

Komparasi SVM dan IndoBERT dalam Klasifikasi Sentimen Program Makanan Bergizi Gratis

Shifatush Shafwah¹, Harun Al Azies^{1,2}

¹Program Studi Teknik Informatika, Fakultas Ilmu Komputer, Universitas Dian Nuswantoro, Semarang, Indonesia

²Research Center for Quantum Computing and Materials Informatics, Fakultas Ilmu Komputer,

Universitas Dian Nuswantoro, Semarang, Indonesia

harun.alazies@dsn.dinus.ac.id

Abstrak

Program Makanan Bergizi Gratis (MBG) memunculkan beragam respons masyarakat di media sosial, khususnya pada platform X (Twitter). Analisis sentimen diperlukan untuk memahami kecenderungan opini publik terhadap program tersebut. Penelitian ini membandingkan kinerja *Support Vector Machine* (SVM) dan IndoBERT dalam mengklasifikasikan sentimen positif dan negatif pada 2.674 tweet terkait MBG. Data diperoleh melalui web scraping dan diproses melalui tahapan cleaning, normalisasi teks, tokenisasi, serta pelabelan menjadi dua kelas sentimen. Ketidakseimbangan data ditangani menggunakan *Synthetic Minority Oversampling Technique* (SMOTE). Model SVM dilatih menggunakan representasi fitur TF-IDF, sedangkan IndoBERT dilatih melalui fine-tuning sebagai model transformer. Evaluasi performa dilakukan menggunakan *10-Fold Cross-Validation*, *confusion matrix*, ROC-AUC, dan uji statistik *paired t-test*. Hasil penelitian menunjukkan bahwa SVM memperoleh akurasi 94,64% dan *F1-Score* 94,63%, sedangkan IndoBERT mencapai akurasi 90,11% dan *F1-Score* 89,92%. Meskipun IndoBERT mencatat nilai AUC sedikit lebih tinggi, kinerja keseluruhan SVM lebih unggul secara konsisten pada data yang telah diseimbangkan dengan SMOTE. Uji *paired t-test* menghasilkan nilai $p < 0,05$, yang menunjukkan bahwa perbedaan performa kedua model bersifat signifikan. SVM lebih efektif digunakan untuk klasifikasi sentimen dua kelas pada dataset MBG yang relatif kecil dan bersifat informal.

Kata kunci: Analisis Sentimen; Program Makanan Bergizi Gratis; SVM; IndoBERT; SMOTE; Media Sosial.

Abstract

The free nutritious meal program (MBG) has generated diverse public responses on social media, particularly on platform X (Twitter). Sentiment analysis is needed to understand public opinion trends towards the program. This study compares the performance of *Support Vector Machine* (SVM) and IndoBERT in classifying positive and negative sentiments in 2,674 tweets related to MBG. Data was obtained through web scraping and processed through cleaning, text normalisation, tokenisation, and labelling stages into two sentiment classes. Data imbalance was handled using the *Synthetic Minority Oversampling Technique* (SMOTE). The SVM model was trained using TF-IDF feature representation, while IndoBERT was trained through fine-tuning as a transformer model. Performance evaluation was conducted using *10-fold cross-validation*, a *confusion matrix*, the ROC-AUC, and a *paired t-test* statistical test. The results showed that the SVM achieved 94.64% accuracy and 94.63% *F1-Score*, while IndoBERT achieved 90.11% accuracy and 89.92% *F1-Score*. Although IndoBERT recorded a slightly higher AUC, the overall performance of SVM consistently outperformed the SMOTE-balanced dataset. A *paired t-test* yielded a p -value < 0.05 , indicating a significant difference in performance between the two models. SVM was more effective for two-class sentiment classification on the relatively small and informal MBG dataset.

Keywords: Sentiment Analysis; The free nutritious meal program; SVM; IndoBERT; SMOTE; Social Media

I. PENDAHULUAN

Perkembangan teknologi digital dan media sosial telah mengubah cara orang menyampaikan pendapat

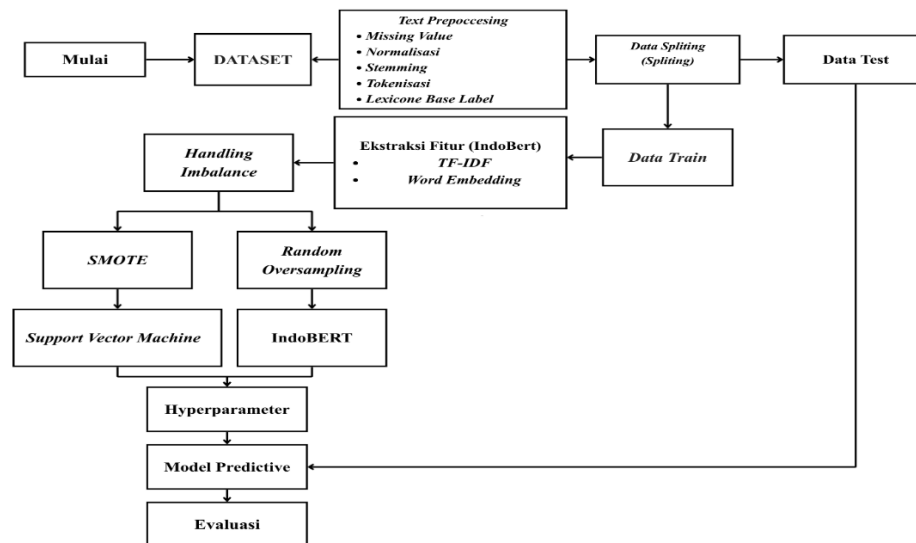
tentang aturan pemerintah. Salah satu topik yang banyak dibicarakan secara online adalah program Makanan Bergizi Gratis yang diadakan pemerintah Indonesia sebagai bantuan sosial untuk meningkatkan gizi anak-anak di sekolah. Tetapi, pelaksanaan aturan ini menimbulkan berbagai reaksi dari masyarakat

