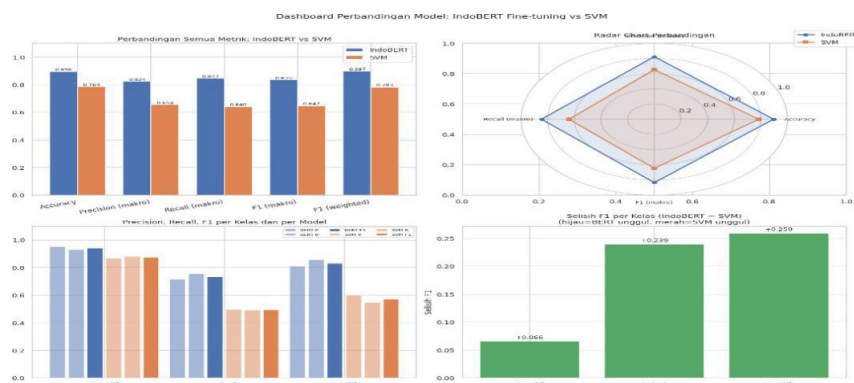


Gambar 8. Pola penggunaan slang dan bahasa informal.

Gambar tersebut menunjukkan bahwa sebesar 47,0% komentar didominasi slang, 31,6% bersifat formal, dan 21,3% merupakan campuran keduanya. Secara absolut, komentar ber-slang berjumlah 14.432, formal 9.603, dan campuran 6.100. Distribusi rasio slang juga memperlihatkan bahwa sebagian besar komentar memiliki rasio slang rendah hingga menengah, namun proporsi komentar dengan slang dominan tetap signifikan. Temuan ini menunjukkan bahwa normalisasi slang diperlukan agar kata-kata tidak baku dapat diseragamkan menjadi bentuk yang lebih standar. Tahap ini membantu mengurangi variasi penulisan dan membuat teks lebih siap digunakan dalam proses pemodelan IndoBERT maupun SVM-TF-IDF. Normalisasi ini juga relevan karena kualitas *preprocessing* dapat memengaruhi hasil analisis sentimen pada data media sosial [23].

2.9 Evaluasi dan Perbandingan Kinerja Model

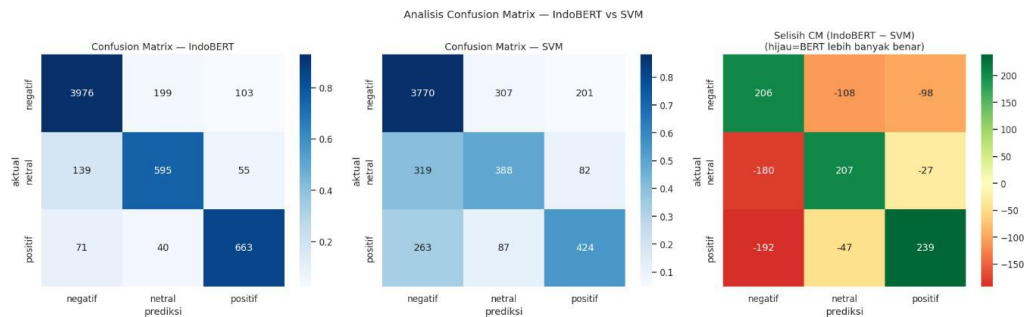
Evaluasi model dilakukan pada data uji sebanyak 5.841 komentar dengan menggunakan beberapa metrik, yaitu *accuracy*, *precision*, *recall*, F1-score, F1-makro, dan AUC-ROC. Penggunaan beberapa metrik diperlukan karena dataset memiliki distribusi kelas yang tidak seimbang.



