

ANALISIS SENTIMEN MASYARAKAT TERHADAP PROGRAM MAKAN BERGIZI GRATIS PADA MEDIA SOSIAL X DENGAN NAÏVE BAYES

Samuel Jason Rain¹⁾, Ratih Ihdahayuningsih²⁾, Sifa Aulia Rahmah³⁾, Hairun Adhari⁴⁾, Yanuar Fahri⁵⁾, Glenny Christo Immanuel Siwu⁶⁾

^{1,2,3,4,5,6} Mahasiswa Program Studi Sistem Informasi, Universitas Bina Sarana Informatika

Corresponding Author: ¹rthihda@gmail.com

Article Info

Article history:

Received: Des 16, 2025

Revised: Des 22, 2025

Accepted: Des 26, 2025

Published: Feb 01, 2026

Keywords:

Analisis Sentimen

Klasifikasi Teks

Naïve Bayes

Media Sosial X

Program Makan Bergizi Gratis(MBG)

ABSTRACT

Program Makan Bergizi Gratis (MBG) merupakan salah satu kebijakan strategis pemerintah yang bertujuan untuk meningkatkan kualitas sumber daya manusia melalui penguatan gizi dan kesehatan bagi siswa sekolah dasar sampai sekolah menengah atas. Meskipun memiliki tujuan positif, program ini memicu berbagai tanggapan dan opini dari masyarakat, terutama setelah munculnya beberapa kasus dugaan keracunan akibat konsumsi makanan dari program tersebut. Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap komentar atau cuitan masyarakat terkait Program MBG yang dikumpulkan dari media sosial X (sebelumnya Twitter). Dataset yang digunakan berupa data tweet yang memuat opini, komentar atau informasi terkait MBG. Metode yang digunakan adalah Naïve Bayes Classifier, metode tersebut dapat mengklasifikasikan teks berbekal asumsi probabilitas kondisional yang kuat. Data teks akan melalui tahapan *preprocessing* seperti *case folding*, *tokenizing*, *filtering*, dan *stemming*, sebelum diboboti menggunakan TF-IDF (*Term Frequency-Inverse Document Frequency*). Hasil klasifikasi sentimen akan dikelompokkan menjadi sentimen Positif dan Negatif. Hasil evaluasi menunjukkan peningkatan performa model dengan nilai akurasi mencapai 0.974. Pada kelas *Negative*, diperoleh *precision* 0.97, *recall* 1.00, dan *f1-score* 0.99, sedangkan pada kelas *Positive* diperoleh *precision* 1.00, *recall* 0.64, dan *f1-score* 0.87. Penelitian ini menunjukkan bahwa algoritma Naïve Bayes yang dikombinasikan dengan metode TF-IDF dan pendekatan self-training mampu melakukan analisis sentimen terhadap tweet mengenai MBG secara efektif dengan hasil klasifikasi akhir menunjukkan bahwa sentimen masyarakat lebih banyak terhadap sentimen negatif, sehingga penelitian ini dapat menjadi dasar evaluasi kebijakan publik berdasarkan analisis sentimen media sosial.



This is an open-access article distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International (CC BY SA 4.0)

1. INTRODUCTION

Masalah kekurangan gizi masih menjadi salah satu penyebab utama berbagai persoalan kesehatan di dunia. Kekurangan gizi dapat memicu berbagai gangguan pertumbuhan anak, salah satunya adalah *stunting*. *Stunting* merupakan kondisi dimana pertumbuhan anak terhambat sehingga anak lebih pendek dari tinggi badan anak seusianya. kondisi ini disebabkan oleh kekurangan gizi kronis dalam jangka waktu yang lama[1]. *Stunting* memiliki dampak jangka panjang yang serius, baik secara individu

maupun bagi masyarakat. Dari sisi kesehatan, *stunting* menghambat pertumbuhan fisik, perkembangan kognitif, meningkatkan resiko penyakit kronis seperti diabetes dan jantung, serta meningkatkan kerentanan terhadap kematian pada masa bayi dan balita. Dari sisi ekonomi *stunting* dapat menurunkan produktivitas tenaga kerja dan status ekonomi[2].

Menurut data Kementerian Republik Indonesia, angka *stunting* pada tahun 2021 mencapai 24%. Meskipun mengalami penurunan menjadi 21% pada tahun 2024, angka ini masih tergolong tinggi

dibandingkan dengan standar *World Health Organization (WHO)*, yang menetapkan batas maksimal di bawah 20%[3].