yang menunjukkan pandangan baik, buruk, dan netral tentang keberhasilan dan keterbukaan program itu. Banyaknya percakapan di platform X (Twitter) membuat media sosial sekarang berfungsi sebagai sumber informasi penting untuk memahami pandangan masyarakat secara luas dan langsung. Ini membuat perlunya analisis mendalam tentang pandangan masyarakat untuk memahami arah pendapat masyarakat tentang program sosial yang dimiliki pemerintah itu [1].

Analisis sentimen adalah cara yang memakai *Natural Language Processing* (NLP) untuk otomatis menemukan perasaan atau pendapat dalam suatu tulisan. Dengan cara ini, para peneliti bisa menggali kecenderungan pendapat masyarakat tentang suatu masalah tanpa perlu melakukan survei biasa, yang seringkali membutuhkan waktu dan biaya yang besar. Proses analisis biasanya dimulai dengan mengumpulkan data dan mempersiapkan tulisan, yang meliputi membersihkan, membagi, dan mengatur tulisan agar data itu siap untuk dipakai di dalam model pengelompokan [2]. Cara ini sangat penting untuk menilai pandangan masyarakat tentang aturan pemerintah seperti program Makanan Bergizi Gratis. Oleh karena itu, analisis sentimen menjadi pilihan hitungan yang efisien untuk menilai bagaimana masyarakat menerima program pemerintah tersebut. Hasil analisis sentimen terhadap kebijakan pemerintah dalam berbagai penelitian menunjukan variabilitas yang signifikan di mana temuan sangat bergantung pada metode pemodelan dan representasi data teks yang digunakan. Sebelumnya, penerapan SVM dengan TF-IDF pada 2.682 tweet mendominasi sentimen positif meski terdapat masalah ketidakseimbangan data [3]. Sebaliknya, IndoBERT justru mengidentifikasi sentimen negatif terkait efektivitas program. Fenomena ini mempertegas pentingnya pemilihan algoritma dan teknik *preprocessing* yang tepat. Penelitian yang dilakukan oleh Alfatah [4]

menerapkan model *Transformer* IndoBERT untuk menganalisis sentimen pada data Twitter berbahasa Indonesia, yang terbukti secara signifikan mengungguli algoritma konvensional seperti *Naïve Bayes* dan *Support Vector Machine* (SVM) dengan mencapai akurasi sebesar 87%, *precision* 86%, serta *F1-score* 85,5%. Berdasarkan penelitian sebelumnya mengkonfirmasi bahwa optimasi SVM dengan SMOTE meningkatkan akurasi klasifikasi sentimen terhadap pemerintah. Studi tersebut juga mengungkap perbedaan platform TikTok didominasi narasi positif dan netral, sedangkan X dipenuhi opini negatif [5]. Pentingnya kombinasi model dan teknik penyeimbangan data untuk kinerja optimal ini juga diperkuat, di mana penerapan IndoBERT yang dikombinasikan dengan teknik *Random Oversampling* berhasil mencapai akurasi 90% dalam analisis sentimen ulasan aplikasi BRImo, sekaligus mengungkap dominasi keluhan teknis dalam sentimen negative [6].

Studi ini dilakukan dengan tujuan untuk menguatkan penelitian sebelumnya dengan menganalisis perbandingan antara model pembelajaran mesin konvensional, *Support Vector Machine* (SVM), dan model yang menggunakan transform, IndoBERT [7]. Pilihan kedua model ini didasarkan pada perbedaan cara kerja: SVM memanfaatkan fitur numerik seperti TF-IDF, sedangkan IndoBERT memiliki pemahaman yang lebih mendalam tentang konteks bahasa Indonesia. Melalui analisis ini, penelitian ini berusaha untuk menemukan cara yang paling efektif dalam mengelompokkan sentimen masyarakat mengenai program Makanan Bergizi Gratis, sehingga dapat membantu dalam analisis opini publik dan menjadi referensi bagi pembuat kebijakan dalam memahami perspektif masyarakat dengan lebih objektif serta mendukung kemajuan penelitian NLP di Indonesia [8].



