

ANALISIS SENTIMEN SOSIAL MEDIA X TERHADAP PROGRAM MAKAN BERGIZI GRATIS MENGGUNAKAN MACHINE LEARNING

Octri Cesar Muda ^{a,1,*}, Nahumi Nugrahaningsih ^{b,2}, Septian Geges ^{c,3}, Jadianan Parhusip ^{d,4}, Tomas Leonardor ^{e,5}

^{a,b,c} Jurusan Teknik Informatika, Fakultas Teknik, Universitas Palangka Raya, Kampus UPR Tunjung Nyaho, Jl. Yos Sudarso, Palangka Raya 73112

¹ cesarmuda@mhs.eng.upr.ac.id*; ² nahumi@it.upr.ac.id; ³ septian.geges@it.upr.ac.id, ⁴ parhusip.jadianan@it.upr.ac.id, ⁵ tomasleonardo@it.upr.ac.id.

* corresponding author

ARTICLE INFO

ABSTRACT

Keywords

Sentiment Analysis, Social Media X, Machine Learning, IndoBERT, Support Vector Machine.

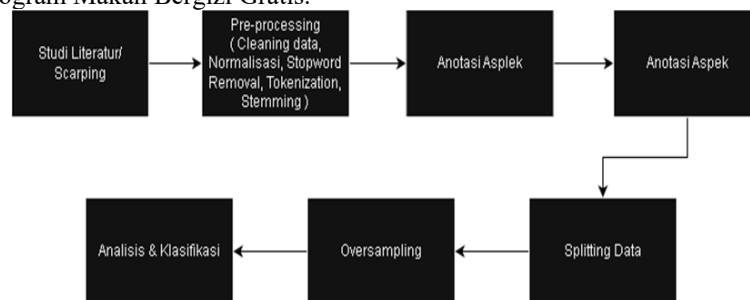
The rapid development of digital technology has made social media an important source of information and a platform for expressing public opinion on government policies, including the Free Nutritious Meal Program which aims to improve public health, quality of life, and reduce nutritional and economic problems; this study analyzes public sentiment toward the program based on 10,381 comments from users of social media platform X (Twitter) collected through web scraping using Google Colab, which were then processed through cleaning, normalization, tokenization, stopword removal, and stemming, and automatically labeled into positive, negative, and neutral categories using IndoBERT, before being classified using four machine learning algorithms, namely Naive Bayes Classifier, Support Vector Machine, K-Nearest Neighbor, and Random Forest, with the results showing that SVM achieved the highest accuracy of 74%, followed by Random Forest at 70%, Naive Bayes at 69%, and KNN at 43%, while neutral and positive sentiments were dominant, indicating that the program has received a generally favorable response from the public, and therefore the findings are expected to serve as evaluation material for the government to improve the effectiveness and implementation of the Free Nutritious Meal Program in the future.

1. Pendahuluan

Program Makan Bergizi Gratis merupakan inisiatif pemerintah atau organisasi untuk memastikan masyarakat, khususnya kelompok rentan, memperoleh akses terhadap makanan sehat guna meningkatkan kesehatan, kualitas hidup, serta mengurangi masalah gizi dan kemiskinan, namun dalam pelaksanaannya program ini memunculkan beragam reaksi masyarakat yang terekam melalui media sosial X (Twitter), baik berupa dukungan, kritik, maupun opini negatif, sehingga analisis sentimen menjadi metode yang relevan untuk mengidentifikasi dan mengelompokkan kecenderungan opini publik berdasarkan polaritasnya dengan bantuan teknologi machine learning, di mana beberapa algoritma yang umum digunakan antara lain Naive Bayes Classifier, Support Vector Machine, K- Nearest Neighbor, dan Random Forest, yang masing-masing memiliki karakteristik dan tingkat akurasi yang berbeda dalam melakukan klasifikasi teks.

2. Metodologi Penelitian

Alur penelitian ini dimulai dari studi literatur dan proses scraping untuk mengumpulkan data komentar dari media sosial X. Data tersebut kemudian melalui tahap pre-processing yang terdiri dari pembersihan data, normalisasi, penghapusan stopwords, tokenisasi, dan stemming agar teks lebih terstruktur dan siap untuk dianalisis. Setelah itu dilakukan anotasi aspek serta anotasi sentimen untuk memberikan label kategori dan polaritas pada setiap data. Data yang telah dilabeli kemudian dibagi menjadi data latih dan data uji melalui proses splitting, serta dilakukan oversampling apabila ditemukan ketidakseimbangan kelas sentimen. Tahap akhir adalah analisis dan klasifikasi menggunakan beberapa algoritma machine learning untuk menghasilkan model terbaik dalam memprediksi sentimen masyarakat terhadap program Makan Bergizi Gratis.



Gambar 1. Tahapan Penelitian

2.1. Pengumpulan Data

Pada tahap pengumpulan data, peneliti terlebih dahulu melakukan studi literatur terhadap penelitian-penelitian terdahulu yang membahas analisis sentimen pada media sosial X serta metode yang digunakan, khususnya algoritma Naive Bayes Classifier (NBC), Support Vector Machine (SVM), K- Nearest Neighbor (KNN), dan Random Forest, kemudian peneliti melakukan proses scraping data dengan mengambil komentar dari media sosial X yang berkaitan dengan program makan bergizi gratis menggunakan Google Colab dan bahasa pemrograman Python, di mana proses pengambilan data dilakukan berdasarkan beberapa kata kunci yang telah ditentukan, yaitu “program makan bergizi gratis”, “program makan bergizi gratis Prabowo”, “dampak program makan bergizi gratis”, “dampak kesehatan program makan bergizi gratis”, dan “kualitas makan bergizi gratis”, sehingga data yang diperoleh relevan dengan tujuan penelitian dan dapat digunakan dalam tahap analisis selanjutnya.

