

KOMPARASI ALGORITMA INDOBERT, SVM, DAN RANDOM FOREST DALAM ANALISIS SENTIMEN PROGRAM MAKAN BERGIZI GRATIS

Saepul Rohman, Ignatius Wiseto Prasetyo Agung

Teknik Informatika, Universitas Adhirajasa Reswara Sanjaya
Jalan Sekolah Internasional No. 1-2, Antapani, Kota Bandung, Jawa Barat, Indonesia
Saepulrohman3445@gmail.com

ABSTRAK

Program Makan Bergizi Gratis merupakan kebijakan pemerintah yang telah dimasukkan dalam Rencana Kerja Pemerintah (RKP) 2025 dan memicu diskusi publik yang luas, terutama di media sosial X. Permasalahan utama dalam penelitian ini adalah bagaimana memahami secara akurat sentimen publik terhadap program tersebut dan menentukan model machine learning mana yang paling akurat untuk mengklasifikasikan opini masyarakat yang diekspresikan dalam bahasa Indonesia yang kompleks di media sosial. Oleh karena itu, penelitian ini bertujuan untuk menganalisis sentimen publik terhadap program tersebut serta membandingkan kinerja model IndoBERT dengan algoritma Support Vector Machine (SVM) dan Random Forest untuk menentukan model yang paling efektif. Metode penelitian yang digunakan meliputi pengumpulan 54.105 tweet dari media sosial X dengan kata kunci "makan bergizi gratis", diikuti oleh tahap pra-pemrosesan data, sebelum akhirnya data tersebut diklasifikasikan menggunakan tiga model yang dibandingkan. Hasil analisis menunjukkan bahwa sentimen publik didominasi oleh respons positif (43,9%) dan netral (41,8%), sementara sentimen negatif hanya sebesar 14,3%. Dari perbandingan kinerja model, IndoBERT terbukti superior dengan mencapai akurasi 87,60%, secara signifikan mengungguli SVM (83,39%) dan Random Forest (70,28%). Keunggulan IndoBERT ini disebabkan oleh kemampuannya dalam memahami konteks bidirectional bahasa Indonesia yang kompleks, menjadikannya model yang paling efektif untuk analisis sentimen pada dataset ini.

Kata kunci : Analisis Sentimen, Makan Bergizi Gratis, IndoBERT, Support Vector Machine, Random Forest, Media Sosial X

1. PENDAHULUAN

Program Makan Bergizi Gratis merupakan salah satu kebijakan strategis yang telah dimasukkan dalam Rencana Kerja Pemerintah (RKP) 2025. Program ini dirancang dengan tujuan utama untuk meningkatkan kualitas gizi pada tiga kelompok prioritas: anak sekolah, ibu hamil, dan balita [1]. Inisiatif ini tidak hanya mencakup penyediaan makanan bergizi dan susu untuk anak-anak dari tingkat prasekolah hingga sekolah menengah atas dan pesantren, tetapi juga memberikan asupan nutrisi tambahan bagi ibu hamil. Pemerintah berharap program ini dapat menjadi katalisator bagi peningkatan perkembangan kognitif anak, percepatan penurunan angka stunting, dan peningkatan angka partisipasi sekolah secara nasional [2].

Meskipun inisiatif ini menunjukkan komitmen kuat pemerintah dalam membangun generasi masa depan yang lebih unggul, implementasinya dihadapkan pada tantangan yang signifikan, terutama dari aspek keberlanjutan fiskal. Dengan estimasi anggaran tahunan yang berkisar antara Rp100 triliun hingga Rp450 triliun, berbagai lembaga kredibel, termasuk Bank Dunia serta lembaga pemeringkat internasional seperti Fitch dan Moody's, telah menyuarakan kekhawatiran mengenai potensi dampaknya terhadap stabilitas ekonomi makro Indonesia. Selain keterbatasan anggaran, tantangan operasional yang kompleks juga membayangi program ini, mencakup logistik distribusi yang merata ke seluruh pelosok negeri, pengawasan standar gizi yang

ketat, serta mitigasi risiko pemborosan makanan dalam skala besar [2].

Kekhawatiran terhadap dampak program ini semakin relevan jika dikaitkan dengan masalah stunting di Indonesia yang masih menjadi isu krusial. Berdasarkan data Survei Status Gizi Indonesia (SSGI) 2023, pemerintah menargetkan penurunan prevalensi stunting hingga 14% pada tahun 2024, sementara angka pada tahun 2022 masih berada di level 21,6% [3]. Kontroversi program semakin meluas dengan adanya wacana penggunaan dana Bantuan Operasional Sekolah (BOS), yang memicu penolakan dari berbagai kalangan, termasuk Federasi Serikat Guru Indonesia (FSGI) [4]. Kondisi ini menimbulkan perdebatan publik yang intens mengenai potensi dampaknya terhadap biaya pendidikan, kesejahteraan guru, dan stabilitas keuangan negara secara keseluruhan [1].

Di tengah kontroversi tersebut, pemahaman mendalam mengenai respons dan persepsi masyarakat menjadi sangat penting. Di era digital saat ini, media sosial, khususnya platform X (sebelumnya Twitter), telah menjadi arena utama bagi diskursus publik. Platform ini tidak hanya berfungsi sebagai saluran penyebaran informasi kebijakan pemerintah, tetapi juga sebagai wadah bagi masyarakat untuk berinteraksi, berdebat, dan menyuarakan aspirasi mereka [5]. Untuk memetakan sentimen publik secara akurat, diperlukan metode analisis yang mampu mengolah data teks dalam volume besar dengan efisien. Analisis sentimen, sebagai salah satu cabang dari *Natural Language Processing* (NLP),

menawarkan kapabilitas untuk mengklasifikasikan opini publik ke dalam kategori positif, netral, atau negatif secara sistematis [3].

Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap Program Makan Bergizi Gratis dan membandingkan kinerja model IndoBERT dengan dua algoritma machine learning klasik, yaitu Support Vector Machine (SVM) dan Random Forest. IndoBERT dipilih karena arsitektur bidirectional berbasis Transformer-nya unggul dalam memahami konteks kalimat secara komprehensif [6]. Kinerja IndoBERT akan dievaluasi dan dibandingkan secara komprehensif dengan dua algoritma machine learning klasik yang telah banyak digunakan dalam penelitian sejenis: *Support Vector Machine* (SVM) dan *Random Forest*. Penelitian sebelumnya telah menunjukkan bahwa SVM mampu mencapai akurasi yang tinggi, berkisar antara 75% hingga 94% [3], [4], sementara *Random Forest* juga menunjukkan performa yang kompetitif dalam tugas klasifikasi teks [5]. Tujuan akhir dari penelitian ini adalah menentukan model yang paling efektif dan akurat dalam mengklasifikasikan sentimen publik berbahasa Indonesia, sehingga hasilnya dapat menjadi masukan berbasis bukti bagi pemerintah dalam evaluasi dan optimalisasi kebijakan. Untuk mencapai tujuan tersebut, penelitian ini akan dilakukan melalui beberapa tahapan. Pertama, pengumpulan data opini publik dari media sosial X yang relevan dengan Program Makan Bergizi Gratis. Kedua, pra-pemrosesan teks yang mencakup pembersihan data, normalisasi bahasa, tokenisasi, dan penghapusan kata tidak relevan untuk mengurangi noise. Ketiga, pelatihan tiga model (IndoBERT, SVM, dan Random Forest) menggunakan dataset yang telah diproses. Keempat, evaluasi kinerja model menggunakan metrik seperti akurasi, presisi, recall, dan F1-score. Terakhir, analisis perbandingan hasil kinerja untuk menentukan model terbaik.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis Sentimen, sering juga disebut opinion mining, merupakan suatu bidang studi komputasi yang berfokus pada identifikasi, ekstraksi, kuantifikasi, dan analisis afek, sentimen, serta subjektivitas dalam data teks [7]. Tujuan utamanya adalah untuk menentukan sikap seorang penulis atau pembicara terhadap suatu topik tertentu atau polaritas kontekstual dari sebuah dokumen. Informasi yang diekstrak dapat diklasifikasikan ke dalam beberapa tingkatan, mulai dari level dokumen (apakah seluruh dokumen mengekspresikan opini positif atau negatif), level kalimat (menentukan sentimen setiap kalimat), hingga level aspek (*aspect-based sentiment analysis*) yang mengidentifikasi opini terhadap fitur atau atribut spesifik dari suatu entitas. Dalam konteks kebijakan publik, analisis sentimen memungkinkan pemerintah untuk mengukur denyut nadi opini publik secara real-

time, mengidentifikasi area keprihatinan, dan mengevaluasi efektivitas komunikasi kebijakan [8].

2.2. Text Mining

Text Mining adalah proses multidisipliner yang menggabungkan teknik dari NLP, statistik, machine learning, dan data mining untuk mengekstraksi informasi dan pengetahuan berkualitas tinggi dari data teks tidak terstruktur [9]. Proses ini mengubah teks menjadi data numerik terstruktur yang dapat dianalisis lebih lanjut. Tahapan dalam *Text Mining* umumnya meliputi pengumpulan data teks, pra-pemrosesan (seperti *tokenisasi*, *Stemming*, dan *stopword removal*), transformasi fitur (misalnya menggunakan model Bag-of-Words atau TF-IDF), dan penerapan algoritma machine learning untuk tugas seperti klasifikasi, klastering, atau pemodelan topik [10].

2.3. Transformer dan IndoBERT

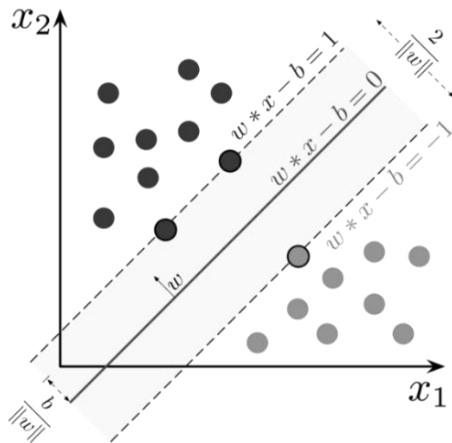
Arsitektur *Transformer*, yang pertama kali diperkenalkan oleh Vaswani et al. pada tahun 2017, merevolusi bidang NLP. Berbeda dengan model sekuensial seperti RNN atau LSTM, *Transformer* sepenuhnya mengandalkan mekanisme yang disebut *self-attention* [11]. Mekanisme ini memungkinkan model untuk menimbang pentingnya kata-kata yang berbeda dalam sebuah urutan ketika memproses kata tertentu. Dengan kata lain, model dapat melihat seluruh kalimat sekaligus dan memahami hubungan kontekstual antara kata-kata yang berjauhan. Arsitektur ini terdiri dari komponen *Encoder* dan *Decoder*, di mana *Encoder* memetakan urutan input menjadi representasi kontinu, dan *Decoder* menghasilkan urutan output kata per kata.