Gambar 1. Diagram Alur Penelitian Sentimen MBG

II. METODE PENELITIAN

Penelitian ini memiliki tujuan untuk mengkaji emosi dari Program Makanan Bergizi Gratis dengan mengintegrasikan dua model SVM dan IndoBERT. Langkah-langkah utama mencakup pengumpulan informasi, persiapan data, tokenisasi, serta pengambilan fitur dengan bantuan IndoBERT, pembagian data penyeimbangan kelas melalui SMOTE, dan klasifikasi menggunakan IndoBERT serta SVM. Gambar 1 merupakan alur dari penelitian.

A. Pengumpulan Data

Dataset yang dimanfaatkan ialah Kumpulan cuitan mengenai program penyediaan makanan sehat tanpa biaya yang didapat dari platfrom media social X dengan kata kunci “Program Makan Bergizi Gratis (MBG)” dan memiliki total 2.682 entri. Setelah dilakukan proses pembersihan termasuk penghapusan nilai yang hilang dan kolom teks yang berulang jumlah data menjadi 2.674.

B. Preprocessing Data

Setelah mengumpulkan informasi dari platform media sosial, Langkah pra-pemrosesan dilakukan untuk memastikan data siap digunakan dalam pembentukan model. Proses ini terdiri dari beberapa Langkah penting, termasuk penanganan nilai yang hilang untuk menghapus atau mengisi informasi yang tidak ada agar tidak mempengaruhi hasil analisis, serta pengkodean untuk mentransformasi data kategorial menjadi format numerik yang dapat dipahami oleh model pembelajaran mesin [9]. Selain itu, data teks juga menjalani tahap pembersihan

dengan menghapus karakter khusus, angka, emoji, tautan, dan tanda baca yang tidak relevan.

C. Ekstraksi Fitur

Tahapan ekstraksi fitur dilakukan untuk mengubah data teks hasil *preprocessing* menjadi bentuk numerik yang dapat dipahami oleh model. Dalam penelitian ini digunakan metode TF-IDF dan *embedding* IndoBERT [10]. TF-IDF memberikan bobot pada kata berdasarkan frekuensinya dalam dokumen dan keseluruhan korpus, sedangkan *embedding* IndoBERT menghasilkan representasi vektor yang memperhatikan konteks antar kata. Dengan tahapan ini, data teks menjadi lebih terstruktur dan siap digunakan dalam proses pelatihan model klasifikasi sentimen.

D. Data Splitting (Pembagian Data)

Setelah menjalani tahap konversi menjadi bentuk *embedding*, dataset yang telah melalui *preprocessing* dibagi menjadi dua segmen utama, yaitu data untuk pelatihan dan data untuk pengujian [11]. Pemisahan ini dilakukan dengan proporsi 80:20 untuk menciptakan keseimbangan dalam proses pembelajaran serta evaluasi dari model. Rasio tersebut memastikan model menerima data terpisah untuk mengetes performanya. Selain itu, pemisahan dilakukan dengan menjaga keseimbangan distribusi antara kelas, sehingga setiap kategori sentiment bisa terwakili secara adil dalam proses pelatihan dan pengujian [12].

E. Penyeimbangan Data

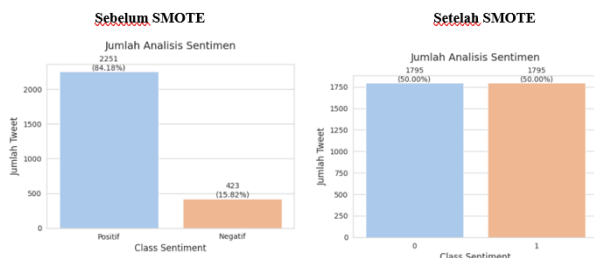
1) SMOTE

Untuk mengatasi ketidakseimbangan data, penelitian ini menggunakan teknik sampel lebih banyak sampel sintetis minoritas SMOTE. Ketidakseimbangan ini dapat mengurangi ketetapan model. Setelah proses pembagian dataset, SMOTE diterapkan pada data latih untuk menghasilkan sampel sintetis dari kelas minoritas melalui interpolasi. Metode ini meningkatkan representasi kelas negatif dan mengurangi bias model selama pelatihan [13].

SMOTE bekerja melalui proses pembuatan sampel sintetis dari kelas minoritas dengan melakukan interpolasi di antara titik data. Pada tahap ini, algoritma memilih sebuah data minoritas (x_i) dan salah satu terdekatnya (x_j) menggunakan pendekatan *k-nearest neighbors*. Berdasarkan kedua data tersebut, sebuah sampel baru (x_{new}) dibentuk menggunakan persamaan berikut:

$$x_{new} = x_i + \lambda (x_j - x_i) \quad (1)$$

Dalam rumus tersebut, x_i merujuk pada data asli dari kelas yang lebih sedikit, x_j adalah terdekatnya, dan λ adalah nilai acak antara 0 dan 1 yang menentukan posisi titik sintetis.



