

ANALISIS SENTIMEN MASYARAKAT INDONESIA TERHADAP PROGRAM MAKAN BERGIZI GRATIS (MBG) DI MEDIA SOSIAL X MENGGUNAKAN PERBANDINGAN ALGORITMA NAIVE BAYES DAN SUPPORT VECTOR MACHINE

Afrissya Azza Putri, Dian Novianto

Teknik Informatika, Institut Sains dan Bisnis Atma Luhur
Jl. Jend. Sudirman, Kel. Selindung, Kec. Pangkalbalam, Kota Pangkalpinang, Indonesia
2211500075@mahasiswa.atmaluhur.ac.id

ABSTRAK

Program Makan Bergizi Gratis (MBG) yang digulirkan pemerintah Indonesia pada awal 2025 memicu perbincangan masif di platform X dengan beragam respons masyarakat. Volume unggahan yang sangat besar menjadikan pendekatan manual tidak efektif untuk memetakan opini publik, sehingga diperlukan metode komputasional berbasis *machine learning*. Penelitian ini bertujuan menganalisis sentimen masyarakat terhadap program MBG serta membandingkan kinerja algoritma Naive Bayes dan Support Vector Machine (SVM) pada klasifikasi sentimen berbahasa Indonesia. Data dikumpulkan melalui *scraping* platform X sebanyak 3.457 tweet periode Januari–Februari 2025. *Preprocessing* teks meliputi *cleaning*, *case folding*, tokenisasi, *stopword removal*, dan *stemming* menggunakan library Sastrawi. Pelabelan sentimen dilakukan dengan InSet Lexicon dan ekstraksi fitur menggunakan TF-IDF dengan pembagian data 80:20. Hasil evaluasi menunjukkan SVM dengan kernel linear mencapai akurasi 92,91%, presisi 92,63%, recall 92,91%, F1-Score 90,49%, dan AUC-ROC 0,9087, sedangkan Naive Bayes memperoleh akurasi 91,73%, presisi 84,40%, recall 91,73%, dan F1-Score 87,91%. SVM terbukti lebih unggul dari Naive Bayes pada seluruh metrik, dengan sentimen masyarakat didominasi positif sebesar 90%.

Kata kunci : analisis sentimen, Makan Bergizi Gratis, Naive Bayes, Support Vector Machine, Twitter

1. PENDAHULUAN

Program Makan Bergizi Gratis (MBG) merupakan salah satu program prioritas nasional yang diluncurkan oleh pemerintah Indonesia di bawah kepemimpinan Presiden Prabowo Subianto pada tahun 2025. Sasaran utama program ini adalah siswa sekolah dan kelompok masyarakat rentan, dengan harapan dapat menekan angka stunting sekaligus mendorong peningkatan kualitas SDM nasional [1], [2]. Sejak diluncurkan, program ini memicu perbincangan masif di berbagai platform media sosial, termasuk X (sebelumnya Twitter), dengan berbagai respons baik yang bersifat mendukung maupun mengkritisi kebijakan tersebut.

Kehadiran platform X sebagai salah satu ruang ekspresi publik utama di Indonesia menjadikannya sumber data yang relevan untuk memahami opini masyarakat terhadap kebijakan pemerintah [3]. Mengingat data yang dihasilkan sangat masif dan tidak berstruktur, pendekatan manual jelas tidak memungkinkan untuk diterapkan dalam skala besar. Oleh karena itu, diperlukan pendekatan komputasi berbasis *text mining* dan *machine learning* untuk mengekstraksi pola sentimen secara otomatis dan sistematis [4].

Untuk menjawab permasalahan tersebut, penelitian ini menerapkan pendekatan *text mining* dan *machine learning* dengan mengumpulkan data tweet berbahasa Indonesia terkait program MBG pada periode implementasi aktual program (Januari–Februari 2025) melalui *scraping* platform X. Data tersebut kemudian dilabeli secara otomatis menggunakan InSet Lexicon yang merupakan kamus sentimen *native* bahasa Indonesia, diekstraksi fiturnya

dengan TF-IDF, serta diklasifikasikan dan dibandingkan kinerjanya menggunakan dua algoritma, yaitu Naive Bayes dan Support Vector Machine (SVM). Perbandingan kinerja dievaluasi berdasarkan metrik *accuracy*, *precision*, *recall*, F1-Score, dan AUC-ROC.