BERT (*Bidirectional Encoder Representations from Transformers*) adalah model yang hanya menggunakan komponen *Encoder* dari arsitektur *Transformer*. Keunikan BERT terletak pada kemampuannya untuk dilatih secara bidirectional, artinya model ini mempelajari konteks sebuah kata berdasarkan semua kata di sekitarnya (baik di kiri maupun di kanan) secara simultan. IndoBERT adalah varian BERT yang secara khusus dilatih dari awal menggunakan korpus bahasa Indonesia yang sangat besar dan beragam, yang dikenal sebagai Indo4B, yang berisi sekitar 4 miliar kata [6]. Pelatihan ekstensif ini memberikan IndoBERT pemahaman mendalam tentang tata bahasa, semantik, dan nuansa pragmatis bahasa Indonesia, menjadikannya sangat efektif untuk berbagai tugas NLP, termasuk analisis sentimen.

2.4. Support Vector Machine (SVM)

SVM adalah algoritma supervised learning yang sangat kuat untuk tugas klasifikasi dan regresi. Prinsip dasar SVM adalah menemukan hyperplane optimal yang dapat memisahkan titik-titik data dari kelas yang berbeda dengan margin (jarak) semaksimal mungkin [5]. Titik-titik data yang paling dekat dengan hyperplane dan menentukan posisinya disebut support vectors. Salah satu kekuatan terbesar SVM adalah

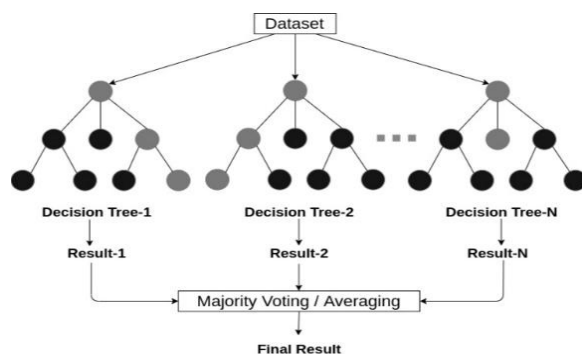
kemampuannya untuk menangani data yang tidak dapat dipisahkan secara linear melalui teknik yang disebut *kernel trick*. *Kernel trick* secara implisit memetakan data input ke ruang fitur berdimensi lebih tinggi di mana pemisah linear dapat ditemukan. Beberapa fungsi kernel yang umum digunakan adalah *Linear*, *Polynomial*, *Radial Basis Function* (RBF), dan *Sigmoid*. Mekanisme kerja algoritma ini secara visual dapat dipahami melalui ilustrasi pada Gambar 1.



Gambar 1. Algoritma Support Vector Machine

2.5. Random Forest

Random Forest adalah salah satu metode ensemble learning yang paling populer dan efektif. Algoritma ini bekerja dengan membangun sejumlah besar decision tree (pohon keputusan) pada saat pelatihan [5]. Untuk tugas klasifikasi, hasil akhir ditentukan oleh majority voting dari semua pohon individual. Dua ide utama di balik *Random Forest* adalah: (1) Bagging (Bootstrap Aggregating), di mana setiap pohon dilatih pada sampel acak dengan penggantian dari dataset asli, dan (2) Random Subspace, di mana pada setiap simpul pohon, pemisahan hanya dipertimbangkan pada subset acak dari total fitur. Kombinasi kedua teknik ini membantu mengurangi varians model dan mencegah overfitting, menghasilkan model yang kuat dan stabil. Representasi visual dari arsitektur *Random Forest* dapat dilihat pada Gambar 2.