Gambar 9. Perbandingan metrik utama IndoBERT dan SVM

F1-makro digunakan karena mampu menghitung performa setiap kelas secara lebih seimbang [12]. IndoBERT memperoleh *accuracy* sebesar 0,8961 dan F1-makro sebesar 0,8347. Sementara itu, SVM-TF-IDF memperoleh *accuracy* sebesar 0,7845 dan F1-makro sebesar 0,6467. Selisih performa ini memperlihatkan bahwa representasi kontekstual IndoBERT lebih sesuai untuk komentar media sosial yang pendek, informal, dan sering mengandung makna implisit. Hal ini sejalan dengan karakteristik model berbasis BERT yang mampu memahami konteks antarkata dalam kalimat [7], [8], [35]. SVM-TF-IDF tetap digunakan sebagai pembanding karena efektif untuk klasifikasi teks berbasis fitur kata [9], [10]. Perbandingan metrik utama memperkuat hasil bahwa IndoBERT lebih unggul dibandingkan SVM-TF-IDF pada seluruh metrik: *accuracy* 0,896 vs 0,783, *precision* (makro) 0,824 vs 0,674, *recall* (makro) 0,848 vs 0,640, dan F1-makro 0,835 vs 0,647. Selisih F1 per kelas menunjukkan keunggulan terbesar pada kelas netral (+0,239) dan positif (+0,259), sedangkan kelas negatif memiliki selisih lebih kecil (+0,066). Radar chart memperlihatkan area IndoBERT yang secara konsisten lebih luas dibandingkan SVM pada semua dimensi evaluasi. Keunggulan ini menunjukkan bahwa model berbasis transformer lebih mampu menangkap konteks kalimat pada komentar TikTok yang pendek, informal, dan tidak selalu eksplisit. Temuan ini juga sejalan dengan beberapa penelitian yang menunjukkan bahwa model berbasis BERT dapat memberikan performa kompetitif atau lebih baik dibandingkan model klasik dalam analisis sentimen

berbahasa Indonesia [17], [29], [33].



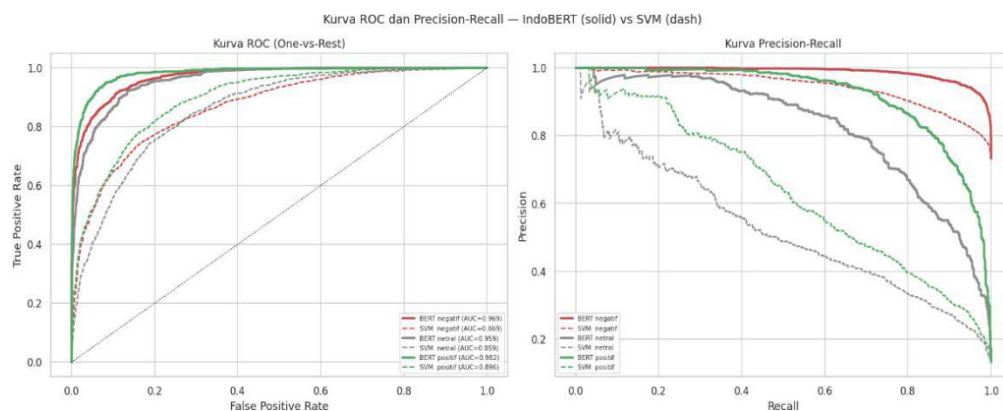
Gambar 10. Confusion Matrix dan Metrik per Kelas.

Pola kesalahan klasifikasi memperlihatkan bahwa pada IndoBERT, kelas negatif berhasil dikenali dengan benar sebanyak 3.976 komentar, kelas netral 595, dan kelas positif 663. Pada SVM-TF-IDF, nilai tersebut lebih rendah, yaitu negatif 3.770, netral 388, dan positif 424. Selisih confusion matrix memperlihatkan IndoBERT lebih banyak benar pada kelas netral (+207) dan positif (+239), sementara SVM lebih banyak salah mengklasifikasikan kelas netral dan positif sebagai negatif. Kelas negatif relatif lebih stabil dikenali oleh kedua model, sedangkan kelas netral lebih sering tertukar dengan kelas lain. Hal ini dapat terjadi karena komentar netral di media sosial tidak selalu benar-benar netral, tetapi dapat berupa pertanyaan, sindiran halus, atau komentar informatif yang masih memiliki nuansa sentimen. Oleh karena itu, evaluasi per kelas diperlukan agar performa model tidak hanya dinilai dari *accuracy* secara umum.