Gambar 2. Perbandingan Distribusi Kelas Sebelum dan Sesudah SMOTE

Metode SMOTE tidak hanya menambah jumlah sampel pada kelas minoritas, tetapi juga menghasilkan variasi sintetis yang membantu model mengenali pola secara lebih efektif. Pendekatan ini mampu mengurangi dominasi kelas mayoritas dan meningkatkan akurasi klasifikasi pada dataset yang tidak seimbang [14]. seperti yang terlihat pada Gambar 2.

2) Random Oversampling

Metode *Random Oversampling* merupakan pendekatan penyeimbangan dataset yang mengatasi masalah ketidakseimbangan kelas dengan mereplikasi sampel acak dari kelas minoritas [15]. Metode ini bekerja dengan mereplikasi instance kelas minoritas hingga mencapai keseimbangan proporsional dengan kelas mayoritas, di mana sampel mayoritas memiliki

N dan M sampel minoritas ($M < N$), maka dilakukan penambahan, $(N - M)$ sampel melalui duplikasi acak sesuai persamaan $x_{\{new\}} = x_i$ dan x_i merepresentasikan sampel terpilih dari kelas minoritas. Mekanisme ini secara efektif mencegah bias model terhadap kelas mayoritas selama proses pelatihan.

F. Pengembangan Model dengan IndoBert

Setelah tahap pra pengolahan selesai, langkah berikutnya adalah menciptakan model klasifikasi sentiment dengan menggunakan IndoBERT dan *Support Vector Machine* (SVM). IndoBERT dipilih karena merupakan model transformer khusus Bahasa Indonesia yang dapat menghasilkan representasi teks dengan konteks yang mendalam lewat mekanisme *self-attention*, sehingga lebih baik dalam memahami hubungan semantik antar kata [16]. Proses pelatihan dilakukan dengan *fine-tuning*, yang meliputi penyesuaian parameter seperti *learning rate*, jumlah epoch, dan ukuran batch untuk mencapai performa optimal pada data yang sudah diproses [17].

Dalam penelitian ini, model dikembangkan secara eksperimental menggunakan lingkungan pemrograman Python dan dilatih dengan data latih yang sudah diseimbangkan melalui *Random Oversampling*. Metode ini diharapkan dapat menghasilkan klasifikasi sentimen yang lebih akurat dan adil, khususnya dalam mengidentifikasi kelas minoritas yang sebelumnya kurang terwakili.

G. Evaluasi Model

Setelah proses pelatihan selesai, model dievaluasi menggunakan data uji untuk mengukur kemampuan dalam melakukan klasifikasi sentimen. Evaluasi ini dilakukan menggunakan beberapa metrik.

$$Akurasi = \frac{TP+TN}{TP+TN+FP+FN} \quad (2)$$

$$Precision = \frac{TP}{TP+FP} \quad (3)$$

$$Recall = \frac{TP}{TP+FN} \quad (4)$$

$$F1 - score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (5)$$

True Positive (TP) adalah jumlah data yang sebenarnya positif dan berhasil dikenali sebagai positif oleh model, menandakan bahwa model telah melakukan klasifikasi data dengan benar. *True Negative* (TN) mencerminkan jumlah data yang benar-benar negatif dan juga diakui sebagai negatif, yang menunjukkan bahwa model tidak melakukan kesalahan dalam mengenali kelas tersebut [18], [19].

Sebaliknya, *False Positive* (FP) terjadi saat model secara keliru menganggap data yang sebenarnya negatif sebagai positif, yang bisa mengakibatkan kesalahan dalam memahami hasil, seperti pada

analisis sentimen ketika ulasan yang netral atau negative dianggap positif. *False Negative* (FN) adalah kondisi di mana model memprediksi data sebagai negatif padahal sebenarnya itu positif, yang dapat menyebabkan hilangnya informasi penting, misalnya dalam kasus ulasan positif yang tidak terdeteksi dengan tepat oleh model.