Berdasarkan latar belakang tersebut, penelitian ini bertujuan untuk: (1) menganalisis distribusi sentimen masyarakat Indonesia terhadap program MBG di platform X pada periode implementasi aktual, dan (2) membandingkan kinerja algoritma Naive Bayes dan SVM dalam klasifikasi sentimen tweet berbahasa Indonesia terkait program MBG.

2. TINJAUAN PUSTAKA

Pada bagian ini diuraikan landasan teori yang menjadi dasar penelitian, meliputi konsep analisis sentimen, algoritma Naive Bayes, Support Vector Machine (SVM), metode ekstraksi fitur TF-IDF, serta InSet Lexicon dan penelitian terdahulu yang relevan.

2.1. Analisis Sentimen

Analisis sentimen merupakan cabang pemrosesan bahasa alami yang bertujuan mengenali dan mengkategorikan polaritas opini dalam sebuah teks, yakni apakah teks tersebut bernada positif, negatif, maupun netral [4]. Penerapan analisis sentimen sangat luas, mencakup evaluasi ulasan produk, pemantauan percakapan di media sosial, hingga pengukuran respons publik terhadap kebijakan pemerintah. Dalam konteks kebijakan pemerintah, analisis sentimen membantu pengambil keputusan dalam memahami persepsi publik secara cepat dan sistematis.

2.2. Naive Bayes

Naive Bayes merupakan salah satu pendekatan klasifikasi berbasis probabilitas yang berlandaskan pada Teorema Bayes, di mana setiap fitur diasumsikan bersifat independen satu sama lain. Keunggulan utama algoritma ini terletak pada efisiensi komputasinya yang tinggi serta kemampuannya menangani dataset teks berdimensi tinggi dengan baik [5]. Di antara berbagai variannya, Multinomial Naive Bayes paling sering dipilih untuk tugas klasifikasi teks karena karakteristiknya yang cocok dengan data bertipe frekuensi kata [6]. Meskipun sederhana, Naive Bayes sering kali memberikan performa yang kompetitif dalam tugas analisis sentimen.

2.3. Support Vector Machine

SVM bekerja dengan cara menemukan batas keputusan terbaik (hyperplane) yang mampu memisahkan data dari kelas berbeda dengan margin semaksimal mungkin dalam ruang fitur berdimensi tinggi [7]. SVM dengan kernel linear terbukti konsisten menghasilkan performa tinggi dalam klasifikasi teks karena kemampuannya menangani data *sparse* berdimensi tinggi [7]. Sejumlah penelitian yang membandingkan kedua metode ini secara konsisten menunjukkan SVM memberikan hasil lebih baik dalam klasifikasi sentimen teks berbahasa Indonesia [5][6].

2.4. TF-IDF

TF-IDF adalah teknik representasi teks yang memberikan bobot pada setiap kata berdasarkan seberapa sering kata tersebut muncul dalam dokumen dibandingkan dengan kemunculannya di seluruh korpus. Metode ini mampu meminimalkan dominasi kata-kata yang terlalu sering muncul dan sekaligus mengangkat bobot kata-kata yang lebih bermakna bagi proses klasifikasi [8]. Dalam penelitian ini, TF-IDF digunakan dengan maksimum 5.000 fitur dan kombinasi unigram-bigram untuk menangkap konteks yang lebih kaya.

2.5. InSet Lexicon

InSet (Indonesian Sentiment Lexicon) adalah kamus sentimen berbahasa Indonesia yang berisi daftar kata positif dan negatif beserta skornya. InSet Lexicon dipilih sebagai metode pelabelan dalam penelitian ini karena memiliki cakupan kosakata bahasa Indonesia yang luas dan mampu menangkap nuansa sentimen lokal tanpa membutuhkan proses penerjemahan [9]. Penggunaan InSet Lexicon dalam penelitian ini menjadi pembeda utama dari penelitian sebelumnya yang menggunakan VADER berbasis bahasa Inggris.