Tabel 2. Perbandingan AUC-ROC per Kelas

Kelas	IndoBERT AUC	SVM AUC	Selisih
Negatif	0,9691	0,8690	0,1001
Netral	0,9586	0,8586	0,1000
Positif	0,9816	0,8963	0,0853

Nilai AUC-ROC IndoBERT lebih tinggi pada seluruh kelas, yaitu 0,9691 untuk negatif, 0,9586 untuk netral, dan 0,9816 untuk positif. Nilai tersebut menunjukkan bahwa IndoBERT memiliki kemampuan pemisahan kelas yang lebih konsisten dibandingkan SVM-TF-IDF pada berbagai ambang keputusan.



Gambar 11. Kurva ROC Multikelas.

Kurva ROC multikelas digunakan untuk melihat kemampuan model dalam membedakan setiap kelas sentimen pada berbagai threshold. Hasil ini memperkuat bahwa IndoBERT memiliki kemampuan diskriminasi kelas yang lebih baik dibandingkan SVM-TF-IDF, tidak hanya pada satu titik evaluasi tertentu. Perbandingan AUC-ROC per kelas memperlihatkan keunggulan IndoBERT secara konsisten, dengan nilai 0,969 untuk kelas negatif, 0,959 untuk kelas netral, dan 0,982 untuk kelas positif, dibandingkan SVM-TF-IDF yang memperoleh 0,869, 0,859, dan 0,896 pada kelas yang sama. Selisih terbesar terdapat pada kelas negatif dan netral (sekitar 0,10), sedangkan kelas positif memiliki selisih sekitar 0,086. Dengan demikian, hasil evaluasi secara keseluruhan mendukung simpulan bahwa IndoBERT lebih sesuai digunakan untuk klasifikasi sentimen komentar TikTok terkait Program MBG, terutama karena data bersifat pendek, informal, tidak seimbang, dan membutuhkan pemahaman konteks.

4. HASIL DAN PEMBAHASAN

Hasil penelitian menunjukkan bahwa preprocessing dan pemilihan model berpengaruh terhadap performa analisis sentimen. Komentar TikTok yang pendek, informal, serta banyak mengandung slang dan emoji membutuhkan pembersihan serta normalisasi agar dapat diproses dengan baik. Tahap preprocessing ini penting karena kualitas teks dapat memengaruhi hasil klasifikasi sentimen [22], [23].

Berdasarkan evaluasi pada data uji, IndoBERT memperoleh *accuracy* sebesar 0,8961 dan F1-makro sebesar 0,8347. Sementara itu, SVM-TF-IDF memperoleh *accuracy* sebesar 0,7845 dan F1-makro sebesar 0,6467. Hasil ini menunjukkan bahwa IndoBERT lebih unggul karena mampu memahami konteks kalimat secara lebih baik melalui arsitektur transformer [7], [8], [35]. SVM-TF-IDF tetap berguna sebagai pembanding karena lebih ringan dan mudah diinterpretasikan melalui bobot kata [9], [10].

Dari sisi sentimen, komentar terkait Program Makan Bergizi Gratis didominasi oleh sentimen negatif sebesar 71,7%, diikuti sentimen netral sebesar 14,3% dan positif sebesar 14,1%. Dominasi negatif menunjukkan bahwa percakapan digital dalam dataset banyak memuat kritik, kekhawatiran, atau ketidakpuasan. Namun, hasil ini tidak dapat langsung dimaknai sebagai penolakan menyeluruh terhadap program karena data media sosial tidak selalu mewakili seluruh populasi.

5. KESIMPULAN

Penelitian ini menunjukkan bahwa IndoBERT memiliki performa lebih baik dibandingkan SVM-TF-IDF dalam analisis sentimen komentar TikTok terkait Program Makan Bergizi Gratis (MBG), dengan *accuracy* sebesar 89,61% dan F1-makro 83,47%, sedangkan SVM-TF-IDF memperoleh *accuracy* 78,45% dan F1-makro 64,67%. Hasil klasifikasi didominasi sentimen negatif (71,7%), diikuti sentimen netral (14,3%) dan positif (14,1%), yang mengindikasikan adanya kritik atau kekhawatiran masyarakat terhadap pelaksanaan program. Namun, temuan ini belum dapat mewakili persepsi masyarakat secara keseluruhan. Penelitian selanjutnya disarankan memperluas sumber data, melakukan validasi label secara manual, serta menerapkan topic modeling atau aspect-based sentiment analysis.