2.2. Preprocessing Data

Tahap berikutnya yaitu pre-processing bertujuan untuk membersihkan dan mempersiapkan data sebelum dilakukan analisis lebih lanjut. Pada penelitian ini setelah melakukan pre-processing masih terdapat beberapa tahapan lainnya untuk mempersiapkan dataset yaitu data cleansing, case folding, tokenization, stopword dan stemming.

2.3. Labeling

Pada tahap selanjutnya dilakukan klasifikasi atau penentuan sentimen setiap tweet melalui proses anotasi otomatis menggunakan metode IndoBERT dan TextBlob, di mana IndoBERT merupakan model bahasa berbasis Transformer yang telah dilatih pada korpus besar Bahasa Indonesia dan dipilih karena kemampuannya memahami konteks kalimat secara dua arah (bidirectional) melalui mekanisme self-attention, sehingga mampu memberikan interpretasi makna yang lebih akurat dibandingkan metode berbasis leksikon maupun algoritma machine learning tradisional, serta lebih mampu menangani karakteristik bahasa informal, singkatan, dan gaya bahasa percakapan yang umum muncul di media sosial; selain itu, pemilihan IndoBERT juga didasarkan pada performanya yang terbukti stabil dan akurat dalam berbagai penelitian sebelumnya, sehingga diharapkan dapat menghasilkan anotasi sentimen yang valid dan berkualitas untuk mendukung proses pelatihan dan pengujian model machine learning pada tahap berikutnya.

2.4. Visualisasi

Tahap berikutnya adalah visualisasi. Visualisasi merupakan langkah penting untuk mempresentasikan hasil analisis sentimen dalam bentuk yang lebih mudah dipahami dan dianalisis.

Visualisasi data memungkinkan untuk menyajikan informasi yang kompleks dalam format grafis yang jelas, sehingga pemangku kepentingan, peneliti, atau pengambil keputusan dapat lebih mudah menginterpretasi hasil analisis dan membuat keputusan berdasarkan data tersebut.

2.5. Evaluasi

Penelitian akan menggunakan rumus confusion matrix untuk evaluasi performa. Confusion matrix adalah tabel yang menyatakan klasifikasi jumlah data uji yang benar dan jumlah data uji yang salah.

3. Hasil dan Pembahasan

3.1. Scraping Data

Pada penelitian ini, tahap awal yang dilakukan adalah pengumpulan data berupa komentar pengguna media sosial X yang berkaitan dengan program Makan Bergizi Gratis guna mengidentifikasi kecenderungan sentimen publik terhadap program tersebut. Pengambilan data dilakukan secara otomatis menggunakan bahasa pemrograman Python melalui teknik web scraping, sehingga diperoleh sebanyak 10.381 komentar yang relevan. Data dikumpulkan dalam rentang waktu Januari hingga November 2025 agar mencerminkan respons masyarakat secara temporal dan representatif terhadap implementasi program tersebut.

Tabel 1. Pengumpulan Data

index	Conversation_id	Created_at	Full_text
1	1966436783370948656	Fri Sep 12 09:40:38 +0000 2025	Jangan remehkan dampak dari satu program. MBG bisa turunkan kemiskinan hingga 5 8% jika kita dukung dan awasi bersama-sama. #gasssnow #satuindonesia
2	1966404066319434083	Fri Sep 12 07:30:38 +0000 2025	Keberagaman latar belakang mitra MBG memperkaya nilai dan dampak dari program ini. #gasssnow #satuindonesia
3	1966400447566225541	Fri Sep 12 07:16:15 +0000 2025	Gabung jadi mitra MBG itu keren loh bisa beri dampak langsung ke masyarakat. #gasssnow #satuindonesia

3.2. Preprocessing

Setelah data dikumpulkan, tahap preprocessing dilakukan, Pada tahapan preprocessing bertujuan untuk membersihkan dan mempersiapkan data sebelum melakukan analisis. Dalam penelitian ini, preprocessing yang dilakukan untuk mempersiapkan dataset, yaitu meliputi data cleansing, normalisasi teks, stopwords removal, tokenizing dan stemming.

Tabel 2. Cleaning Data

index	Full_text	Clean_twitter
1	Jangan remehkan dampak dari satu program. MBG bisa turunkan kemiskinan hingga 5 8% jika kita dukung dan awasi bersama-sama. #gasssnow #satuindonesia https://t.co/0vZk71V1Wz	jangan remehkan dampak dari satu program mbg bisa turunkan kemiskinan hingga jika kita dukung dan awasi bersama sama



index	Full_text	Clean_twitter
2	Keberagaman latar belakang mitra MBG memperkaya nilai dan dampak dari program ini. #gasssnow #satuindonesia https://t.co/Rv2Ipa9vs	keberagaman latar belakang mitra mbg memperkaya nilai dan dampak dari program ini
3	Gabung jadi mitra MBG itu keren loh bisa beri dampak langsung ke masyarakat. #gasssnow #satuindonesia https://t.co/Uhvun0NfjS	gabung jadi mitra mbg itu keren loh bisa beri dampak langsung ke masyarakat