2.6. Penelitian Terdahulu

Analisis sentimen terhadap topik kebijakan publik di Indonesia telah banyak dilakukan dengan membandingkan algoritma Naive Bayes dan SVM. Yunita & Kamayani menganalisis sentimen terhadap

kebijakan penghapusan kewajiban skripsi dengan 700 tweet, dan memperoleh akurasi SVM 80% lebih tinggi dibandingkan Naive Bayes 74,3%. Pada topik kebijakan kendaraan listrik di Indonesia, Ningsih juga membuktikan bahwa SVM memberikan kinerja lebih baik dibanding Naive Bayes [10]. Supian menganalisis sentimen terhadap pemindahan Ibu Kota Nusantara (IKN) dengan hasil yang konsisten, di mana SVM mengungguli Naive Bayes pada seluruh metrik evaluasi [6].

Rabbani menguji beberapa kernel SVM untuk klasifikasi sentimen kenaikan harga BBM dengan 1.200 data dan mencatat akurasi tertinggi 85,2% pada kernel linear [7]. Munandar juga menunjukkan bahwa SVM mampu mengklasifikasikan sentimen ulasan aplikasi belajar online dengan akurasi tinggi [11]. Konsistensi hasil dari berbagai studi tersebut menegaskan bahwa Naive Bayes dan SVM merupakan dua algoritma yang paling banyak digunakan dan efektif dalam penelitian analisis sentimen teks berbahasa Indonesia.

Secara khusus terkait program Makan Bergizi Gratis, Saleh & Imanda telah melakukan analisis sentimen pada periode kampanye menggunakan 401 tweet dengan metode pelabelan VADER yang berbasis bahasa Inggris sehingga memerlukan proses translasi, dan memperoleh akurasi Naive Bayes 67,9% serta SVM 86,42% [3]. Perbedaan penelitian ini dengan penelitian tersebut terletak pada dua aspek utama: pertama, penggunaan data dari periode implementasi aktual program MBG (Januari–Februari 2025) dengan volume yang jauh lebih besar, yaitu 3.457 tweet; dan kedua, penggunaan InSet Lexicon yang merupakan kamus sentimen native bahasa Indonesia sebagai metode pelabelan, sehingga tidak memerlukan proses penerjemahan dan lebih mampu menangkap konteks bahasa lokal. Dengan pembaruan tersebut, penelitian ini diharapkan dapat menghasilkan relevansi konteks yang lebih kuat dan akurasi klasifikasi yang lebih tinggi dibandingkan penelitian terdahulu.

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif eksperimental berbasis *text mining* dan *machine learning*. Tahapan penelitian meliputi pengumpulan data, *preprocessing* teks, pelabelan sentimen, ekstraksi fitur, pemodelan, dan evaluasi model yang dijelaskan secara sistematis sebagai berikut.

3.1. Pengumpulan Data

Proses pengumpulan data dilakukan dengan mengambil tweet berbahasa Indonesia secara otomatis dari platform X, menggunakan dua kata kunci pencarian yaitu "Makan Bergizi Gratis" dan "MBG" pada rentang waktu Januari hingga Februari 2025, bertepatan dengan fase implementasi aktual program MBG oleh pemerintah. Total data yang berhasil dikumpulkan adalah 3.459 tweet berbahasa Indonesia.

3.2. Preprocessing Teks

Sebelum data teks dapat digunakan untuk melatih model, diperlukan tahap preprocessing untuk membersihkan dan menyeragamkan format data sehingga model dapat bekerja secara optimal. *Preprocessing* dilakukan melalui lima tahapan berurutan sebagaimana ditampilkan pada Tabel 1.

Tabel 1. Tahapan Preprocessing Teks

Tahapan	Sebelum	Sesudah
Cleaning	@lenteradata Semoga program makan bergizi gratis...	Semoga program makan bergizi gratis...
Case Folding	Semoga Program Makan Bergizi Gratis...	semoga program makan bergizi gratis...
Tokenisasi	semoga program makan bergizi gratis	[semoga, program, makan, bergizi, gratis]
Stopword Removal	[semoga, program, makan, bergizi, gratis]	[program, makan, bergizi, gratis]
Stemming	[program, makan, bergizi, gratis]	[program, makan, gizi, gratis]

Tabel 1 menunjukkan, setiap tahapan preprocessing memberikan transformasi yang spesifik pada data teks. Tahap *cleaning* menghapus elemen-elemen yang tidak relevan seperti URL, mention, dan hashtag. *Case folding* menyeragamkan kapitalisasi teks, tokenisasi memecah kalimat menjadi token kata, *stopword removal* menghapus kata-kata umum yang tidak bermakna, dan *stemming* mengubah kata ke bentuk dasarnya menggunakan library Sastrawi. Setelah seluruh proses preprocessing, data valid yang tersisa berjumlah 3.457 tweet dari total 3.459 tweet awal.