6. SARAN

Penelitian selanjutnya disarankan menerapkan aspect-based sentiment analysis atau topic modeling untuk mengidentifikasi aspek-aspek yang mendominasi sentimen negatif, seperti anggaran, kualitas makanan, distribusi, dan komunikasi publik. Selain itu, penggunaan model lain seperti IndoBERTweet, XGBoost, Random Forest, maupun pendekatan hybrid berbasis lexicon dan machine learning dapat menjadi pembanding guna memperoleh metode analisis sentimen yang lebih komprehensif untuk teks media sosial berbahasa Indonesia.

DAFTAR PUSTAKA

- [1] D. M. Boyd and N. B. Ellison, "Social network sites: Definition, history, and scholarship," *Journal of Computer-Mediated Communication*, vol. 13, no. 1, pp. 210–230, 2007. doi: <http://dx.doi.org/10.1111/j.1083-6101.2007.00393.x>
- [2] B. Liu, "Sentiment Analysis and Opinion Mining." Morgan & Claypool Publishers, 2012. [Online]. Available: <https://www.cs.uic.edu/~liub/FBS/SentimentAnalysis-and-OpinionMining.pdf>
- [3] H. S. 'Adilah and R. Alit, "Analisis sentimen masyarakat Twitter terhadap kebijakan pemerintah dalam menaikkan harga bahan bakar minyak dengan menggunakan metode Support Vector Machine," *JINACS*, vol. 5, no. 2, pp. 201–215, 2023. doi: <http://dx.doi.org/10.26740/jinacs.v5n02.p201-215>
- [4] M. R. Manoppo et al., "Analisis sentimen publik di media sosial terhadap kenaikan PPN 12% di Indonesia menggunakan IndoBERT," *Jurnal Kecerdasan Buatan dan Teknologi Informasi*, vol. 4, no. 2, pp. 152–163, 2025. doi: <http://dx.doi.org/10.69916/jkbt.v4i2.322>
- [5] W. Romadhona and A. R. Isnain, "Analisis sentimen pengguna media sosial terhadap kebijakan kenaikan pajak hiburan menggunakan metode SVM," *JIPi: Jurnal Ilmiah Penelitian dan Pembelajaran Informatika*, vol. 9, no. 4, pp. 2185–2195, 2024. doi: <http://dx.doi.org/10.29100/jipi.v9i4.5603>
- [6] E. Syahrul and D. Fatharani, "Hybrid sentiment analysis of Maxim app users using Support Vector Machine and lexicon-based approach," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 3S1, 2025. doi: <http://dx.doi.org/10.23960/jitet.v13i3S1.8148>
- [7] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, "IndoLEM and IndoBERT: A benchmark dataset and

pre-trained language model for Indonesian NLP," in Proc. COLING, 2020. [Online]. Available: <https://aclanthology.org/2020.coling-main.66/>

[8] B. Wilie et al., "IndoNLU: Benchmark and resources for evaluating Indonesian natural language understanding," in Proc. AACL-IJCNLP, 2020. [Online]. Available: <https://aclanthology.org/2020.aacl-main.85/>

[9] C. Cortes and V. Vapnik, "Support-vector networks," Machine Learning, vol. 20, pp. 273–297, 1995. doi: <http://dx.doi.org/10.1007/BF00994018>

[10] G. Salton and C. Buckley, "Term-weighting approaches in automatic text retrieval," Information Processing & Management, vol. 24, no. 5, pp. 513–523, 1988. doi: [http://dx.doi.org/10.1016/0306-4573\(88\)90021-0](http://dx.doi.org/10.1016/0306-4573(88)90021-0)

[11] A. C. Najib, A. Irsyad, G. A. Qandi, and N. A. Rakhmawati, "Perbandingan metode lexicon-based dan SVM untuk analisis sentimen berbasis ontologi pada kampanye Pilpres Indonesia tahun 2019 di Twitter," Fountain of Informatics Journal, vol. 4, no. 2, pp. 41–48, 2019. doi: <http://dx.doi.org/10.21111/fij.v4i2.3573>