Tabel 3. Normalisasi

index	Clean_twitter	Normalisasi
1	jangan remehkan dampak dari satu program mbg bisa turunkan kemiskinan hingga jika kita dukung dan awasi bersama sama	jangan remehkan dampak dari satu program makan bergizi gratis bisa turunkan kemiskinan hingga jika kita dukung dan awasi bersama sama
2	keberagaman latar belakang mitra mbg memperkaya nilai dan dampak dari program ini	keberagaman latar belakang mitra makan bergizi gratis memperkaya nilai dan dampak dari program ini
3	gabung jadi mitra mbg itu keren loh bisa beri dampak langsung ke masyarakat	gabung jadi mitra makan bergizi gratis itu keren loh bisa beri dampak langsung ke masyarakat

Tabel 4. Stopword Removal

index	Normalisasi	Stopword
1	jangan remehkan dampak dari satu program makan bergizi gratis bisa turunkan kemiskinan hingga jika kita dukung dan awasi bersama sama	jangan remehkan dampak satu program makan bergizi gratis turunkan kemiskinan hingga dukung awasi bersama sama
2	keberagaman latar belakang mitra makan bergizi gratis memperkaya nilai dan dampak dari program ini	keberagaman latar belakang mitra makan bergizi gratis memperkaya nilai dampak program
3	gabung jadi mitra makan bergizi gratis itu keren loh bisa beri dampak langsung ke masyarakat	gabung jadi mitra makan bergizi gratis keren beri dampak langsung masyarakat

Tabel 5. Tokenisasi

index	Tokenisasi
1	[jangan, remehkan, dampak, satu, program, makan, bergizi, gratis, turunkan, kemiskinan, hingga, dukung, awasi, bersamasama]
2	[keberagaman, latar, belakang, mitra, makan, bergizi, gratis, memperkaya, nilai, dampak, program]
3	[gabung, jadi, mitra, makan, bergizi, gratis, keren, loh, beri, dampak, langsung, masyarakat]

Setelah itu, tahap berikutnya adalah stemming. Stemming adalah proses mengubah kata – kata ke bentuk dasarnya (stem) dengan menghapus imbuhan. Tujuannya adalah untuk memperlakukan kata – kata dengan akar yang sama sebagai kata yang sama, meskipun memiliki bentuk morfologis yang berbeda. Proses stemming dimulai dengan mengambil daftar token dari teks ulasan yang telah melalui proses tokenisasi.

3.4 Anotasi Aspek

Pemberian label aspek pada teks secara otomatis menggunakan pendekatan berbasis kata kunci.

Tabel 6. Anotasi Aspek

Aspek	Kata Kunci	Jumlah Data
Kualitas Program	"masalah", "rusak", "positif", "negatif", "keren", "buruk", "baik", "nikmat", "enak", "layak", "bermanfaat", "berguna", "berkualitas", "tepat guna", "aman dikonsumsi", "membantu anak sekolah"	2.288
Distribusi	"distribusi", "merata", "tidak merata", "pelaksanaan", "terlambat", "lancar", "tepat sasaran", "tidak sampai", "daerah", "sekolah", "wilayah"	647
Dampak Kesehatan	"sehat", "muntah", "psikologis", "fisik", "pintar", "stunting", "mental", "akademik", "nutrisi", "racun", "tidak sehat", "daya tahan tubuh", "sakit", "bergizi", "kuat", "bertenaga", "kurang gizi", "perut sakit".	1.444
Kebijakan Pemerintah	"evaluasi", "investasi", "sistem", "inovasi", "infrastruktur", "negara", "program", "prabowo", "efektif", "strategi", "politik", "ekonomi", "pencitraan", "kampanye", "dukung", "layan", "nasional", "pemerintah", "bidang", "program", "bijak", "langkah".	1.190
Anggaran	"anggaran", "biaya", "korupsi", "transparan", "uang", "anggar", "hemat", "mahal", "duit", "pajak", "efisien", "dana"	185
Dampak Sosial	"miskin", "orang", "manfaat", "umat", "siswa", "membantu masyarakat", "hidup", "tidak rasa", "curiga", "lapang kerja", "kesejahteraan", "rakyat", "membantu ekonomi", "beban", "setuju", "dukung", "apresiasi", "bagus", "tidak setuju", "kecewa", "kritik", "protes", "netral", "puas", "tidak puas", "nyata", "keluarga"	956
Lainnya		541
Total		7.251

3.5. Anotasi Sentimen

- IndoBERT*

Proses anotasi sentimen pada komentar ini menggunakan IndoBERT, yaitu model bahasa berbasis transformer yang telah dilatih khusus pada korpus teks Bahasa Indonesia.

Tabel 7. Anotasi Sentimen IndoBERT

Sentimen	Jumlah
Positif	1.717
Netral	2.930
Negatif	2.604
Total	7.251

- Textblob*

Pada penelitian ini juga menggunakan Textblob untuk memberikan label sentimen pada setiap komentar. TextBlob yang merupakan metode otomatis dalam analisis sentimen berbasis bahasa Inggris yang relatif sederhana.