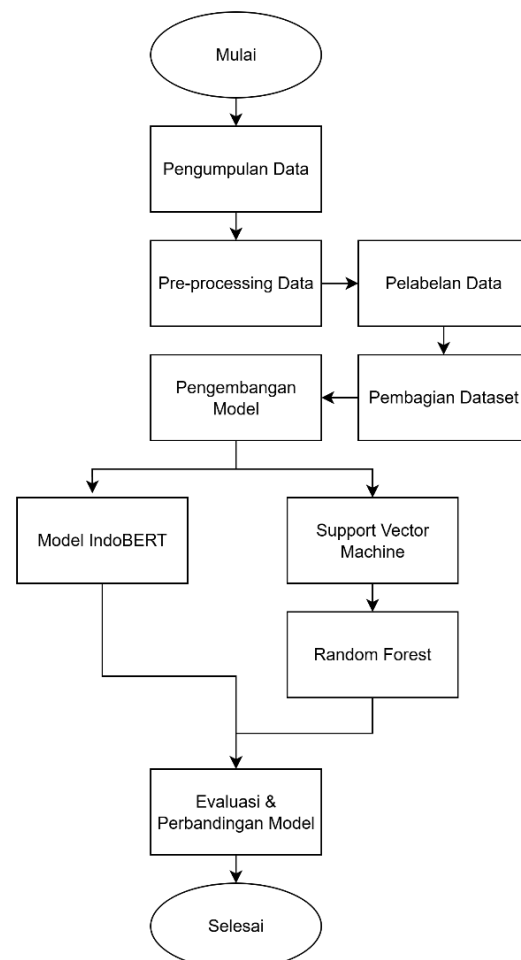


Gambar 2. Algoritma Random Forest

2.6. Tinjauan Studi

Berbagai penelitian telah mengaplikasikan algoritma-algoritma ini untuk analisis sentimen. Studi oleh [12] menunjukkan bahwa model BERT mencapai akurasi 83% dalam analisis sentimen kemajuan AI, mengungguli model lain. Secara spesifik untuk bahasa Indonesia, [13] mengonfirmasi keunggulan IndoBERT dengan akurasi 87% untuk analisis sentimen pembelajaran daring. Dalam konteks kebijakan publik, [3] berhasil menerapkan SVM untuk menganalisis sentimen program makan siang gratis dengan akurasi impresif sebesar 94%. Sementara itu, [4] juga menggunakan SVM dan mencapai akurasi 75,39%. Studi komparatif oleh [5] menemukan bahwa *Random Forest* (akurasi 91%) sedikit lebih unggul dari SVM (akurasi 90%) untuk analisis sentimen Metaverse, menyoroti bahwa performa algoritma dapat bervariasi tergantung pada domain dan karakteristik data. Penelitian ini dibangun di atas fondasi studi-studi tersebut dengan melakukan perbandingan langsung antara model *Transformer* canggih (IndoBERT) dengan dua algoritma klasik yang terbukti andal (SVM dan *Random Forest*) pada dataset kebijakan publik yang relevan dan berskala besar.

3. METODE PENELITIAN



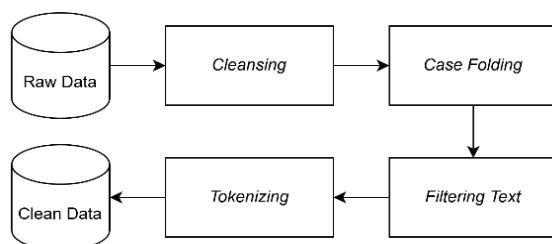
Gambar 3. Alur Penelitian

Metodologi penelitian ini disusun secara sistematis untuk memastikan analisis yang valid dan komprehensif. Alur penelitian mencakup beberapa tahapan utama yang dijelaskan di bawah ini. Gambar 3 di bawah ini menunjukkan alur penelitian secara keseluruhan.

3.1. Pengumpulan dan Karakteristik Data

Objek penelitian ini adalah sentimen publik yang diekspresikan di media sosial X mengenai Program Makan Bergizi Gratis. Proses pengumpulan data dilakukan dengan teknik crawling menggunakan alat *Tweet Harvest*. Kata kunci pencarian yang digunakan adalah "makan bergizi gratis" untuk memastikan relevansi data. Periode data yang dikumpulkan mencakup rentang waktu dari 1 Juni 2024 hingga 31 Maret 2025. Proses ini berhasil mengumpulkan total 54.105 *tweet* berbahasa Indonesia. Dataset mentah ini memiliki 15 kolom yang mencakup informasi seperti ID *tweet*, ID pengguna, teks lengkap *tweet*, waktu pembuatan, jumlah suka, jumlah balasan, dan jumlah *retweet*.

3.2. Pra-pemrosesan Data



Gambar 4. Tahapan Alur Pra-pemrosesan Data

Data mentah dari media sosial dikenal sangat noisy dan tidak terstruktur. Oleh karena itu, tahap pra-pemrosesan yang teliti sangat krusial untuk meningkatkan kualitas data dan performa model [14]. Proses ini dilakukan melalui serangkaian langkah berikut:

a. Data Cleaning

Tahap awal ini berfokus pada penghapusan elemen-elemen yang tidak informatif untuk analisis sentimen. Ini termasuk penghapusan URL (misalnya, <http://...>), mention pengguna (misalnya, @username), hashtag (misalnya, #program), karakter non-alfanumerik (tanda baca, simbol), dan angka [15].

b. Case Folding

Seluruh teks dalam dataset dikonversi menjadi huruf kecil. Langkah ini bertujuan untuk menstandarisasi data sehingga kata-kata seperti "Program", "program", dan "PROGRAM" dianggap sebagai token yang sama [16].

c. Normalisasi

Kata-kata slang, singkatan, dan typo yang umum di media sosial dinormalisasi ke bentuk standarnya. Misalnya, 'yg' diubah menjadi 'yang', 'ga' atau 'gak' menjadi 'tidak', dan 'bgt' menjadi 'banget' [17].

d. Tokenisasi

Teks yang telah dibersihkan kemudian dipecah menjadi unit-unit yang lebih kecil yang disebut token (biasanya kata). Proses ini menggunakan pustaka NLTK (Natural Language Toolkit) [18].

e. Stopword Removal

Kata-kata yang sering muncul tetapi memiliki sedikit makna semantik (kata henti) dihapus. Penelitian ini menggunakan daftar kata henti bahasa Indonesia dari pustaka Sastrawi dan memperkayanya dengan daftar kata henti tambahan yang spesifik untuk konteks media sosial [17].

f. Stemming

Langkah terakhir adalah mereduksi kata-kata berimbuhan ke bentuk kata dasarnya. Proses ini menggunakan algoritma *Stemming* Nazief dan Adriani yang diimplementasikan dalam pustaka Sastrawi. Sebagai contoh, kata-kata seperti "memberikan", "diberikan", dan "pemberian" semuanya akan direduksi menjadi kata dasar "beri" [16].

3.3. Pelabelan dan Pembagian Dataset

Mengingat volume data yang besar, pelabelan dilakukan secara otomatis menggunakan model IndoBERT pre-trained yang telah dioptimalkan untuk klasifikasi sentimen bahasa Indonesia (mdhugol/indonesia-bert-sentiment-classification). Model ini mengklasifikasikan setiap *tweet* ke dalam salah satu dari tiga kategori: positif, netral, atau negatif.

