

ANALISIS SENTIMEN TERHADAP MAKAN BERGIZI GRATIS PADA SOSIAL MEDIA X MENGGUNAKAN METODE DECISION TREE DENGAN PEMBANDING NAIVE BAYES

Kurniawan Yoga Pratama, Sentot Achmadi, Karina Aulia Sari

Teknik Informatika, Institut Teknologi Nasional Malang

Jalan Raya Karanglo km 2 Malang, Indonesia

2118043@scholar.itn.ac.id

ABSTRAK

Media sosial X (sebelumnya Twitter) menjadi wadah populer bagi masyarakat untuk menyuarakan opini, termasuk terkait program makan bergizi gratis yang bertujuan meningkatkan kualitas gizi anak sekolah dan keluarga kurang mampu. Program ini memunculkan beragam reaksi, baik positif, negatif, maupun netral, sehingga penting dilakukan analisis sentimen untuk memahami persepsi publik. Permasalahan yang diangkat dalam penelitian ini adalah bagaimana mengidentifikasi dan mengklasifikasikan sentimen masyarakat di media sosial X terhadap program makan bergizi gratis dengan metode yang tepat. Tujuan penelitian ini adalah menganalisis sentimen masyarakat, mengelompokkannya ke dalam tiga kategori (positif, negatif, netral), serta membandingkan performa algoritma Decision Tree sebagai metode utama dengan Naïve Bayes sebagai pembanding. Metode penelitian dilakukan melalui tahapan pengumpulan data sebanyak 643 tweet dengan kata kunci “Makan Bergizi Gratis”, preprocessing (case folding, cleansing, tokenisasi, stopword removal, normalisasi, stemming), pembobotan TF-IDF, serta klasifikasi menggunakan Decision Tree dan Naïve Bayes. Evaluasi kinerja model dilakukan menggunakan confusion matrix dengan metrik akurasi, presisi, recall, dan f1-score. Hasil penelitian menunjukkan bahwa Decision Tree mencapai akurasi 76,38% dengan kinerja terbaik pada sentimen positif, sementara Naïve Bayes mencapai 75,59% namun cenderung bias terhadap kelas mayoritas. Secara keseluruhan, Decision Tree dinilai lebih proporsional untuk klasifikasi multi-kelas pada analisis sentimen program makan bergizi gratis.

Kata kunci : Analisis Sentimen, Makan Bergizi Gratis, Decision Tree, Media Sosial X, Naïve Bayes

1. PENDAHULUAN

Media sosial telah berkembang menjadi salah satu saluran penting bagi masyarakat untuk menyampaikan pendapat mereka, memberikan tanggapan, dan berbicara tentang berbagai masalah kontemporer. X (sebelumnya Twitter) adalah platform media sosial yang sangat populer saat ini yang memungkinkan pengguna berbagi informasi dengan cepat dan mudah melalui tweet, retweet, dan komentar. Dalam situasi seperti ini, media sosial X menjadi tempat yang tepat untuk menggali pendapat publik tentang berbagai program sosial yang dilaksanakan oleh organisasi dan pemerintah.[5]

Permasalahan utama dalam penelitian ini adalah bagaimana mengidentifikasi dan mengklasifikasikan sentimen masyarakat terhadap program makan bergizi gratis pada media sosial X. Sentimen masyarakat yang tereksprei dalam bentuk teks cenderung bervariasi, sehingga memerlukan metode analisis yang tepat. Penelitian ini bertujuan menganalisis sentimen tersebut, membaginya ke dalam kategori positif, negatif, dan netral, serta membandingkan kinerja metode Decision Tree dengan Naïve Bayes untuk mengetahui model yang lebih optimal.[5]

Tujuan program makan bergizi gratis, yang sering diperdebatkan, adalah untuk meningkatkan asupan gizi masyarakat, terutama bagi anak-anak dan keluarga yang kurang mampu dengan memastikan bahwa semua orang dapat mendapatkan makanan bergizi tanpa harus membayar. Program makan bergizi gratis,

seperti program sosial lainnya, seringkali menghasilkan berbagai reaksi, mulai dari kritik dan ketidakpuasan.[6]

Untuk menyelesaikan masalah tersebut, penelitian dilakukan melalui beberapa tahap, yaitu pengumpulan data dari media sosial X, preprocessing teks (case folding, tokenisasi, stopword removal, normalisasi, dan stemming), serta ekstraksi fitur menggunakan TF-IDF. Data yang diperoleh kemudian diklasifikasikan dengan algoritma Decision Tree dan dibandingkan hasilnya dengan Naïve Bayes. Evaluasi performa kedua model diharapkan dapat menunjukkan metode yang paling efektif dalam analisis sentimen terkait program makan bergizi gratis.