[12] M. Sokolova and G. Lapalme, "A systematic analysis of performance measures for classification tasks," Information Processing & Management, vol. 45, no. 4, pp. 427–437, 2009. doi: <http://dx.doi.org/10.1016/j.ipm.2009.03.002>

[13] I. Daulay, M. Firdaus, B. Ardiansyah, R. Hutagaol, and R. Rahmaddeni, "Analisis sentimen opini publik terhadap penerimaan bantuan subsidi upah (BSU) menggunakan algoritma Naive Bayes," SENTIMAS, vol. 1, no. 1, pp. 155–162, 2023. [Online]. Available: <https://journal.irpi.or.id/index.php/sentimas/article/view/602>

[14] F. A. Ahda and M. Zainuddin, "Prediksi kepuasan pelayanan perpustakaan menggunakan algoritma Decision Tree (C4.5)," Jurnal Teknologi Informasi, vol. 10, no. 2, pp. 143–150, 2019. doi: <http://dx.doi.org/10.36382/jti-tki.v10i2.368>

[15] W. P. Anggraini and M. S. Utami, "Klasifikasi sentimen masyarakat terhadap kebijakan Kartu Pekerja di Indonesia," Faktor Exacta, vol. 13, no. 4, pp. 255, 2021. doi: <http://dx.doi.org/10.30998/faktorexacta.v13i4.7964>

[16] A. Nurhasananda and M. Akbar, "Analisis sentimen masyarakat terhadap kebijakan Undang-Undang Tentara Nasional Indonesia (UU TNI) menggunakan Support Vector Machine," Jurnal Komputer Informasi dan Teknologi, vol. 5, no. 1, pp. 14–14, 2025. doi: <http://dx.doi.org/10.53697/jkomitek.v5i1.2603>

[17] N. N. A. Aryanti and O. Suria, "Analisis sentimen terhadap keputusan hubungan kerja di Indonesia: Komparasi IndoBERT dengan SVM, Random Forest, dan Decision Tree dengan optimasi TF-IDF," Rabbit: Jurnal Teknologi dan Sistem Informasi Univrab, vol. 10, no. 2, pp. 1158–1176, 2025. doi: <http://dx.doi.org/10.36341/rabit.v10i2.6364>

[18] A. R. Isnain, A. I. Sakti, D. Alita, and N. S. Marga, "Sentimen analisis publik terhadap kebijakan lockdown pemerintah Jakarta menggunakan algoritma SVM," Jurnal Data Mining dan Sistem Informasi, vol. 2, no. 1, pp. 31–37, 2021. doi: <http://dx.doi.org/10.33365/jdmsi.v2i1.1021>

[19] R. Yunita and M. Kamayani, "Perbandingan algoritma SVM dan Naïve Bayes pada analisis sentimen penghapusan kewajiban skripsi," The Indonesian Journal of Computer Science, vol. 12, no. 5, 2023. doi: <http://dx.doi.org/10.33022/ijcs.v12i5.3415>

[20] L. Rofiqi and M. Akbar, "Analisis sentimen terkait RUU Perampasan Aset dengan Support Vector Machine," JEKIN: Jurnal Teknik Informatika, vol. 4, no. 3, pp. 529–538, 2024. doi: <http://dx.doi.org/10.58794/jekin.v4i3.824>

[21] F. A. Indriyani, A. Fauzi, and S. Faisal, "Analisis sentimen aplikasi TikTok menggunakan algoritma Naïve Bayes dan Support Vector Machine," TEKNOSAINS: Jurnal Sains, Teknologi dan Informatika, vol. 10, no. 2, pp. 176–184, 2023. doi: <http://dx.doi.org/10.37373/tekno.v10i2.419>

[22] P. A. Permatasari, L. Linawati, and L. Jasa, "Survei tentang analisis sentimen pada media sosial," Majalah Ilmiah Teknologi Elektro, vol. 20, no. 2, pp. 177–186, 2021. doi: <http://dx.doi.org/10.24843/MITE.2021.v20i02.P01>