Sebagai bentuk langkah strategis, pemerintah saat ini telah meng-implementasikan Program Makan Bergizi Gratis (MBG) di berbagai daerah. Program ini merupakan kebijakan prioritas Presiden dan Wakil Presiden periode 2024-2029, yang ditujukan untuk meningkatkan kesehatan serta membantu perekonomian keluarga, mulai dari ibu hamil hingga pelajar tingkat sekolah menengah atas dan pondok pesantren[4].

Seiring munculnya rencana program ini hingga peng-implementasiannya sekarang, banyak menuai berbagai tanggapan publik di media sosial khususnya platform X(Twitter), baik dalam bentuk dukungan maupun kritik. Kontroversi muncul terutama terkait besarnya perkiraan biaya anggaran yang mencapai Rp.450 triliun. Selain itu, perhatian masyarakat juga tertuju pada rencana penggunaan dana dari APBN pendidikan dan dana BOS, yang memicu penolakan dari federasi Serikat Guru Indonesia (FSGI). Hal ini menimbulkan kekhawatiran masyarakat mengenai dampak terhadap biaya pendidikan, gaji guru, dan stabilitas keuangan negara[5].

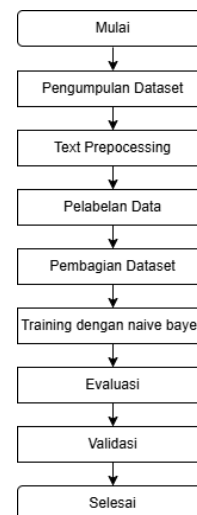
Untuk memahami bagaimana masyarakat menilai kebijakan tersebut, diperlukan pendekatan analisis sentimen. Analisis sentimen adalah metode yang digunakan untuk memahami, mengekstrak, dan memproses data teks secara otomatis guna mengidentifikasi sentimen yang terkandung dalam suatu opini [6].

Beberapa penelitian terdahulu telah menerapkan analisis sentimen berbasis media sosial untuk mengevaluasi kebijakan publik. Penelitian berjudul “*Analisis Sentimen Vaksin Sinovac Pada Twitter Menggunakan Algoritma Naïve Bayes*” yang ditulis oleh Sri Lestari dan Sudin Saepudin membuktikan bahwa penggunaan algoritma Naïve Bayes menghasilkan performa klasifikasi yang sangat baik dengan tingkat akurasi sebesar 92,96%.

Dalam klasifikasi sentimen, penelitian ini menggunakan algoritma *naïve bayes* yang didukung oleh metode *TF-IDF (Term Frequency-Inverse Document Frequency)* untuk pembobotan kata. Proses analisis data dilakukan menggunakan *google colab*. Penelitian ini menerapkan pendekatan semi-supervised learning melalui teknik self-training sebagai tahap validasi model untuk memanfaatkan data berlabel dan tidak berlabel secara optimal. Pendekatan ini diharapkan mampu meningkatkan kinerja model sehingga menghasilkan klasifikasi sentimen yang lebih akurat dan representatif.

Melalui analisis ini diharapkan dapat memberikan gambaran yang objektif mengenai kecenderungan sentiment masyarakat terhadap program Makan Bergizi Gratis (MBG), serta menjadi masukan bagi pemerintah dalam mengevaluasi dan meningkatkan strategi kebijakan gizi nasional di masa mendatang.

2. MATERIALS AND METHODS



Gambar 1. Tahapan Penelitian

2.1. Pengumpulan Dataset

Dataset yang digunakan berupa data tweet yang memuat opini, komentar atau informasi mengenai pelaksanaan Program Makan Bergizi Gratis di platform X, dengan periode pengambilan data dari November 2024 hingga Maret 2025. Data diperoleh dari platform *Kaggle* sebanyak 10.000 entri. Data tersebut mencakup beberapa atribut, yaitu *created_at* (tanggal dan jam komentar dikirim), *full_text* (isi komentar), dan *username* (nama pengguna).