Setelah seluruh dataset dilabeli, data dibagi menjadi tiga subset untuk keperluan pelatihan dan evaluasi model. Pembagian ini menggunakan metode stratified sampling untuk memastikan bahwa distribusi proporsi setiap kelas sentimen (positif, netral, negatif) tetap konsisten di semua subset. Proporsi pembagiannya adalah:

- Data Pelatihan (*Training Set*): 80% dari total data, digunakan untuk melatih model.
- Data Validasi (*Validation Set*): 10% dari total data, digunakan untuk fine-tuning hyperparameter dan memantau overfitting selama proses pelatihan.
- Data Pengujian (*Testing Set*): 10% dari total data, digunakan untuk evaluasi akhir performa model pada data yang belum pernah dilihat sebelumnya.

3.4. Pengembangan dan Evaluasi Model

Penelitian ini mengembangkan dan membandingkan tiga model klasifikasi:

a. Model IndoBERT

Menggunakan arsitektur pre-trained indolem/indobert-base-uncased. Model ini di-fine-tuning pada dataset pelatihan untuk tugas klasifikasi tiga kelas sentimen. Proses fine-tuning melibatkan penyesuaian bobot di seluruh lapisan model agar sesuai dengan karakteristik data spesifik penelitian ini.

b. Model SVM

Sebelum melatih SVM, teks yang telah diproses diubah menjadi representasi numerik menggunakan TF-IDF (*Term Frequency-Inverse Document Frequency*). TF-IDF menimbang pentingnya sebuah kata tidak hanya berdasarkan frekuensinya dalam sebuah dokumen tetapi juga seberapa jarang kata itu muncul di seluruh korpus. Model SVM kemudian dilatih menggunakan kernel linear yang dikenal efisien untuk data teks berdimensi tinggi.

c. Model *Random Forest*

Sama seperti SVM, model ini juga menggunakan representasi fitur TF-IDF. Model *Random Forest* dikonfigurasi dengan 100 decision tree untuk memastikan ensemble yang cukup beragam dan kuat.

Kinerja ketiga model dievaluasi secara kuantitatif pada testing set. Metrik evaluasi utama yang digunakan adalah *Accuracy*, *Precision*, *Recall*, dan *F1-Score*, yang dihitung dari confusion matrix.

a. *Accuracy*

Mengukur persentase prediksi yang benar secara keseluruhan [19].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Dengan keterangan sebagai berikut:

TP = *True Positive*

TN = *True Negative*

FP = *False Positive*

FN = *False Negative*

b. *Precision*

Mengukur tingkat keakuratan prediksi positif [20].

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Dengan keterangan sebagai berikut:

TP = *True Positive*

FP = *False Positive*

c. *Recall*

Mengukur kemampuan model untuk menemukan kembali semua sampel positif yang relevan [20].

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

Dengan keterangan sebagai berikut:

TP = *True Positive*

FN = *False Negative*

d. *F1-Score*

Rata-rata harmonik dari *Precision* dan *Recall*, memberikan ukuran keseimbangan antara keduanya [20].

$$F1\ Score = \frac{2 \cdot P \cdot R}{P + R} \quad (4)$$

Dengan keterangan sebagai berikut:

P = *Precision*

R = *Recall*

4. HASIL DAN PEMBAHASAN

4.1. Hasil Analisis Data

Dari total 54.105 *tweet* yang dianalisis, hasil pelabelan otomatis menunjukkan distribusi sentimen yang menarik. Proses pelabelan mengklasifikasikan

23.847 *tweet* (setara dengan 44,08%) sebagai sentimen positif, yang mencerminkan adanya dukungan dan apresiasi yang signifikan dari publik. Selanjutnya, sebanyak 22.306 *tweet* (41,23%) dikategorikan sebagai netral, yang umumnya berisi penyampaian informasi faktual, pertanyaan, atau diskusi objektif tanpa kecenderungan emosional yang kuat. Sementara itu, hanya sebagian kecil, yaitu 7.952 *tweet* (14,70%), yang teridentifikasi sebagai sentimen negatif. Distribusi ini secara jelas mengindikasikan bahwa diskursus publik di media sosial X mengenai Program Makan Bergizi Gratis cenderung didominasi oleh respons yang mendukung dan informatif. Volume kritik atau penolakan secara eksplisit berada pada porsi yang jauh lebih kecil, menunjukkan bahwa program ini memiliki penerimaan awal yang cukup baik di kalangan pengguna media sosial.

4.2. Hasil Kinerja Model

Evaluasi komparatif pada dataset pengujian menunjukkan perbedaan performa yang signifikan antara ketiga model. Hasil lengkapnya disajikan pada Tabel 1 divisualisasikan pada Gambar 5 dan 6.