III. HASIL DAN PEMBAHASAN

A. Deskripsi Dataset dan Pra-Pemrosesan

Dataset yang digunakan dalam penelitian ini diperoleh dari pengambilan data komentar pengguna pada aplikasi X mengenai program Makanan Bergizi jumlah keseluruhan data yang berhasil dikumpulkan mencapai 2.682 komentar, seperti yang terlihat pada Gambar 3.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 2682 entries, 0 to 2681
Data columns (total 15 columns):
#   Column                                Non-Null Count  Dtype
---  -
0   conversation_id_str                    2682 non-null   int64
1   created_at                            2682 non-null   object
2   favorite_count                        2682 non-null   int64
3   full_text                             2682 non-null   object
4   id_str                                2682 non-null   int64
5   image_url                             1702 non-null   object
6   in_reply_to_screen_name                335 non-null   object
7   lang                                  2682 non-null   object
8   location                              0 non-null      float64
9   quote_count                           2682 non-null   int64
10  reply_count                           2682 non-null   int64
11  retweet_count                          2682 non-null   int64
12  tweet_url                             2682 non-null   object
13  user_id_str                           2682 non-null   int64
14  username                              0 non-null      float64
```

Gambar 3. Dataset Program Makanan Bergizi

Setelah dilakukan pemeriksaan kualitas data, ditemukan adanya *missing value* serta duplikasi data, sehingga proses pembersihan data perlu dilakukan untuk memastikan kualitas dataset yang digunakan dalam pemodelan [20]. Proses pembersihan ini menghasilkan dataset akhir sebanyak 2.674 data valid. Dataset terdiri dari komentar teks berbahasa Indonesia dengan dua kategori sentimen yaitu positif dan negatif, dengan distribusi positif 2.251 data dan negatif 423 data. Distribusi tersebut mengindikasikan adanya ketidakseimbangan yang signifikan [21]. Untuk menjawab kemungkinan adanya bias dalam model terhadap kelas yang dominan, penelitian ini menerapkan *class weight* pada fungsi *loss* saat pelatihan model [22].

B. Hasil Preprocessing

Tahapan pra-pemrosesan telah berhasil mengubah data teks mentah menjadi data yang siap untuk diproses lebih lanjut. Serangkaian langkah yang dilakukan meliputi pengubahan format huruf, pembersihan karakter yang tidak diperukan (seperti URL, emoji, dan symbol), penyesuaian istilah yang

tidak baku, penghapusan kata umum, dan pengurangan bentuk kata, yang terbukti meningkatkan mutu korpus teks secara signifikan [23]. Implementasi yang dilakukan menunjukkan bahwa penyesuaian istilah berhasil mengubah kata-kata tidak baku seperti “gk”, “ga”, “bgt”, dan “udh” ke bentuk baku, sehingga kesesuaian dengan kosakata Bahasa Indonesia yang dikenali oleh model dapat tercipta [24].

Di samping itu, pengurangan bentuk kata terbukti efektif dalam mengurangi variabilitas morfologis dengan mengubah kata berawalan menjadi bentuk dasar [25]. Pengurangan kompleksitas ini menghasilkan representasi token yang lebih seragam dan efisien, memungkinkan dimensi fitur diperkecil tanpa menghilangkan informasi makna yang penting.

Seperti yang terlihat pada Gambar 4 menunjukkan perbandingan distribusi kata sebelum dan sesudah dilakukan tahap pra-pemrosesan.



Gambar 4. Distribusi kata sebelum (a) dan sesudah (b) dilakukan tahap pra-pemrosesan

C. Model Train

Pada tahap pemodelan, penelitian ini memanfaatkan algoritma *transformer*-based IndoBERT untuk melakukan klasifikasi sentimen pada komentar berbahasa Indonesia [26]. Hasil analisis awal memperlihatkan adanya ketidakseimbangan distribusi antara kelas positif dan negatif, yang berpotensi menurunkan kinerja model akibat bias terhadap kelas mayoritas [27]. Untuk mengatasi permasalahan tersebut, dilakukan teknik *Random Oversampling* guna menambah jumlah sampel pada kelas minoritas sehingga proporsi kedua