3.3. Pelabelan Sentimen

Proses pemberian label sentimen pada setiap tweet dilakukan secara otomatis dengan memanfaatkan InSet Lexicon. Berbeda dengan pendekatan menggunakan VADER yang berbasis bahasa Inggris dan memerlukan proses translasi, InSet Lexicon bekerja langsung pada teks bahasa Indonesia sehingga lebih akurat dalam menangkap konteks lokal. Nilai sentimen setiap tweet diperoleh dengan menjumlahkan skor masing-masing kata yang terdeteksi dalam kamus. Tweet dengan skor > 0 dilabeli positif, skor < 0 dilabeli negatif, dan skor = 0 dilabeli netral. Tweet berlabel netral (74 tweet) dikeluarkan dari dataset training, menghasilkan 3.383 tweet untuk pemodelan.

3.4. Ekstraksi Fitur TF-IDF

Fitur teks diekstraksi menggunakan TF-IDF dengan maksimum 5.000 fitur dan kombinasi unigram-bigram (ngram_range=1,2) untuk menangkap konteks kata yang lebih kaya dalam teks berbahasa Indonesia.

3.5. Pembagian Data dan Pemodelan

Dataset dibagi dengan rasio 80:20 menggunakan *stratified split* untuk mempertahankan proporsi kelas, menghasilkan 2.706 data training dan 677 data testing. Pada tahap pemodelan, dua algoritma diuji dan dibandingkan performanya, yaitu Multinomial Naive Bayes yang berperan sebagai acuan dasar, serta SVM dengan konfigurasi kernel linear. Evaluasi menggunakan lima metrik: Accuracy, Precision, Recall, F1-Score, dan AUC-ROC.

4. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil dari setiap tahapan penelitian mulai dari distribusi sentimen hasil pelabelan, hasil klasifikasi masing-masing algoritma, hingga perbandingan performa antara Naive Bayes dan Support Vector Machine berdasarkan lima metrik evaluasi yang digunakan.

4.1. Distribusi Sentimen

Hasil pelabelan menggunakan InSet Lexicon menunjukkan distribusi sentimen seperti yang disajikan pada Tabel 2.

Tabel 2. Distribusi Sentimen Hasil Pelabelan

Kelas Sentimen	Jumlah Tweet	Persentase	Keterangan
Positif	3.110	90,0%	Data training
Negatif	273	7,9%	Data training
Netral	74	2,1%	Dikeluarkan
Total	3.457	100%	-

Hasil pada Tabel 2 memperlihatkan, dari 3.457 tweet yang diproses terdapat 3.110 tweet positif (90%), 273 tweet negatif (7,9%), dan 74 tweet netral (2,1%). Tingginya proporsi sentimen positif mengindikasikan bahwa masyarakat Indonesia pada umumnya memberikan respons yang mendukung terhadap program MBG selama periode Januari–Februari 2025. Hasil ini sejalan dengan penelitian Saleh & Imanda (2025) yang juga menemukan dominasi sentimen positif terhadap program serupa. Namun, distribusi yang tidak seimbang antara kelas positif (90%) dan negatif (7,9%) ini berpotensi mempengaruhi kemampuan model dalam mendeteksi kelas minoritas.

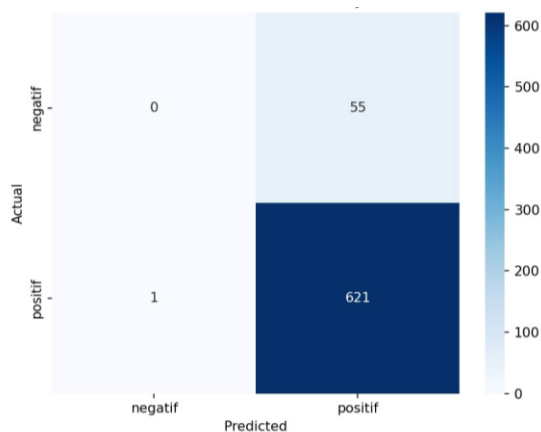
4.2. Hasil Klasifikasi Naive Bayes

Hasil evaluasi model Naive Bayes per kelas disajikan pada Tabel 3.