Tabel 1. Perbandingan Kinerja Model

Model	Akurasi	Presisi	Recall	F1-Score
IndoBERT	87,60%	87,56%	87,60%	87,58%
SVM	83,39%	83,40%	83,39%	83,38%
Random Forest	70,28%	72,10%	70,28%	67,61%

a. Model IndoBERT

Dengan akurasi keseluruhan 87,60%, IndoBERT secara jelas menjadi model dengan performa terbaik. Keunggulannya terlihat konsisten di semua metrik. Secara khusus, model ini menunjukkan kinerja yang sangat kuat dalam mengidentifikasi sentimen netral (F1-score 0,90) dan positif (F1-score 0,88). Performa pada sentimen negatif sedikit lebih rendah (F1-score 0,79), yang mungkin disebabkan oleh jumlah sampel yang lebih kecil dan variasi ekspresi negatif yang lebih kompleks.

b. Model SVM

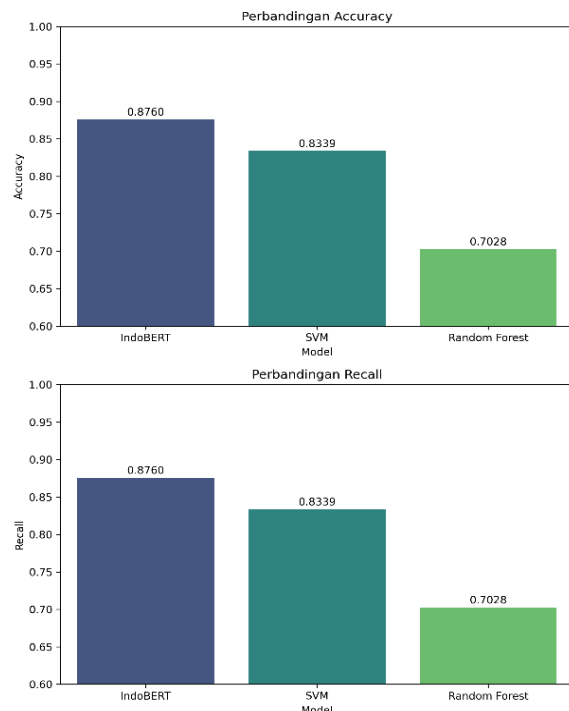
Menempati posisi kedua dengan akurasi 83,39%. SVM menunjukkan performa yang solid dan seimbang di semua kelas, menjadikannya baseline yang sangat kuat. Meskipun kalah dari IndoBERT, efisiensi komputasinya yang jauh lebih tinggi (waktu pelatihan hanya 3,53 menit dibandingkan 71,45 menit untuk IndoBERT) menjadikannya alternatif yang menarik.

c. Model *Random Forest*

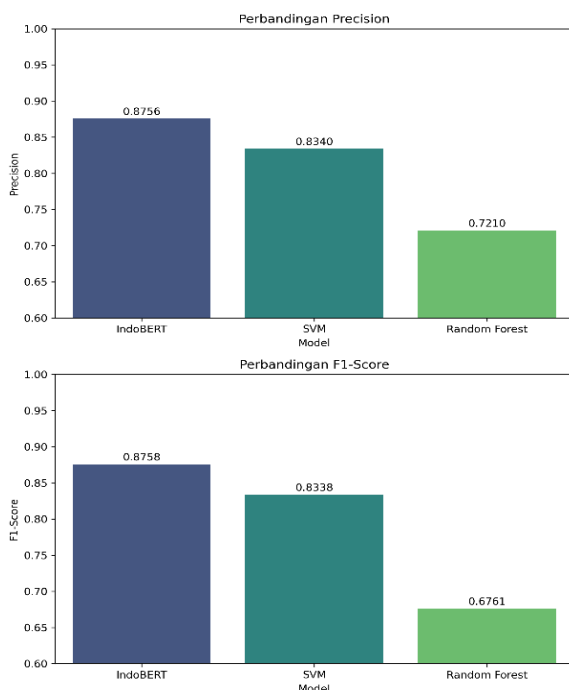
Menunjukkan kinerja yang paling lemah dengan akurasi 70,28%. Yang paling menonjol adalah kegagalannya dalam mengidentifikasi sentimen negatif, dengan nilai recall yang sangat rendah yaitu 0,18. Ini berarti model hanya berhasil menemukan 18% dari semua *tweet* negatif yang sebenarnya. Tingginya presisi (0,85) untuk kelas

ini menunjukkan bahwa ketika model memprediksi sesuatu sebagai negatif, prediksinya cenderung benar, tetapi ia sangat jarang melakukannya. Ketidakseimbangan ini membuat *Random Forest* tidak cocok untuk dataset ini.

Gambar 5 dan 6. Grafik Perbandingan Akurasi, *recall*, *precision* dan *F1-Score* dari Ketiga Model



Gambar 5. Grafik Perbandingan *Accuracy* dan *Recall* dari Ketiga Model



Gambar 6. Grafik Perbandingan *Precision* dan *F1-Score* dari Ketiga Model

4.3. Pembahasan

Hasil penelitian ini secara meyakinkan menunjukkan superioritas model berbasis *Transformer* (IndoBERT) dibandingkan dengan algoritma machine learning klasik (SVM dan *Random Forest*) untuk tugas analisis sentimen pada teks media sosial berbahasa Indonesia. Keunggulan IndoBERT dengan selisih akurasi 4,21% dari SVM bukanlah hal yang sepele. Perbedaan ini berasal dari kemampuan fundamental arsitektur *Transformer* untuk memodelkan dependensi jangka panjang dalam teks melalui mekanisme *self-attention*. Tidak seperti SVM atau *Random Forest* yang memperlakukan teks sebagai *bag-of-words* (kehilangan urutan kata), IndoBERT mampu memahami bahwa urutan dan posisi kata sangat mempengaruhi makna. Misalnya, IndoBERT dapat membedakan antara "program ini tidak baik" dan "tidak, program ini baik", sebuah nuansa yang sering kali hilang dalam model klasik.