kelas menjadi lebih seimbang dan model dapat mempelajari pola secara lebih representatif [28]. Seperti pada Tabel 1 yang menunjukkan hasil perbandingan *K-FOLD*.

Hasil evaluasi dengan validasi silang *K-Fold* sebagaimana terlihat pada Tabel 1 menunjukkan bahwa, meskipun IndoBERT telah dioptimalkan dengan penanganan ketidakseimbangan data, model SVM justru mencatat kinerja yang lebih unggul secara konsisten di seluruh metrik evaluasi.

Tabel 1. Hasil Perbandingan K-FOLD

Model	Accuracy	F1-Score	Precision	Recall
IndoBERT	0.90113	0.9009	0.90647	0.90111
SVM	0.94640	0.9463	0.94910	0.9464

Akurasi yang dicapai SVM mencapai 94,64%, melampaui IndoBERT yang sebesar 90,11%. Hasil ini mengindikasikan bahwa dalam konteks data dan skenario ini, pendekatan SVM mampu memberikan hasil klasifikasi yang lebih andal dan komprehensif dibandingkan model berbasis *transformer* yang lebih kompleks.

D. Model Test

Evaluasi kinerja model menggunakan validasi silang 5-fold mengungkapkan bahwa model *Support Vector Machine* (SVM) menunjukkan performa akurasi yang lebih unggul secara konsisten pada setiap fold apabila dibandingkan dengan model IndoBERT. Pola keunggulan ini tergambar stabil, dengan selisih akurasi yang berkisar antara 0.04 hingga 0.05 di seluruh proses validasi. Seperti pada Tabel 2 yang menunjukkan hasil data perbandingan per Fold.

Tabel 2. Data Perbandingan Per Fold

Fold	SVM Accuracy	IndoBERT Accuracy	Selisih (SVM–IndoBERT)
1	0.945908	0.905687	0.040221
2	0.936111	0.886111	0.050000
3	0.941667	0.894444	0.047223
4	0.943056	0.895833	0.047223
5	0.965278	0.923611	0.041667

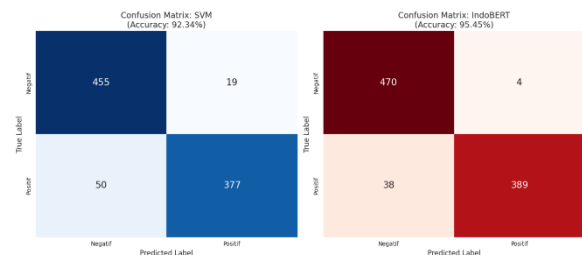
Keunggulan konsisten tersebut kemudian dikonfirmasi lebih lanjut melalui perhitungan rata-rata akurasi akhir. Analisis menunjukkan bahwa akurasi rata-rata SVM mampu mencapai 94.64%, sementara IndoBERT berada pada angka 90.11%. SVM unggul signifikan dengan selisih akhir sebesar 4.53%, memperkuat temuan awal mengenai efektivitasnya pada dataset yang digunakan.

Tabel 3. Hasil Uji Statistik (*Paired t-Test*)

Parameter	Nilai
T-Statistic	24.4640
P-Value	0.00002

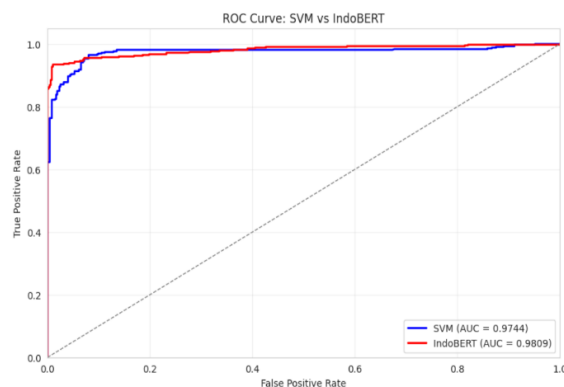
Pada Tabel 3 Uji *paired t-test* menunjukkan nilai *t-statistic* sebesar 24.4640 dan *p-value* sebesar 0.00002. Nilai *p-value* yang jauh di bawah ambang signifikansi 0.05 menunjukkan bahwa perbedaan akurasi antara model SVM dan IndoBERT bersifat signifikan secara statistik. Keunggulan SVM pada evaluasi *k-fold* tidak terjadi secara acak, tetapi merupakan efek nyata dari perbedaan performa model.

Hasil evaluasi pada Gambar 5 tersebut menunjukkan perbandingan kinerja model SVM dan IndoBERT dalam klasifikasi sentimen dua kelas. Pada model SVM, sebanyak 455 data negatif dan 377 data positif berhasil diprediksi dengan benar, sementara 19 data negatif salah diprediksi sebagai positif dan 50 data positif salah diklasifikasikan sebagai negatif, menghasilkan akurasi total 92,34%.



Gambar 5. Confusion Matrix SVM dan IndoBERT

Sebaliknya, IndoBERT menunjukkan performa yang lebih unggul dengan 470 prediksi benar pada kelas negatif dan 389 pada kelas positif, serta hanya menghasilkan 4 *false positive* dan 38 *false negative*, sehingga mencapai akurasi 95,45%. Perbedaan ini mengindikasikan bahwa IndoBERT memiliki kemampuan pemahaman konteks yang lebih baik dibandingkan SVM, terutama dalam menangkap nuansa linguistik pada teks berbahasa Indonesia. Dengan demikian, IndoBERT terbukti lebih efektif dalam mengurangi kesalahan klasifikasi dan memberikan hasil prediksi yang lebih akurat pada tugas analisis sentimen.



Gambar 6. ROC Curve: SVM vs IndoBERT