Table 8. Anotasi Sentimen Textblob

Sentimen	Jumlah
Positif	6.611
Netral	140
Negatif	500
Total	7.251

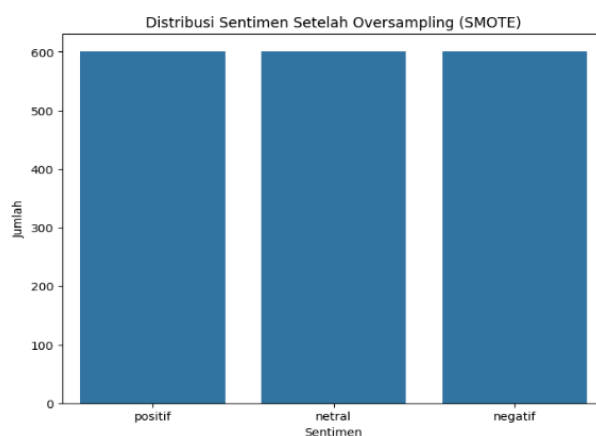
3.6. Pemodelan

- Splitting Data*

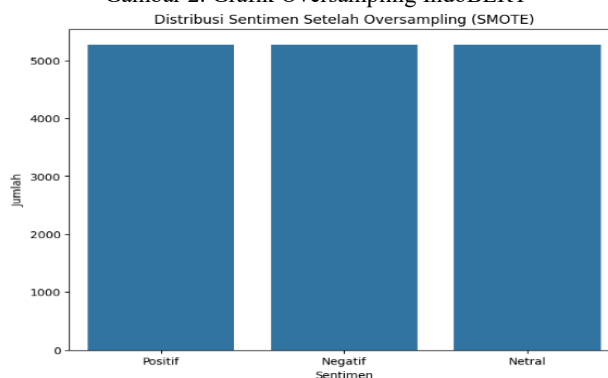
Pada tahap ini, data yang telah diberikan label akan dibagi menjadi dua bagian, yaitu data latih dan data uji. Pembagian data dilakukan dengan perbandingan 80% : 20%, di mana 80% dari total data digunakan sebagai data latih dan 20% sisanya digunakan sebagai data uji.

- Oversampling*

Dalam penelitian ini, terdapat kelas yang tidak seimbang maka peneliti menggunakan teknik oversampling diterapkan untuk mengatasi ketidakseimbangan distribusi kelas pada dataset.



Gambar 2. Grafik Oversampling IndoBERT

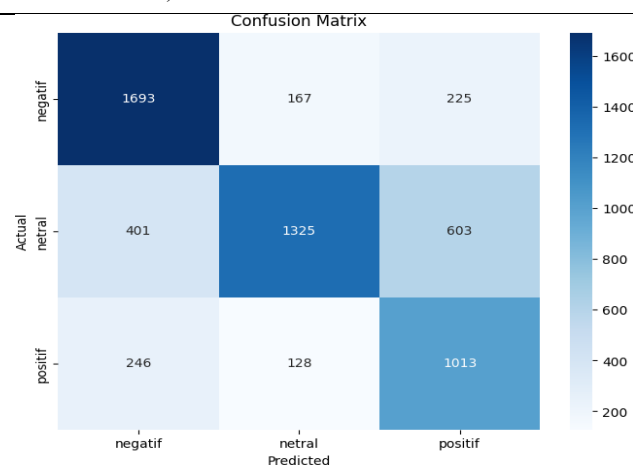


Gambar 3. Grafik Oversampling Textblob

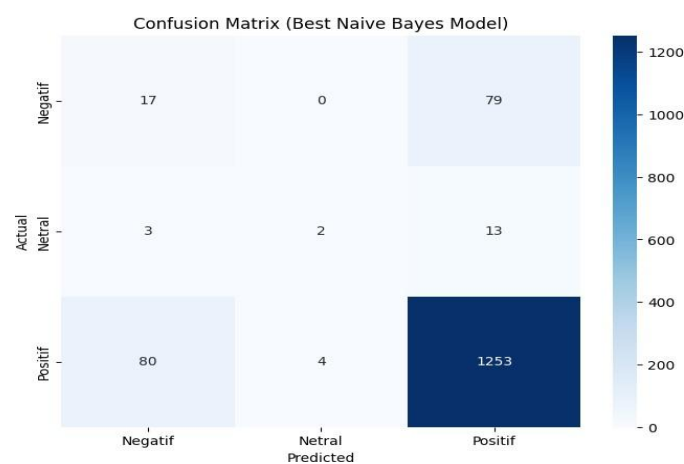
- *Hyperparameter Tuning*
Dalam penelitian ini, dilakukan hyperparameter tuning pada model Naive Bayes, Support Vector Machine (SVM), K-Nearest Neighbors (KNN), dan Random Forest untuk mengoptimalkan kinerja prediksi sentimen.
- *Naïve Bayes Classifier*
Pada penelitian ini, model Multinomial Naïve Bayes digunakan dengan parameter $\alpha = 0.001$.