2.2. Text Pre-Processing

Tahap *text preprocessing* merupakan langkah dalam proses *text mining*. Pada tahap ini dilakukan serangkaian prosedur dan pemrosesan untuk menyiapkan data teks sebelum digunakan pada proses *knowledge discovery* dalam sistem *text mining*[7]. Proses *pre-procesing* yang dilakukan meliputi lima tahap,yaitu:

1. Tokenize: proses teks akan dipecah menjadi beberapa unit kata yang lebih kecil,yang disebut “token”[8]. Tujuannya agar komputer dapat memproses teks secara terstruktur,karena komputer tidak bisa langsung memahami kalimat panjang.
2. Case Folding: langkah dalam mengubah kata dalam kalimat menjadi huruf kecil,tujuannya untuk mengurangi variasi kata akibat perbedaan kapitalisasi kata[9].
3. Filter tokens: menghapus semua karakter non alphabet,seperti simbol, spasi, angka dan lain lain yang berulang[10].
4. Stopword: menghapus kata tidak penting misalnya kata penghubung seperti “dan”,”atau”,”kemudian” dan lainnya yang tidak berpengaruh pada klasifikasi
5. Stemming: menghilangkan imbuhan pada masing kata sehingga menjadi kata dasar. Untuk tahap stemmingini, digunakan library sastrawi yang memang dirancang untuk bahasa Indonesia[11].

Selanjutnya, teks direpresentasikan dalam bentuk numerik menggunakan metode *TF-IDF* (*Term Frequency-Inverse Document Frequency*). Pembobotan TF-IDF merupakan proses mengubah data teks menjadi nilai numerik agar setiap kata atau fitur dapat dihitung tingkat kepentingannya[12].

2.3. Pelabelan Data

Di tahap ini, data akan ditandai secara manual. Label ini dipisahkan menjadi dua kategori, yaitu sentimen positif dan sentimen negatif. Sentimen positif diberikan pada tweet yang menunjukkan dukungan atau pandangan positif terhadap Program MBG, sedangkan sentimen negatif diberikan pada tweet yang memuat kritik, penolakan atau pandangan negatif.

2.4. Pembagian Dataset

Setelah melalui pelabelan data dipecah menjadi 2, data uji dan data latih. Pada analisis ini data uji 20% dan banyak data latih 80%. *Data training* merupakan sekumpulan data yang digunakan untuk melatih model agar mampu mengenali pola sesuai yang diharapkan. Sementara itu, *data testing* berfungsi untuk menguji performa model berdasarkan hasil pelatihan yang telah dilakukan.[13]

2.5. Naive Bayes

Metode ini didasarkan pada Teorema Bayes yang diperkenalkan oleh Thomas Bayes pada abad ke-18. *Naive Bayes* (NB) merupakan teknik klasifikasi berbasis statistik yang digunakan untuk memperkirakan probabilitas suatu data termasuk ke dalam kelas tertentu [14]. Naive Bayes dianalisis melalui dua pendekatan, yaitu subjektif dan *frequentist*. Pendekatan subjektif menggambarkan bagaimana tingkat keyakinan suatu hipotesis dapat diperbarui secara rasional ketika terdapat informasi atau bukti baru. Sementara pendekatan *frequentist* memandangnya sebagai peluang terbalik (*inverse probability*) yang diperoleh dari perbandingan dua kejadian berbeda.[15]

2.6. Evaluasi

Evaluasi model dilakukan untuk mengukur kinerja klasifikasi sentimen menggunakan data uji. Metrik evaluasi yang digunakan meliputi *accuracy*, *precision*, *recall*, dan *F1-score*, serta *confusion matrix* untuk melihat peforma masing-masing sentimen.

2.7. Validasi

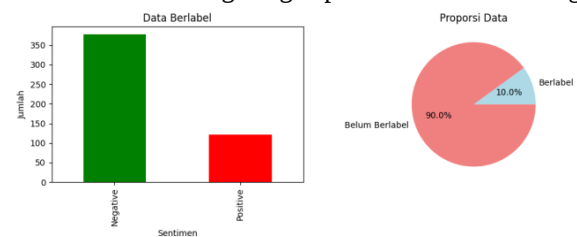
Validasi dilakukan dengan menerapkan pendekatan *self-training* pada data tidak berlabel. Model awal digunakan untuk memberikan label otomatis pada data dengan tingkat kepercayaan di atas ambang tertentu. Data hasil *auto-labeling* kemudian digabungkan dengan data berlabel awal untuk melatih ulang model dan memperoleh mode akhir yang lebih optimal.

3. RESULTS AND DISCUSSION

3.1. Pengumpulan Dataset

Dari total 10.000 dataset awal, penelitian ini menggunakan 5000 data pertama sebagai sampel analisis. Dataset tersebut kemudian dibagi menjadi dua kelompok, yaitu data berlabel dan tidak berlabel.