Berdasarkan analisis kurva ROC pada Gambar 6 yang disajikan, dapat diobservasi bahwa kedua model SVM dan IndoBERT menunjukkan performa diskriminasi yang sangat kuat, sebagaimana tercermin dari nilai AUC (Area Under Curve) yang mendekati 1. Kurva ROC IndoBERT (AUC = 0.9809) terletak sedikit di atas kurva ROC SVM (AUC = 0.9744), mengindikasikan kemampuan klasifikasi yang sedikit lebih unggul dalam memisahkan kelas positif dan negatif. Meskipun demikian, selisih AUC yang sangat kecil, yaitu hanya 0.0065, menandakan bahwa dari sudut pandang kurva ROC, kedua model memiliki kapabilitas yang setara. Kedua kurva tersebut berada di sudut kiri atas, yang merefleksikan kombinasi *True Positive Rate* (TPR) yang tinggi dan *False Positive Rate* (FPR) yang rendah, sebuah karakteristik ideal dari model klasifikasi yang andal. Secara visual, kedekatan kedua kurva ini memperkuat temuan sebelumnya bahwa meskipun terdapat perbedaan dalam metrik akurasi, kedua model tetap konsisten dalam kinerja secara keseluruhan.

IV. KESIMPULAN

Penelitian ini mengevaluasi efektivitas teknik pra-pemrosesan tekstual dan penanganan ketidakseimbangan data terhadap kinerja model *Support Vector Machine* (SVM) dan IndoBERT dalam analisis sentimen Program Makanan Bergizi Gratis, dan hasilnya menunjukkan bahwa SVM memiliki performa klasifikasi lebih unggul dengan akurasi 94,64% dibandingkan IndoBERT (90,11%) yang dikonfirmasi melalui uji statistik *paired t-test*. Temuan tersebut mengungkapkan bahwa meskipun IndoBERT memiliki kemampuan diskriminatif lebih baik berdasarkan nilai AUC 0,9809, SVM menunjukkan stabilitas dan efisiensi yang lebih tinggi dalam konteks dataset berbahasa Indonesia, membuktikan bahwa kompleksitas model tidak selalu berkorelasi positif dengan akurasi klasifikasi. Pengembangan selanjutnya dapat diarahkan pada

eksplorasi teknik augmentasi data yang lebih variatif, implementasi metode *embedding* kontekstual untuk bahasa informal Indonesia, serta perluasan cakupan dataset untuk meningkatkan kemampuan model dalam menangkap nuansa sentimen masyarakat terhadap kebijakan publik.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih sebesar-besarnya kepada Universitas Dian Nuswantoro atas dukungan dan fasilitas yang diberikan selama pelaksanaan penelitian ini. Ucapan terima kasih yang tulus juga penulis sampaikan kepada dosen pembimbing atas bimbingan, arahan, dan ilmu yang sangat berharga, serta kepada keluarga dan rekan-rekan atas dukungan moral dan semangat yang tak ternilai hingga terselesaikannya penelitian ini.

REFERENSI

- [1] N. Harahap, “Analisis Framing Netizen Di X Dalam Pelaksanaan Program Makan Bergizi Gratis,” vol. 4, 2025.
- [2] H. Hairani, “Peningkatan Kinerja Metode SVM Menggunakan Metode KNN Imputasi dan K-Means-Smote untuk Klasifikasi Kelulusan Mahasiswa Universitas Bumigora,” *J. Teknol. Inf. Dan Ilmu Komput.*, vol. 8, no. 4, pp. 713–718, July 2021, doi: 10.25126/jtiik.2021843428.
- [3] W. Anggriyani and M. Fakhriza, “Analisis Sentimen Program Makan Gratis Pada Media Sosial X Menggunakan Metode NLP,” *J. Comput. Syst. Inform. JoSYC*, vol. 5, no. 4, pp. 1033–1042, Aug. 2024, doi: 10.47065/josyc.v5i4.5826.
- [4] D. Alfatah, “Penerapan Model Transformer Untuk Deteksi Sentimen Pada Data Twitter Berbahasa Indonesia”.
- [5] D. Pateman, T. F. Prasetyo, and H. Sujadi, “SENTIMENT ANALYSIS OF GOVERNMENT ON TIKTOK AND X PLATFORMS WITH SVM AND SMOTE APPROACH,” *JITK J. Ilmu Pengetah. Dan Teknol. Komput.*, vol. 10, no. 4, pp. 900–908, June 2025, doi: 10.33480/jitk.v10i4.6645.
- [6] A. A. P. Simarmata and T. B. Sasongko, “Sentiment Analysis on BRImo Application Reviews Using IndoBERT,” vol. 9, no. 3.
- [7] Mutiara Sintia Dewi and A. H. Hasugian, “PENERAPAN SUPPORT VECTOR MACHINE UNTUK ANALISIS SENTIMEN PADA TANGGAPAN MASYARAKAT DI MEDIA SOSIAL TERHADAP PROGRAM MAKAN SIANG GRATIS,” *Rabit J. Teknol. Dan*

- Sist. Inf. Univrab*, vol. 10, no. 2, pp. 911–922, July 2025, doi: 10.36341/rabit.v10i2.6425.
- [8] F. Fatkhurrohmah, B. I. Nugroho, and N. Fadillah, “Analisis Sentimen Program Makan Bergizi Gratis Pemerintah RI Melalui Twitter Menggunakan Metode SVM,” *RIGGS J. Artif. Intell. Digit. Bus.*, vol. 4, no. 3, pp. 3906–3917, Aug. 2025, doi: 10.31004/riggs.v4i3.2533.
- [9] T. Thomas and E. Rajabi, “A systematic review of machine learning-based missing value imputation techniques,” *Data Technol. Appl.*, vol. 55, no. 4, pp. 558–585, Aug. 2021, doi: 10.1108/DTA-12-2020-0298.