Tabel 3. Hasil Evaluasi Naive Bayes per Kelas

Kelas	Precision	Recall	F1-Score	Support
Negatif	0,00	0,00	0,00	55
Positif	0,92	1,00	0,96	622
Accuracy	91,73%	-	-	677
Weighted Avg	0,84	0,92	0,88	677

Merujuk pada Tabel 3, model Naive Bayes menghasilkan akurasi keseluruhan 91,73%. Untuk kelas positif, model mencapai Precision 0,92, Recall 1,00, dan F1-Score 0,96 yang menunjukkan performa sangat baik. Sebaliknya, kelas negatif sama sekali tidak berhasil dideteksi dengan nilai Precision, Recall, dan F1-Score masing-masing 0,00, yang mengindikasikan kegagalan model dalam mengenali sentimen negatif akibat dominasi kelas mayoritas (*class imbalance*) di mana model hanya belajar mengklasifikasikan kelas mayoritas. *Confusion matrix* hasil klasifikasi Naive Bayes ditampilkan pada Gambar 1.



Gambar 1. *Confusion Matrix* Naive Bayes

Gambar 1 memperlihatkan, dari 677 data testing model Naive Bayes mengklasifikasikan 621 tweet positif dengan benar (True Positive) dan hanya 1 yang salah (False Negative). Seluruh 55 tweet negatif keliru diklasifikasikan sebagai positif (False Positive = 55) dengan True Negative = 0, yang mengonfirmasi bahwa model gagal sepenuhnya dalam mengidentifikasi sentimen negatif.

4.3. Hasil Klasifikasi SVM

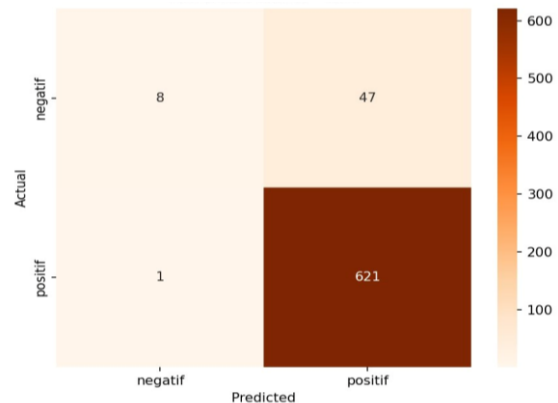
Hasil evaluasi model SVM per kelas disajikan pada Tabel 4.

Tabel 4. Hasil Evaluasi SVM per Kelas

Kelas	Precision	Recall	F1-Score	Support
Negatif	0,89	0,15	0,25	55
Positif	0,93	1,00	0,96	622
Accuracy	92,91%	-	-	677
Weighted Avg	0,93	0,93	0,90	677

Tabel 4 menunjukkan model SVM dengan kernel linear menghasilkan akurasi 92,91% dengan AUC-ROC 0,9087. Untuk kelas positif, model mencapai Precision 0,93, Recall 1,00, dan F1-Score 0,96. Dibandingkan Naive Bayes, SVM menunjukkan perbaikan signifikan pada kelas negatif dengan Precision 0,89 dan Recall 0,15, meskipun Recall masih rendah akibat ketidakseimbangan data. Nilai AUC-

ROC sebesar 0,9087 menunjukkan kemampuan diskriminasi model yang baik. Dibandingkan hasil penelitian Saleh & Imanda (2025) yang menghasilkan akurasi SVM 86,42%, penelitian ini menghasilkan akurasi lebih tinggi berkat volume data yang jauh lebih besar. *Confusion matrix* SVM ditampilkan pada Gambar 2.