Tabel 9. Hasil klasifikasi NBC menggunakan IndoBERT

	Precision	Recall	F1-Score
Negatif	0,72	0,81	0,77
Netral	0,82	0,57	0,67
Positif	0,55	0,73	0,63
Accuracy	0,69		



Gambar 4. Confusion matrix NBC IndoBERT



Gambar 5. Confusion matrix NBC TextBlob

Tabel 10. Hasil klasifikasi NBC menggunakan Textblob

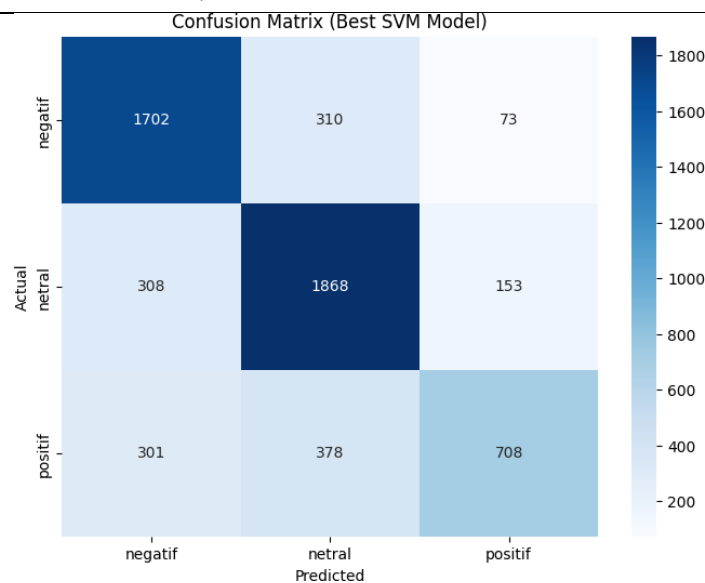
	Precision	Recall	F1-Score
Negatif	0,17	0,18	0,17
Netral	0,33	0,11	0,17
Positif	0,93	0,94	0,93
Accuracy	0,88		

- *Support Vector Machine*

Algoritma Support Vector Machine (SVM), digunakan kernel linear dengan parameter $C = 1.0$ dan $\gamma = 'scale'$. Pemilihan kernel linear Pada dilakukan karena data hasil representasi TF-IDF umumnya memiliki dimensi tinggi dan bersifat linier.

Tabel 11. Hasil klasifikasi SVM menggunakan IndoBERT

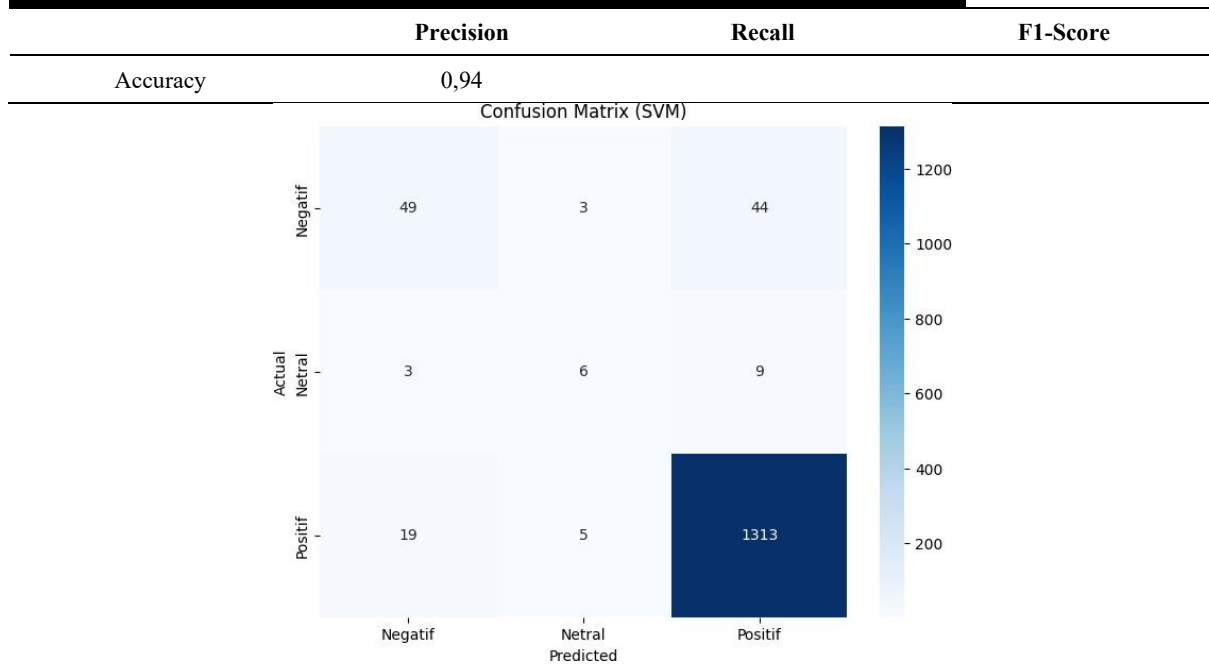
	Precision	Recall	F1-Score
Negatif	0,74	0,82	0,77
Netral	0,73	0,80	0,76
Positif	0,76	0,51	0,61
Accuracy	0,74		



Gambar 6. Confusion matrix SVM IndoBERT

Tabel 12. Hasil klasifikasi SVM menggunakan Textblob

	Precision	Recall	F1-Score
Negatif	0,69	0,51	0,59
Netral	0,43	0,33	0,38
Positif	0,96	0,98	0,97