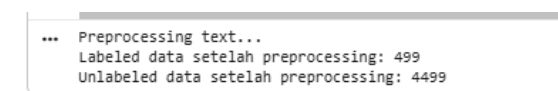
Sebanyak 500 data diberi label sentimen secara manual dan digunakan sebagai data latih awal untuk membangun model klasifikasi, sedangkan 4.500 data lainnya dimanfaatkan dalam proses validasi melalui metode auto labeling dengan pendekatan self-training.



Gambar 2. Visualisasi data labeled dan unlabeled

Berdasarkan visualisasi yang ditampilkan diagram batang menunjukkan distribusi sentimen pada data berlabel, yang terdiri dari sentimen negatif dan positif. Sementara itu, diagram lingkaran menggambarkan proporsi antara data berlabel dan tidak berlabel. Visualisasi ini digunakan untuk memberikan gambaran kondisi dataset sebelum dilakukan proses text preprocessing dan pelatihan model klasifikasi.

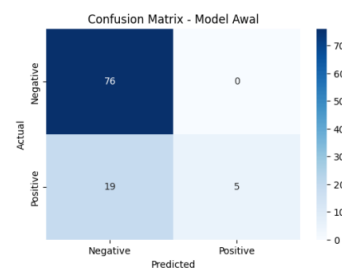
3.2. Text Preprocessing



Gambar 3. Hasil Text Preprocessing

Hasil text preprocessing menunjukkan bahwa jumlah data berlabel berkurang dari 500 menjadi 499 data, sedangkan jumlah data tidak berlabel berkurang dari 4500 menjadi 4.499 data. Pengurangan jumlah data disebabkan adanya data yang kosong atau tidak relevan terhadap konteks. Hasil ini menunjukkan proses text processing berhasil membersihkan data tanpa menghilangkan sebagian besar informasi yang digunakan dalam proses pelatihan dan validasi model.

3.3. Training Dengan Naive Bayes



Gambar 4. Confusion Matrix Awal

Confusion matrix pada gambar menunjukkan hasil klasifikasi sentimen menggunakan model naïve bayes terhadap data uji. Berdasarkan confusion matrix tersebut, dapat dijelaskan bahwa:

1. Sebanyak 76 data sentimen negatif berhasil diklasifikasikan dengan benar sebagai sentimen negatif (*true negative*).

2. Tidak terdapat data sentimen negatif yang salah diklasifikasikan sebagai sentimen positif (*false positive* = 0).
3. Sebanyak 5 data sentimen positif berhasil diklasifikasikan dengan benar sebagai sentimen positif (*true positive*).
4. Sebanyak 19 data sentimen positif salah diklasifikasikan sebagai sentimen negatif (*false negative*).

3.4. Evaluasi

Pada tahap ini dilakukan evaluasi terhadap model klasifikasi sentimen yang telah dibangun untuk mengetahui performa model berdasarkan hasil pengujian pada data uji. Berikut hasil evaluasi model awal:

```
*** === TRAINING INITIAL MODEL ===
Akurasi Model Awal: 0.8100

Classification Report Model Awal:
              precision    recall  f1-score   support

 Negative      0.80      1.00      0.89        76
 Positive      1.00      0.21      0.34        24

 accuracy      0.81      0.81      0.81       100
 macro avg     0.90      0.60      0.62       100
 weighted avg  0.85      0.81      0.76       100
```

Gambar 5.Hasil Evaluasi Training Model Awal

Berdasarkan hasil evaluasi model awal, diperoleh nilai akurasi sebesar 0.810. Untuk kelas *Negative*, model menghasilkan nilai *precision* sebesar 0.80, *recall* 1.00, dan *f1-score* 0.89, sedangkan untuk kelas *Positive* diperoleh *precision* 1.00, *recall* 0.21, dan *f1-score* 0.34. Hasil ini menunjukkan bahwa model cenderung lebih baik dalam mengklasifikasikan sentimen negatif dibandingkan sentimen positif, yang terlihat dari nilai *recall* kelas positif yang masih rendah.

```
EVALUASI MODEL FINAL
Data high confidence untuk auto-label: 4496
Dataset final: 4995 data

Distribusi label:
label
Negative  4873
Positive  122
Name: count, dtype: int64

Akurasi Model Final: 0.9740

Classification Report:
              precision    recall  f1-score   support

 Negative      0.97      1.00      0.99       972
 Positive      1.00      0.04      0.07         27

 accuracy      0.99      0.52      0.53       999
 macro avg     0.99      0.52      0.53       999
 weighted avg  0.97      0.97      0.96       999
```