Gambar 2. *Confusion Matrix* SVM

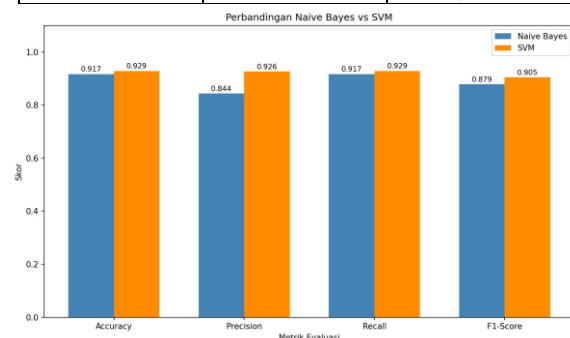
Dari Gambar 2 terlihat bahwa, *confusion matrix* SVM menunjukkan hasil lebih baik dibandingkan Naive Bayes. Model SVM berhasil mengklasifikasikan 621 tweet positif (True Positive) dan mendeteksi 8 tweet negatif dengan benar (True Negative), meskipun masih terdapat 47 tweet negatif yang salah diklasifikasikan sebagai positif (False Positive = 47). Peningkatan nilai True Negative dari 0 pada Naive Bayes menjadi 8 pada SVM membuktikan bahwa SVM lebih unggul dalam menangani ketidakseimbangan kelas.

4.4. Perbandingan Naive Bayes dan SVM

Perbandingan performa kedua algoritma secara keseluruhan disajikan pada Tabel 5 dan Gambar 3.

Tabel 5. Perbandingan Hasil Evaluasi Model

Metrik	Naive Bayes	SVM
Accuracy	91,73%	92,91%
Precision	84,40%	92,63%
Recall	91,73%	92,91%
F1-Score	87,91%	90,49%
AUC-ROC	-	0,9087



Gambar 3. Grafik Perbandingan Naive Bayes vs SVM

Tabel 5 dan Gambar 3 secara bersamaan menunjukkan, SVM mengungguli Naive Bayes pada seluruh metrik evaluasi. Kesenjangan terbesar terlihat pada nilai Precision, di mana SVM mencapai 92,63% sementara Naive Bayes hanya 84,40%, terpaut selisih 8,23%. Untuk Accuracy, SVM menghasilkan 92,91% berbanding 91,73% Naive Bayes. F1-Score SVM sebesar 90,49% juga lebih tinggi dari Naive Bayes 87,91%. Keunggulan SVM ini disebabkan kemampuannya menemukan *hyperplane* optimal pada ruang fitur berdimensi tinggi yang dihasilkan TF-IDF [7]. Kedua model menunjukkan keterbatasan dalam mendeteksi kelas negatif akibat ketidakseimbangan kelas [3].

4.5. Perbandingan dengan Penelitian Terdahulu

Untuk memperkuat validitas hasil penelitian, dilakukan perbandingan dengan penelitian terdahulu yang menggunakan metode serupa sebagaimana disajikan pada Tabel 6.

Tabel 6. Perbandingan dengan Penelitian Terdahulu

Peneliti	Algoritma	Akurasi NB	Akurasi SVM	Data
Saleh & Imanda (2025)	NB & SVM	67,90%	86,42%	401
Yunita & Kamayani (2023)	NB & SVM	74,30%	80,00%	700
Rabbani et al. (2023)	SVM Kernel	-	85,20%	1.200
Penelitian Ini	NB & SVM	91,73%	92,91%	3.457

Tabel 6 memperlihatkan bahwa penelitian ini menghasilkan akurasi Naive Bayes (91,73%) dan SVM (92,91%) tertinggi di antara seluruh studi pembandingan. Peningkatan akurasi ini didorong oleh jumlah data yang jauh lebih besar (3.457 tweet) dibandingkan penelitian terdahulu, serta penggunaan InSet Lexicon sebagai metode pelabelan yang lebih tepat untuk teks berbahasa Indonesia. Temuan ini menegaskan bahwa volume dan kualitas data pelabelan berkontribusi signifikan terhadap performa model klasifikasi sentimen.

5. KESIMPULAN DAN SARAN

Studi ini berhasil memetakan persepsi masyarakat Indonesia terhadap program Makan Bergizi Gratis (MBG) melalui analisis data dari platform X pada periode implementasi nyata program (Januari–Februari 2025). Hasil analisis menunjukkan sentimen masyarakat didominasi oleh sentimen positif sebesar 90%, mengindikasikan penerimaan yang baik terhadap program MBG pada fase awal implementasinya. Dari sisi perbandingan model, SVM dengan kernel linear terbukti lebih unggul dibandingkan Naive Bayes pada semua metrik yang diukur, dengan akurasi 92,91%, presisi 92,63%, F1-

Score 90,49%, dan AUC-ROC 0,9087, sehingga algoritma ini lebih direkomendasikan untuk klasifikasi sentimen teks berbahasa Indonesia. Penelitian ini memiliki keterbatasan berupa ketidakseimbangan kelas (*class imbalance*) di mana sentimen positif mendominasi sebesar 90% sehingga kemampuan model dalam mendeteksi kelas negatif masih terbatas. Penelitian selanjutnya disarankan untuk menerapkan teknik penyeimbangan kelas seperti SMOTE dan mengeksplorasi model *deep learning* seperti IndoBERT untuk meningkatkan kemampuan deteksi kelas minoritas.