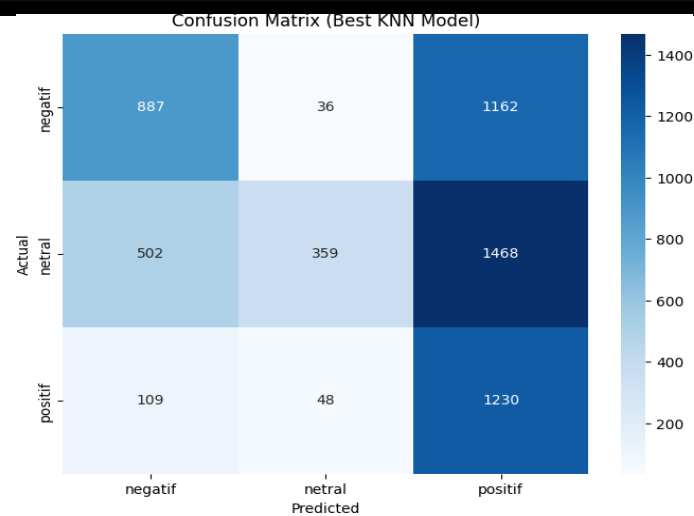
Gambar 7. Confusion matrix SVM Indobert

- *K-Nearest Neighbor*

Model K-Nearest Neighbor (KNN) dalam penelitian ini menggunakan $n_neighbors = 3$, dengan fungsi jarak Minkowski ($p = 2$) atau Euclidean Distance Nilai $k = 3$

Tabel 13. Hasil klasifikasi KNN menggunakan IndoBERT

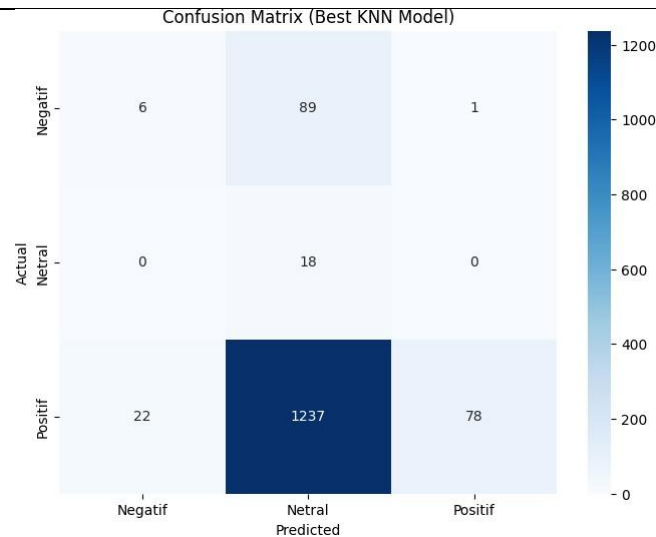
	Precision	Recall	F1-Score
Negatif	0,59	0,43	0,50
Netral	0,81	0,15	0,26
Positif	0,32	0,89	0,47
Accuracy	0,43		



Gambar 8. Confusion matrix KNN IndoBERT

Tabel 14. Hasil klasifikasi KNN menggunakan Textblob

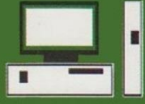
	Precision	Recall	F1-Score
Negatif	0,21	0,06	0,10
Netral	0,01	1,00	0,03
Positif	0,99	0,06	0,11
Accuracy	0,07		



Gambar 9. Confusion matrix KNN Textblob

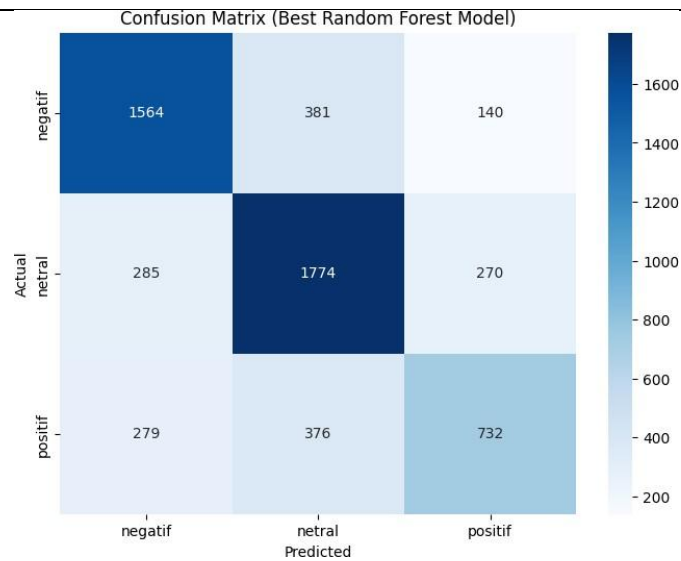
- *Random Forest*

Pada algoritma Random Forest, digunakan parameter `n_estimators = 200`, `max_depth = none`, `min_samples_leaf = 1`, `min_samples_split = 2`, dan `random_state = 42`.



Tabel 15. Hasil klasifikasi Random Forest menggunakan IndoBERT

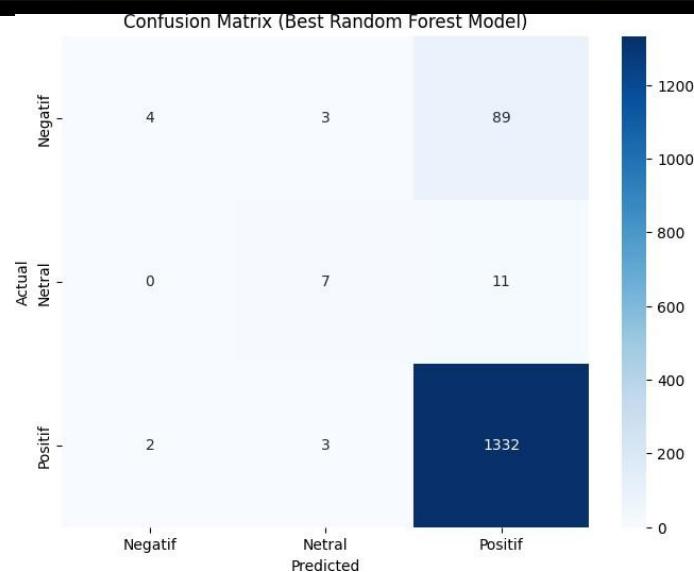
	Precision	Recall	F1-Score
Negatif	0,74	0,75	0,74
Netral	0,70	0,75	0,72
Positif	0,63	0,54	0,58
Accuracy	0,70		