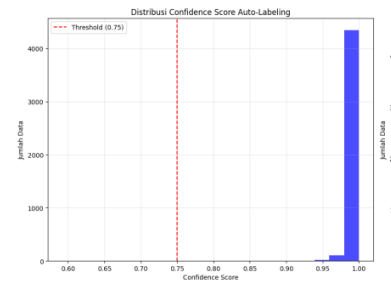
Selanjutnya, dilakukan evaluasi pada model final menggunakan data berjumlah 4.995 dengan distribusi kelas 4.873 data negatif dan 122 data positif. Hasil evaluasi menunjukkan peningkatan performa model dengan nilai akurasi mencapai 0,974. Pada kelas *Negative*, diperoleh *precision* 0.97, *recall* 1.00, dan *f1-*

score 0.99, sedangkan pada kelas *Positive* diperoleh *precision* 1.00, *recall* 0.64, dan *f1-score* 0.87.

Nilai *macro average* dan *weighted average* yang tinggi menunjukkan bahwa model memiliki performa yang baik secara keseluruhan meskipun terdapat ketidakseimbangan jumlah data antar kelas. Dengan demikian, model yang dibangun dapat dikatakan mampu melakukan klasifikasi sentimen dengan baik dan layak digunakan untuk analisis sentimen pada data penelitian ini.

3.5. Validasi

a. Distribusi Confidence Score Auto-Labeling

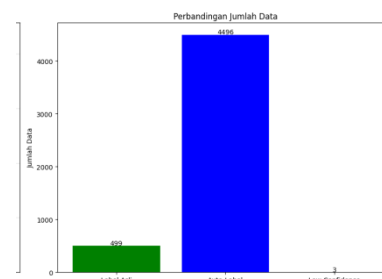


Gambar 7.Distribusi Confidence Score Auto-Labeling

Diagram di atas menampilkan distribusi nilai *confidence score* dari proses *auto-labeling* menggunakan metode *self-training*. Sumbu horizontal menunjukkan nilai *confidence score*, sedangkan sumbu vertikal menunjukkan jumlah data. Garis vertikal putus-putus berwarna merah menunjukkan ambang batas *confidence* sebesar 0,75 yang digunakan dalam penelitian.

Berdasarkan gambar tersebut, dapat dilihat bahwa sebagian besar data memiliki nilai *confidence score* yang sangat tinggi, yaitu berada pada rentang 0,95 hingga 1,00. Hal ini menunjukkan bahwa model Naive Bayes memiliki tingkat keyakinan yang tinggi dalam memberikan label otomatis pada data tidak berlabel.

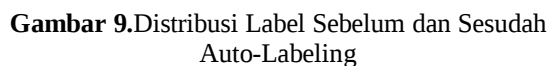
b. Perbandingan Jumlah Data Berlabel dan Auto-Label



Gambar 8.Perbandingan Jumlah Data Berlabel dan Auto-Label

Diagram di atas menunjukkan perbandingan jumlah data berdasarkan status pelabelan, yaitu data berlabel manual, data hasil auto-labeling, dan data dengan *confidence* rendah.

c. Distribusi Label Sebelum dan Sesudah Auto-Labeling



Gambar 9.Distribusi Label Sebelum dan Sesudah Auto-Labeling

Berdasarkan gambar tersebut, terlihat bahwa jumlah data pada kedua kelas sentimen mengalami peningkatan signifikan setelah proses auto-labeling. Peningkatan ini lebih dominan pada kelas sentimen negatif, yang menunjukkan bahwa opini negatif terhadap Program Makan Bergizi Gratis lebih banyak muncul pada data tidak berlabel.

a. **Sentiment Positive (evaluasi)**

[illegible]

Kata-kata seperti “bergizi gratis”, “makan siang”, “program pemerintah”, dan “anak” menunjukkan bahwa sentimen positif banyak berkaitan dengan manfaat program bagi siswa dan masyarakat. Word

b. Sentiment Negative



4. CONCLUSION

Model akhir menunjukkan peningkatan kinerja yang signifikan dengan nilai akurasi sebesar 97.40%, atau meningkat sebesar 16,40% dibandingkan model awal. Hasil klasifikasi akhir menunjukkan bahwa sentimen masyarakat lebih banyak terhadap sentimen negatif, sehingga penelitian ini dapat menjadi dasar evaluasi kebijakan publik berdasarkan analisis sentimen media sosial.