DAFTAR PUSTAKA

- [1] P. R. Indonesia, Peraturan Presiden (Perpres) Nomor 83 Tahun 2024 tentang Badan Gizi Nasional, no. 211799. Jakarta, Indonesia, 2024. [Online]. Available: <https://peraturan.bpk.go.id/Details/295857/perpres-no-83-tahun-2024>
- [2] I. Febryanti, I. Indiaty, M. A. Pane, and P. Astuti, "Implementasi Kebijakan Makan Bergizi Gratis (MBG) (Studi Kasus Pada SDN 3 Kepanjen Kabupaten Malang)," *Dialogue : Jurnal Ilmu Administrasi Publik*, vol. 7, no. 1, pp. 067–079, Jun. 2025, doi: 10.14710/dialogue.v7i1.26628.
- [3] M. F. Saleh and R. Imanda, "Public Sentiment Analysis of the Free Meal Program: A Comparison of Naive Bayes and Support Vector Machine Methods on the Twitter (X) Social Media Platform," *Journal of Applied Informatics and Computing*, vol. 9, no. 1, pp. 131–139, Jan. 2025, doi: 10.30871/jaic.v9i1.8895.
- [4] P. Fremmuzar and A. Baita, "Uji Kernel SVM dalam Analisis Sentimen Terhadap Layanan Telkomsel di Media Sosial Twitter," *Komputika : Jurnal Sistem Komputer*, vol. 12, no. 2, pp. 57–66, Sep. 2023, doi: 10.34010/komputika.v12i2.9460.
- [5] R. Yunita and M. Kamayani, "Perbandingan Algoritma SVM Dan Naive Bayes Pada Analisis Sentimen Penghapusan Kewajiban Skripsi," *Indonesian Journal of Computer Science*, vol. 12, no. 5, Oct. 2023, doi: 10.33022/ijcs.v12i5.3415.
- [6] A. Supian, B. Tri Revaldo, N. Marhadi, L. Efrizoni, and R. Rahmaddeni, "Perbandingan Kinerja Naive Bayes Dan SVM Pada Analisis Sentimen Twitter Ibukota Nusantara," *Jurnal Ilmiah Informatika*, vol. 12, no. 01, pp. 15–21, Mar. 2024, doi: 10.33884/jif.v12i01.8721.
- [7] S. Rabbani, D. Safitri, N. Rahmadhani, A. A. F. Sani, and M. K. Anam, "Perbandingan Evaluasi Kernel SVM untuk Klasifikasi Sentimen dalam Analisis Kenaikan Harga BBM," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 3, no. 2, pp. 153–160, Oct. 2023, doi: 10.57152/malcom.v3i2.897.
- [8] B. A. Maulana, M. J. Fahmi, A. M. Imran, and N. Hidayati, "Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes

- dan Support Vector Machine (SVM),” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, pp. 375–384, Feb. 2024, doi: 10.57152/malcom.v4i2.1206.
- [9] F. Koto and G. Y. Rahmaningtyas, “Inset lexicon: Evaluation of a word list for Indonesian sentiment analysis in microblogs,” in *2017 International Conference on Asian Language Processing (IALP)*, IEEE, Dec. 2017, pp. 391–394. doi: 10.1109/IALP.2017.8300625.
- [10] W. Ningsih, B. Alfianda, R. Rahmadden, and D. Wulandari, “Comparison of Naive Bayes and SVM Algorithms in Twitter Sentiment Analysis on Electric Car Use in Indonesia,” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, pp. 556–562, Feb. 2024, doi: 10.57152/malcom.v4i2.1253.
- [11] A. A. Munandar, F. Farikhin, and C. E. Widodo, “Sentimen Analisis Aplikasi Belajar Online Menggunakan Klasifikasi SVM,” *JOINTECS (Journal of Information Technology and Computer Science)*, vol. 8, no. 2, p. 77, Jul. 2023, doi: 10.31328/jointecs.v8i2.4747