- [23] S. Khairunnisa, A. Adiwijaya, and S. Al Faraby, "Pengaruh text preprocessing terhadap analisis sentimen komentar masyarakat pada media sosial Twitter," *Jurnal Media Informatika Budidarma*, vol. 5, no. 2, pp. 406, 2021. doi: <http://dx.doi.org/10.30865/mib.v5i2.2835>
- [24] P. Arsi and R. Waluyo, "Sentiment analysis of discourse on moving the Indonesian capital city using the Support Vector Machine (SVM) algorithm," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 8, no. 1, pp. 147–156, 2021, doi: <https://doi.org/10.25126/jtiik.0813944>.
- [25] T. A. Dewi and E. Mailoa, "Perbandingan implementasi metode SMOTE pada algoritma Support Vector Machine (SVM) dalam analisis sentimen opini masyarakat tentang Mixue," *Jurnal Indonesia: Manajemen Informatika dan Komunikasi*, vol. 4, no. 3, pp. 849–855, 2023. doi: <http://dx.doi.org/10.35870/jimik.v4i3.289>
- [26] S. Juanita, "Analisis sentimen persepsi masyarakat terhadap Pemilu 2019 pada media sosial Twitter menggunakan Naive Bayes," *Jurnal Media Informatika Budidarma*, vol. 4, no. 3, 2020. doi: <http://dx.doi.org/10.30865/mib.v4i3.2140>
- [27] L. B. Ilmawan and M. A. Mude, "Perbandingan metode klasifikasi Support Vector Machine dan Naive Bayes pada analisis sentimen Twitter," *SMATIKA Jurnal*, vol. 10, no. 2, 2020. doi: <http://dx.doi.org/10.32664/smatika.v10i02.455>
- [28] D. Sugiarto, E. Utami, and A. Yaqin, "Perbandingan kinerja model TF-IDF dan BOW untuk klasifikasi opini publik tentang kebijakan BLT minyak goreng," *Jurnal Teknik Industri*, vol. 12, no. 3, 2022. doi: <http://dx.doi.org/10.25105/jti.v12i3.15669>
- [29] S. A. Salsabila, B. Priyatna, A. Hananto, and Tukino, "Komparasi kinerja model Naive Bayes, SVM, dan BERT dalam klasifikasi sentimen ulasan pada aplikasi YUMMY," *STORAGE: Jurnal Ilmiah Teknik dan Ilmu Komputer*, vol. 4, no. 2, pp. 42–47, 2025. doi: <http://dx.doi.org/10.55123/storage.v4i2.5120>
- [30] F. F. Mailoa, "Analisis sentimen data Twitter menggunakan metode text mining tentang masalah obesitas di Indonesia," *Journal of Information Systems for Public Health*, vol. 6, no. 1, pp. 44, 2021. doi: <http://dx.doi.org/10.22146/jisph.44455>
- [31] A. N. Syafia, M. F. Hidayattullah, and W. Sutеды, "Studi komparasi algoritma SVM dan Random Forest pada analisis sentimen komentar YouTube BTS," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 8, no. 3, pp. 207–212, 2023. doi: <http://dx.doi.org/10.30591/jpit.v8i3.5064>
- [32] I. Septiana and D. Alita, "Perbandingan Random Forest dan SVM dalam analisis sentimen quick count Pemilu 2024," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 9, no. 3, pp. 224–233, 2024. doi: <http://dx.doi.org/10.30591/jpit.v9i3.6640>
- [33] F. Alvin and N. A. S. Winarsih, "Perbandingan kinerja model IndoBERT, IndoBERTweet, dan algoritma klasik pada analisis sentimen isu Indonesia Gelap," *Building of Informatics, Technology and Science (BITS)*, vol. 7, no. 3, pp. 1601–1613, 2025. doi: <https://doi.org/10.47065/bits.v7i3.8636>
- [34] A. Vaswani et al., "Attention is all you need," in *Advances in Neural Information Processing Systems*, 2017. doi: <https://doi.org/10.5555/3295222.3295349>
- [35] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. NAACL-HLT*, pp. 4171–4186, 2019. [Online]. Available: <https://aclanthology.org/N19-1423/>