- [1] H. A. Rahmah, A. Anggraini, Y. P. Nilasari, and E. P. Salsabil, "Analisis Efektivitas Program Makan Bergizi Gratis Di Sekolah Dasar Indonesia Tahun 2025," *Integrative Perspectives of Social and Science Journal*, vol. 2, no. 2, pp. 2855–2866, 2025.
- [2] Maryuni, L. Handayani, and H. Trustisari, *Butating(Buku Pintar Cegah Stunting)*.
- [3] N. Azzahra, A. D. Dharmawan, A. F. Mardatilah, and M. I. Habibi, "Pelaksanaan Uji Coba Program Makan Bergizi Gratis di SMP Negeri 4 Tangerang," vol. 3, no. 4, pp. 5036–5044, 2025.
- [4] R. Saputra and F. N. Hasan, "Analisis Sentimen Terhadap Program Makan Siang & Susu Gratis Menggunakan Algoritma Naive Bayes," *Jurnal Teknologi Dan Sistem Informasi Bisnis*, vol. 6, no. 3, pp. 411–419, 2024, doi: 10.47233/jteksis.v6i3.1378.

- [5] E. Triningsih, "Analisis Sentimen Terhadap Program Makan Bergizi Gratis Menggunakan Algoritma MACHINE LEARNING PADA SOSIAL MEDIA X," 2025.
- [6] D. N. Sari, F. Adelia, F. Rosdiana, B. B. Butar, and M. Hariyanto, "ANALISA SENTIMEN TERHADAP REVIEW PRODUK KECANTIKAN kebutuhan konsumen dan hal ini yang harus sebaiknya konsumen mengetahui dengan detail produk yang akan dibeli , hal ini dapat dipelajari dari Beberapa review tentang produk kosmetik dapat membantu konsume," no. November, pp. 109–118, 2020.
- [7] S. Lestari and S. Saepudin, "Analisis Sentimen Vaksin Sinovac Pada Twitter Menggunakan Algoritma Naive Bayes," 2021.
- [8] A. Pramita Widyassari *et al.*, "Analisis sentimen publik di twitter terhadap pelantikan presiden Prabowo menggunakan algoritma Naïve Bayes Analysis of public sentiment on twitter towards president Prabowo's inauguration using the Naïve Bayes algorithm."
- [9] A. Sentimen *et al.*, "Analisis Sentimen Masyarakat Di Media Sosial X Terhadap Kemenkes Dengan Naive Bayes dan SVM," *Jurnal Sains dan Teknologi*, vol. 7, no. 1, pp. 1–6, 2025, doi: 10.55338/saintek.v7i1.4615.
- [10] I. F. Wijaya, D. Dionisia, and B. Rarasati, "Analisis Sentimen Pengguna Aplikasi Byond by BSI Menggunakan Algoritma Naive Bayes," *R2J*, vol. 7, no. 6, 2025, doi: 10.38035/rj.v7i6.
- [11] Mukhammad Hafiz Bima Ibrahim, Anisa Dzulkarnain, and Alqis Rausanfitra, "Penggunaan Metode Naïve Bayes untuk Klasifikasi Topik Tweet Bidang dan Non-Bidang Rektorat Telkom University pada Akun Telyufess," *Bulletin of Computer Science Research*, vol. 5, no. 5, pp. 876–885, Aug. 2025, doi: 10.47065/bulletincsr.v5i5.705.
- [12] E. Febriyani and H. Februariyanti, "Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Algoritma Naive Bayes Classifier Di Twitter," vol. 17, no. 1, pp. 25–38, 2022.
- [13] Fieryando and B. Kristianto, "Analisis Sentimen Terhadap TikTok Shop Dengan K-Nearest Neighbor , Decision Tree , dan Naive Bayes," pp. 21–29.
- [14] B. Rahmatullah, S. A. Saputra, P. Budiono, and D. P. Wigandi, "Sentimen Analisis Makan Bergizi Gratis Menggunakan Algoritma Naive Bayes," *Journal of Information Technology*, vol. 5, no. 1, 2025, doi: 10.46229/jifotech.v5i1.978.
- [15] M. C. Rani, F. D. Azkia, R. A. Dewi, M. Wahyudi, Sumanto, and A. S. Budiman, "Perbandingan Algoritma Random Forest, Naive Bayes, dan Neural Network dalam Klasifikasi Penyakit Jantung," vol. 4, no. 2, pp. 77–84, 2025.