Gambar 10. Confusion matrix Random Forest IndoBERT

Tabel 16. Hasil klasifikasi Random Forest menggunakan Textblob

	Precision	Recall	F1-Score
Negatif	0,67	0,04	0,08
Netral	0,54	0,39	0,45
Positif	0,93	1,00	0,96
Accuracy	0,93		



Gambar 11. Confusion matrix Random Forest Textblob

- Hasil klasifikasi seluruh model

Tabel 17. Hasil klasifikasi seluruh model menggunakan IndoBERT

	Accuracy	Precision	Recall
Naïve Bayes Classifier	0,69	0,69	0,69
Support Vector Machine	0,74	0,74	0,73
K-Nearest Neighbor	0,43	0,43	0,39
Random Forest	0,70	0,70	0,70

Pada tabel 17 dapat dilihat bahwa dari data hasil labeling menggunakan IndoBERT memberikan hasil klasifikasi menggunakan metode Naïve Bayes Classifier mendapatkan nilai akurasi sebanyak 0,69, Berikutnya klasifikasi menggunakan metode Support Vector Machine mendapatkan nilai akurasi sebanyak 0,74, Berikutnya klasifikasi menggunakan metode K-Nearest Neighbor mendapatkan nilai akurasi sebanyak 0,43, Berikutnya klasifikasi menggunakan metode Random Forest mendapatkan nilai akurasi sebanyak 0,70.

Tabel 18. Hasil klasifikasi seluruh model menggunakan Textblob

	Accuracy	Precision	Recall
Naïve Bayes Classifier	0,88	0,88	0,88
Support Vector Machine	0,94	0,94	0,94
K-Nearest Neighbor	0,07	0,07	0,11
Random Forest	0,93	0,93	0,90

Pada tabel 18 dapat dilihat bahwa dari data hasil labeling menggunakan IndoBERT memberikan hasil klasifikasi menggunakan metode Naïve Bayes Classifier mendapatkan nilai akurasi sebanyak 0,88, Berikutnya klasifikasi menggunakan metode Support Vector Machine mendapatkan nilai akurasi sebanyak 0,94, Berikutnya klasifikasi menggunakan metode K-Nearest Neighbor mendapatkan nilai akurasi sebanyak 0,07, Berikutnya klasifikasi menggunakan metode Random Forest mendapatkan nilai akurasi sebanyak 0,93.

4. Kesimpulan

Berdasarkan hasil penelitian, data berupa komentar pengguna media sosial X mengenai program Makan Bergizi Gratis berhasil dikumpulkan dan diproses melalui tahapan preprocessing serta pelabelan sentimen untuk diklasifikasikan menggunakan metode Naive Bayes Classifier, Support Vector Machine, K-Nearest Neighbor, dan Random Forest. Hasil pengujian menunjukkan bahwa Support Vector Machine memberikan performa terbaik, baik pada skema pelabelan IndoBERT maupun TextBlob, namun pelabelan menggunakan IndoBERT menghasilkan distribusi kelas yang lebih seimbang dan tidak bias karena kemampuannya memahami konteks Bahasa Indonesia secara lebih komprehensif, sedangkan TextBlob cenderung menghasilkan bias akibat keterbatasannya terhadap bahasa non-Inggris. Analisis sentimen menunjukkan dominasi sentimen positif terhadap program tersebut, yang mengindikasikan penerimaan masyarakat yang cukup baik. Secara keseluruhan, penelitian ini membuktikan bahwa pendekatan machine learning, khususnya SVM dengan pelabelan berbasis IndoBERT, efektif untuk menganalisis respons publik terhadap kebijakan sosial dan dapat dimanfaatkan sebagai bahan evaluasi dalam perumusan kebijakan yang lebih tepat sasaran.

Ucapan Terimakasih

Penulis menyampaikan apresiasi dan terima kasih kepada dosen pembimbing atas bimbingan akademik, arahan metodologis, serta dukungan yang diberikan selama proses penelitian, kepada seluruh dosen dan staf Program Studi Teknik Informatika Fakultas Teknik Universitas Palangka Raya atas kontribusi keilmuan dan fasilitas yang menunjang kegiatan akademik, kepada orang tua dan keluarga atas dukungan moral dan material yang berkelanjutan, kepada rekan-rekan mahasiswa atas diskusi dan kolaborasi yang konstruktif, serta kepada seluruh pihak lain yang telah memberikan kontribusi secara langsung maupun tidak langsung dalam penyelesaian penelitian ini.