- [10] I. A. Wisky, S. Defit, and G. W. Nurcahyo, “Development of extraction features for Detecting Adolescent Personality with Machine Learning Algorithms,” *JOIV Int. J. Inform. Vis.*, vol. 8, no. 3–2, p. 1606, Nov. 2024, doi: 10.62527/joiv.8.3-2.3091.
- [11] “Jurnal Zonasi Marwika Rifattul Iffa - MARWIKARIFATTULIFFA Teknik Informatika.”
- [12] T. Safitri, Y. Umaidah, and I. Maulana, “Analisis Sentimen Pengguna Twitter Terhadap BTS Menggunakan Algoritma Support Vector Machine,” vol. 7, no. 1.
- [13] S. P. Octarini, A. Y. Zakiyyah, and K. Purwandari, “Comparison of IndoBERT and SVM Algorithm to Perform Aspect Based Sentiment Analysis using Hierarchical Dirichlet Process”.
- [14] M. Mukherjee and M. Khushi, “SMOTE-ENC: A Novel SMOTE-Based Method to Generate Synthetic Data for Nominal and Continuous Features,” *Appl. Syst. Innov.*, vol. 4, no. 1, p. 18, Mar. 2021, doi: 10.3390/asi4010018.
- [15] I. A. Rahma and L. H. Suadaa, “Penerapan Text Augmentation untuk Mengatasi Data yang Tidak Seimbang pada Klasifikasi Teks Berbahasa Indonesia,” *J. Teknol. Inf. Dan Ilmu Komput.*, vol. 10, no. 6, pp. 1329–1340, Dec. 2023, doi: 10.25126/jtiik.2023107325.
- [16] M. F. Kono, I. N. Fajri, and Y. Pristyanto, “Public Sentiment Analysis on Corruption Issues in Indonesia Using IndoBERT Fine-Tuning, Logistic Regression, and Linear SVM,” *Logist. Regres.*, vol. 9, no. 5.
- [17] A. Wafda, “Aspect-Based Sentiment Analysis terhadap Cuitan Platform X tentang Kurikulum Merdeka Menggunakan IndoBERT”.
- [18] S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad, M. Chenaghlu, and J. Gao, “Deep Learning Based Text Classification: A Comprehensive Review,” Jan. 04, 2021, *arXiv:arXiv:2004.03705*. doi: 10.48550/arXiv.2004.03705.
- [19] M. Kumar, L. Khan, and H.-T. Chang, “Evolving techniques in sentiment analysis: a comprehensive review,” *PeerJ Comput. Sci.*, vol. 11, p. e2592, Jan. 2025, doi: 10.7717/peerj-cs.2592.
- [20] A. N. Ma’aly, D. Pramesti, A. D. Fathurahman, and H. Fakhurroja, “Exploring Sentiment Analysis for the Indonesian Presidential Election Through Online Reviews Using Multi-Label Classification with a Deep Learning Algorithm,” *Information*, vol. 15, no. 11, p. 705, Nov. 2024, doi: 10.3390/info15110705.
- [21] Z. Sitorus, M. Iqbal, D. Nasution, and R. F. Wijaya, “Penerapan Deep Learning dan Analisis Sentimen terhadap Gap Kompetensi Lulusan Lembaga Pendidikan dan Pelatihan Vokasi terhadap Dunia Kerja dengan Metode Long Short-Term Memory (LSTM),” vol. 6, no. 2, 2025.
- [22] M. A. Alfaridzy, E. Haerani, and L. Oktavia, “KLASIFIKASI SENTIMEN MASYARAKAT TERHADAP EFISIENSI ANGGARAN PEMERINTAH MENGGUNAKAN METODE NAÏVE BAYES CLASSIFIER,” vol. 5, 2024.
- [23] J. Khan, K. Ahmad, S. K. Jagatheesaperumal, and K.-A. Sohn, “Textual variations in social media text processing applications: challenges, solutions, and trends,” *Artif. Intell. Rev.*, vol. 58, no. 3, p. 89, Jan. 2025, doi: 10.1007/s10462-024-11071-z.
- [24] M. I. Raif, N. N. Hidayati, and T. Matulatan, “Otomatisasi Pendeteksi Kata Baku Dan Tidak Baku Pada Data Twitter Berbasis KBBI,” *J. Teknol. Inf. Dan Ilmu Komput.*, vol. 11, no. 2, pp. 337–348, Apr. 2024, doi: 10.25126/jtiik.20241127404.
- [25] V. E. Priyanti, “ANALISIS SENTIMEN PADA ULASAN APLIKASI SMART CITY MENGGUNAKAN PEMODELAN LATENT DIRICHLET ALLOCATION DAN KLASIFIKASI NAÏVE BAYES”.
- [26] Y. Tay, M. Dehghani, D. Bahri, and D. Metzler, “Efficient Transformers: A Survey,” *ACM Comput. Surv.*, vol. 55, no. 6, pp. 1–28, June 2023, doi: 10.1145/3530811.
- [27] W. Chen, K. Yang, Z. Yu, Y. Shi, and C. L. P. Chen, “A survey on imbalanced learning: latest research, applications and future directions,” *Artif. Intell. Rev.*, vol. 57, no. 6, p. 137, May 2024, doi: 10.1007/s10462-024-10759-6.
- [28] D. Elreedy, A. F. Atiya, and F. Kamalov, “A theoretical distribution analysis of synthetic minority oversampling technique (SMOTE) for imbalanced learning,” *Mach. Learn.*, vol. 113, no. 7, pp. 4903–4923, July 2024, doi: 10.1007/s10994-022-06296-4.