Konsistensi temuan ini dengan penelitian lain seperti [12] dan [13] memperkuat argumen bahwa untuk tugas pemahaman bahasa alami yang kompleks, investasi komputasi pada model *deep learning* seperti IndoBERT memberikan hasil yang sepadan. Kemampuan model ini untuk memahami konteks bidirectional memungkinkannya menangkap fenomena linguistik yang rumit seperti ironi atau sarkasme dengan lebih baik, meskipun ini tetap menjadi tantangan.

Performa SVM yang solid (83,39%) menunjukkan bahwa algoritma klasik masih sangat relevan. Kemampuannya untuk bekerja dengan baik pada data teks berdimensi tinggi (setelah transformasi TF-IDF) dan efisiensinya menjadikannya pilihan yang pragmatis untuk aplikasi yang tidak menuntut akurasi tertinggi atau memiliki keterbatasan sumber daya. Di sisi lain, kegagalan *Random Forest* dalam penelitian ini, terutama pada kelas minoritas (negatif), menyoroti kelemahannya dalam menangani dataset yang tidak seimbang (*imbalanced dataset*) dan sangat *sparse*. Meskipun teknik seperti SMOTE (*Synthetic Minority Over-sampling Technique*) dapat membantu, hasil ini menunjukkan bahwa secara inheren, arsitektur berbasis pohon keputusan mungkin kurang cocok untuk data tekstual dibandingkan model berbasis margin seperti SVM atau model kontekstual seperti IndoBERT.

Dari perspektif kebijakan publik, distribusi sentimen yang didominasi oleh respons positif dan netral (total 85,91%) memberikan modal sosial yang kuat bagi pemerintah untuk melanjutkan program. Namun, sentimen negatif sebesar 14,3%, meskipun minoritas, tidak boleh diabaikan. Analisis konten menunjukkan bahwa kritik ini sangat terfokus dan substantif, yaitu pada masalah anggaran dan keberlanjutan. Hal ini sejalan dengan kekhawatiran yang diungkapkan oleh para ekonom dan lembaga keuangan [2]. Oleh karena itu, temuan ini menyiratkan bahwa strategi komunikasi pemerintah ke depan harus secara proaktif dan transparan menjawab

kekhawatiran ini. Menjelaskan secara rinci sumber pendanaan, mekanisme pengawasan, dan proyeksi dampak ekonomi jangka panjang dapat membantu meredakan skeptisisme publik dan membangun kepercayaan. Keterlibatan para pemangku kepentingan, seperti yang disarankan oleh [21], menjadi krusial untuk memastikan keberhasilan implementasi.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil melakukan analisis sentimen komprehensif terhadap Program Makan Bergizi Gratis di media sosial X dan mengevaluasi kinerja tiga algoritma klasifikasi yang berbeda. Berdasarkan analisis terhadap 54.105 tweet, dapat ditarik beberapa kesimpulan utama. Pertama, dari segi persepsi publik, program ini diterima dengan baik oleh mayoritas pengguna media sosial X, dengan sentimen positif (43,9%) dan netral (41,8%) secara kolektif mencapai 85,71%. Sentimen negatif yang ada (14,3%) cenderung terfokus pada isu substantif terkait anggaran dan keberlanjutan program. Kedua, dari segi metodologi, model IndoBERT menunjukkan superioritas yang jelas dengan akurasi 87,60%, membuktikan bahwa arsitektur *Transformer* dengan pemahaman konteks bidirectional sangat efektif untuk menganalisis teks media sosial berbahasa Indonesia yang kompleks. Kinerjanya secara signifikan melampaui SVM (83,39%) dan *Random Forest* (70,28%).

Berdasarkan kesimpulan tersebut, beberapa saran diajukan. Bagi peneliti selanjutnya, terdapat beberapa arah pengembangan yang menjanjikan. Pertama, memperluas cakupan analisis ke platform media sosial lain (seperti Facebook, Instagram, atau TikTok) untuk mendapatkan pandangan publik yang lebih holistik. Kedua, mengintegrasikan analisis sentimen dengan data demografis atau geografis yang lebih rinci untuk memetakan persepsi di berbagai segmen masyarakat. Ketiga, dari sisi teknis, eksplorasi lebih lanjut pada model *Transformer* yang lebih baru (misalnya, IndoBART atau GPT varian Indonesia) dan penerapan teknik penanganan data tidak seimbang pada algoritma klasik dapat menjadi fokus penelitian di masa depan untuk terus mendorong batas akurasi dalam analisis sentimen.