Daftar Pustaka

- [1] Z. Firmansyah dan N. F. Puspitasari, "analisis sentimen masyarakat terhadap vaksinasi covid-19 berdasarkan opini pada twitter menggunakan algoritma naive bayes," *Jurnal Teknik Informatika*, vol. 14, no. 2, 2021, doi: 10.15408/jti.v14i2.24024.
- [2] I. Juventius, T. Gurning, P. P. Adikara, dan R. S. Perdana, "analisis sentimen dokumen twitter menggunakan metode naïve bayes dengan seleksi fitur gu metric," 2023. [Daring]. Tersedia pada: <http://j-ptiik.ub.ac.id>
- [3] A. Mondaref Jon dan I. Vitra Paputungan, "analisis sentimen pada media sosial instagram klub Persija Jakarta menggunakan metode naive bayes." [Daring]. Tersedia pada: <https://www.instagram.com/persija/>
- [4] M. Hadi Arfian dkk., "Analisis Sentimen Pada Media Sosial Menggunakan Metode Support Vector Machine," *Jurnal Ilmu Teknik dan Komputer*, vol. 09, no. 01, hlm. 1–6, 2025, doi: 10.22441/jitkom.v9i1.001.
- [5] F. Fridom Mailo dkk., "Analisis Sentimen Data Twitter Menggunakan Metode Text Mining Tentang Masalah Obesitas di Indonesia," 2019.
- [6] H. A. Rahmah, A. Anggraini, Y. P. Nilasari, dan E. P. Salsabilla, "analisis efektivitas program makan bergizi gratis di sekolah dasar Indonesia tahun 2025," *Integrative Perspectives of Social and Science Journal*, vol. 2, no. 2, hlm. 2855, 2025.
- [7] "makan bergizi gratis dan dampak bagi pertumbuhan ekonomi," Suardi Suardi, Amy Syah Purmadani, 2023.
- [8] A. Albaburrahim, A. P. A. Putikadyanto, A. N. Efendi, M. A. Alatas, S. Romadhon, dan L. R. Wachidah, "Program Makan Bergizi Gratis: Analisis Kritis Transformasi Pendidikan Indonesia Menuju Generasi Emas 2045," Entita: *Jurnal Pendidikan Ilmu Pengetahuan Sosial dan Ilmu-Ilmu Sosial*, hlm. 767–780, Mei 2025, doi: 10.19105/ejpis.v1i.19191.
- [9] S. Satriajati, S. Bagus Panuntun, dan S. Pramana, "Seminar Nasional Official Statistics 2020: Statistics in the New Normal: A Challenge of Big Data and Official Statistics IMPLEMENTASI WEB SCRAPING DALAM PENGUMPULAN BERITA KRIMINAL PADA MASA PANDEMI COVID-19 Studi Kasus: Situs Berita detik.com."

- [10] R. Gelar Guntara, "Pemanfaatan Google Colab Untuk Aplikasi Pendeteksian Masker Wajah Menggunakan Algoritma Deep Learning YOLOv7," *Jurnal Teknologi Dan Sistem Informasi Bisnis*, vol. 5, no. 1, hlm. 55–60, Feb 2023, doi: 10.47233/jteksis.v5i1.750.
- [11] S. Diantika, "penerapan teknik random oversampling untuk mengatasi imbalance class dalam klasifikasi website phishing menggunakan algoritma lightgbm," 2023.
- [12] A. Wibowo dan A. R. Isnain, "Implementasi Algoritma Machine Learning untuk Klasifikasi Suara Lingkungan," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 5, no. 2, hlm. 616–625, Mar 2025, doi: 10.57152/malcom.v5i2.1712.
- [13] M. Novitri Waruwu dkk., "implementasi algoritma machine learning untuk deteksi performa akademik mahasiswa (implementation of machine learning algorithm for student academic performance detection)."
- [14] R. S. Nurhalizah, R. Ardianto, dan P. Purwono, "Analisis Supervised dan Unsupervised Learning pada Machine Learning: Systematic Literature Review," *Jurnal Ilmu Komputer dan Informatika*, vol. 4, no. 1, hlm. 61–72, Agu 2024, doi: 10.54082/jiki.168.
- [15] F. V. Sari dan A. Wibowo, "analisis sentimen pelanggan toko online jd.id menggunakan metode naïve bayes classifier berbasis konversi ikon emosi," *Jurnal SIMETRIS*, vol. 10, no. 2, 2019.
- [16] D. Normawati dan S. A. Prayogi, "Implementasi Naïve Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," 2021.
- [17] J. Homepage dkk., "MALCOM: Indonesian Journal of Machine Learning and Computer Science Implementation of Decision Tree Algorithm and Support Vector Machine for Lung Cancer Classification Implementasi Algoritma Decision Tree dan Support Vector Machine untuk Klasifikasi Penyakit Kanker Paru," vol. 3, hlm. 15–19, 2023.
- [18] L. N. Aziza, R. Y. Astuti, B. A. Maulana, dan N. Hidayati, "Penerapan Algoritma K-Nearest Neighbor untuk Klasifikasi Ketahanan Pangan di Provinsi Jawa Tengah," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, hlm. 404–412, Feb 2024, doi: 10.57152/malcom.v4i2.1201.
- [19] R. S. Daulay, "Analisis Kritis dan Pengembangan Algoritma K-Nearest Neighbor (KNN): Sebuah Tinjauan Literatur," *Jurnal Pendidikan Sains dan Komputer*, vol. 4, no. 02, hlm. 131–141, Des 2024, doi: 10.47709/jpsk.v4i02.5055.
- [20] Suci Amaliah, M. Nusrang, dan A. Aswi, "Penerapan Metode Random Forest Untuk Klasifikasi Varian Minuman Kopi di Kedai Kopi Konijiwa Bantaeng," *VARIANSI: Journal of Statistics and Its application on Teaching and Research*, vol. 4, no. 3, hlm. 121–127, Des 2022, doi: 10.35580/variensiunm31