DAFTAR PUSTAKA

- [1] P. A. Maharani, A. R. Namira, Dan T. V. Chairunnisa, "Peran Makan Siang Gratis Dalam Janji Kampanye Prabowo Gibran Dan Realisasinya," *J. Law Soc. Soc.*, Vol. 1, No. 1, Hlm. 1–10, Jun 2024, Doi: 10.70656/Jolasos.V1i1.79.
- [2] A. A. Haikal Dan H. H. Anbiya, "Pengaruh Program Makan Siang Dan Susu Gratis Prabowo Gibran Terhadap Sektor Industri Manufaktur," *Natl. Conf. Electr. Inform. Ind. Technol. Neit*, Vol. 1, No. 1, Art. No. 1, Agu 2024.
- [3] Z. Purwanti Dan Sugiyono, "Pemodelan *Text Mining* Untuk Analisis Sentimen Terhadap Program Makan Siang Gratis Di Media Sosial X Menggunakan Algoritma *Support Vector Machine* (Svm)," *J. Indones. Manaj. Inform. Dan Komun.*, Vol. 5, No. 3, Art. No. 3, Sep 2024, Doi: 10.35870/Jimik.V5i3.1001.
- [4] - Annisa Ramadhani, "Analisis Sentimen Tanggapan Publik Terhadap Program Makan Siang Gratis Menggunakan Algoritma Nbc Dan Svm," *Build. Inform. Technol. Sci. Bits*, Vol. 6, No. No.3, Art. No. No.3, Des 2024.
- [5] P. K. Sari Dan R. R. Suryono, "Komparasi Algoritma *Support Vector Machine* Dan *Random Forest* Untuk Analisis Sentimen Metaverse," *J. Mnemon.*, Vol. 7, No. 1, Hlm. 31–39, 2024.
- [6] B. Wilie Dkk., "Indonlu: Benchmark And Resources For Evaluating Indonesian Natural Language Understanding," *Arxiv Prepr. Arxiv200905387*, 2020.
- [7] R. N. Ikhsani Dan F. F. Abdulloh, "Optimasi Svm Dan Decision Tree Menggunakan Smote Untuk Mengklasifikasi Sentimen Masyarakat Mengenai Pinjaman Online," *J. Media Inform. Budidarma*, Vol. 7, No. 4, Hlm. 1667–1677, 2023.
- [8] F. Fathonah Dan A. Herlina, "Text Mining Analisis Sentimen Mengenai Vaksin Covid-19 Menggunakan Metode Naive Bayes," 2021.
- [9] R. Robbi Nanda, "Klasifikasi Berita Menggunakan Metode *Support Vector Machine* (Svm)," *Klasifikasi Ber. Menggunakan Metode Support Vector Mach. Svm*, Vol. 5, No. 2, Hlm. 269–278, 2022.
- [10] K. Anwar, "Analisa Sentimen Pengguna Instagram Di Indonesia Pada Review Smartphone Menggunakan Naive Bayes," *Klik Kaji. Ilm. Inform. Dan Komput.*, Vol. 2, No. 4, Hlm. 148–155, 2022.
- [11] Y. Yang Dan X. Cui, "Bert-Enhanced Text Graph Neural Network For Classification," *Entropy*, Vol. 23, No. 11, Art. No. 11, Nov 2021, Doi: 10.3390/E23111536.
- [12] N. A. R. Putri Dan A. Ardiansyah, "Analisis Sentimen Terhadap Kemajuan Kecerdasan Buatan Di Indonesia Menggunakan Bert Dan Roberta," *J. Sains Dan Inform.*, Vol. 9, No. 2, Art. No. 2, Nov 2023.
- [13] M. N. Hidayat Dan R. Pramudita, "Analisis Sentimen Terhadap Pembelajaran Secara Daring Pasca Pandemi Covid-19 Menggunakan Metode Indobert," *Inf. Manag. Educ. Prof. J. Inf. Manag.*, Vol. 8, No. 2, Hlm. 161, Jan 2024, Doi: 10.51211/Imbi.V8i2.2719.
- [14] E. P. Mualim, "Pendekatan Oversampling Smote Untuk Imbalanced Dataset Aksara Lampung Dan Klasifikasi Menggunakan Svm," 2022.
- [15] A. Ilham, "Analisis Sentimen Masyarakat Terhadap Kesehatan Mental Pada Twitter Menggunakan Algoritme K-Nearest Neighbor,"

- Dipresentasikan Pada Prosiding Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (Senafti), 2023, Hlm. 539–547.
- [16] M. D. Al Fahreza, A. Luthfiarta, M. Rafid, Dan M. Indrawan, “Analisis Sentimen: Pengaruh Jam Kerja Terhadap Kesehatan Mental Generasi Z,” *J. Appl. Comput. Sci. Technol.*, Vol. 5, No. 1, Hlm. 16–25, 2024.
- [17] A. Nuraini, A. Faqih, G. Dwilestari, N. D. Nuris, Dan R. Narasati, “Analisis Sentimen Terhadap Review Aplikasi Brimo Di Google Play Store Menggunakan Algoritma Naïve Bayes,” *Jati J. Mhs. Tek. Inform.*, Vol. 7, No. 6, Hlm. 3661–3666, 2023.
- [18] J. Jumadi, D. Maylawati, L. Pratiwi, Dan M. Ramdhani, “Comparison Of Nazief-Adriani And Paice-Husk Algorithm For Indonesian Text Stemming Process,” Dipresentasikan Pada Iop Conference Series: Materials Science And Engineering, Iop Publishing, 2021, Hlm. 032044.
- [19] R. R. R. Arisandi, B. Warsito, Dan A. R. Hakim, “Aplikasi Naïve Bayes Classifier (Nbc) Pada Klasifikasi Status Gizi Balita Stunting Dengan Pengujian K-Fold Cross Validation,” *J. Gaussian*, Vol. 11, No. 1, Hlm. 130–139, 2022.
- [20] R. S. Passa, S. Nurmaini, Dan D. P. Rini, “Deteksi Tumor Otak Pada Magnetic Resonance Imaging Menggunakan YOLOv7,” *J. Ilm. Matrik*, Vol. 25, No. 2, Hlm. 116–121, 2023.
- [21] - Elsa Triningsih, “Analisis Sentimen Terhadap Program Makan Bergizi Gratis Menggunakan Algoritma Machine Learning Pada Sosial Media X,” *Build. Inform. Technol. Sci. Bits*, Vol. 6, No. 4, Art. No. 4, Jan 2025, Diakses: 20 Maret 2025. [Daring]. Tersedia Pada: <https://Repository.Uin-Suska.Ac.Id/86626/>