2. TINJAUAN PUSTAKA

2.1. Makan Bergizi Gratis

Tujuan program makan siang gratis yang diluncurkan pemerintah adalah untuk meningkatkan nutrisi anak sekolah, ibu hamil, dan balita. Tapi program ini menghadapi masalah besar, seperti masalah distribusi dan risiko pemborosan. Dengan menganalisis sentimen platform X, penelitian ini mencoba memahami reaksi masyarakat terhadap program tersebut.[2]

2.2. Analisis Sentimen

Analisis sentimen menggunakan text analytics untuk mendapatkan data dari berbagai sumber dari internet dan berbagai platform media sosial. Tujuan

analisis sentimen adalah untuk mengetahui perasaan, sikap, atau pendapat pengguna tentang topik atau entitas tertentu. Ini melakukannya dengan mengklasifikasikan teks menjadi sentimen positif, negatif, atau netral. Teknik dan algoritma komputasional seperti machine learning digunakan dalam analisis sentimen. [12]

2.3. Text Pre-processing

Ini merupakan langkah awal dalam menyiapkan data agar sesuai dengan format yang diperlukan untuk analisis. Tahapan ini bertujuan untuk mengekstrak, mengolah, dan menyusun informasi, serta mengevaluasi hubungan antar teks baik pada data terstruktur maupun tidak terstruktur. Proses ini juga membantu mereduksi dimensi data agar menjadi lebih ringkas dan terorganisir, sehingga lebih mudah untuk dianalisis pada tahap selanjutnya. Beberapa tahapan dalam pra-pemrosesan meliputi pembersihan data (cleansing), normalisasi bahasa, penghapusan *stopwords*, dan tokenisasi. [12]

2.4. Metode Decision Tree

merupakan sebuah model prediktif yang menggunakan struktur berbentuk pohon atau hierarki. Prinsip utama dari metode ini adalah mengubah data menjadi bentuk pohon keputusan yang berisi serangkaian aturan untuk pengambilan keputusan [3] Struktur Decision Tree terdiri dari beberapa komponen utama:

- Root Node (Simpul Akar): Merupakan simpul paling atas yang mewakili keseluruhan dataset. Dari simpul ini, pohon akan mulai terbagi.
- Internal Node (Simpul Internal): Merupakan simpul percabangan yang merepresentasikan sebuah tes pada suatu atribut atau fitur data.
- Branch (Cabang): Mewakili hasil dari pengujian pada simpul internal, yang menghubungkan ke simpul selanjutnya.
- Leaf Node (Simpul Daun): Merupakan simpul akhir yang tidak memiliki cabang lagi. Simpul ini merepresentasikan keputusan atau hasil klasifikasi akhir (misalnya: "positif", "negatif", atau "netral").

Alur Metode Decision Tree Cara kerja Decision Tree adalah membangun pohon keputusan dengan cara memecah data secara berulang berdasarkan "pertanyaan" (fitur) terbaik. Proses ini bertujuan untuk membuat kelompok data yang sejenis (murni).

Langkah 1: Hitung Entropi Total

Pertama, hitung tingkat ketidakpastian dari keseluruhan dataset (S). Ini disebut Entropi(S).

$$Entropy(S) = -\sum_{i=1}^c p_i \log_2(p_i) \quad (1)$$

Langkah 2: Hitung Information Gain untuk Setiap Atribut Selanjutnya, untuk setiap atribut (fitur), hitung seberapa besar "keuntungan informasi" yang didapat

jika kita memecah data menggunakan atribut tersebut. Ini disebut Information Gain.

$$Gain(S, A) = Entropy(S) - \sum_{values(A)} \frac{|S_v|}{|S|} Entropy(S_v) \quad (2)$$

Langkah 3: Pilih Atribut Terbaik sebagai Node

Bandingkan nilai Information Gain dari semua atribut.

Atribut dengan Information Gain tertinggi dipilih sebagai simpul pemecah (decision node). Ini adalah pertanyaan terbaik untuk memisahkan data pada tahap ini.

Langkah 4: Ulangi Proses untuk Setiap Cabang (Rekursif)

Setelah data dipecah, ulangi Langkah 1 hingga 3 untuk setiap cabang (subset data) yang baru. Proses ini terus berjalan hingga semua data dalam satu cabang sudah sejenis (Entropi = 0) atau tidak ada lagi atribut yang bisa memisahkan data. Cabang yang berhenti akan menjadi simpul daun (leaf node) yang berisi keputusan akhir.

2.5. Evaluasi

Evaluasi dilakukan untuk mengukur kualitas dari suatu model klasifikasi. Kinerja sistem klasifikasi perlu dianalisis guna mengetahui seberapa akurat prediksi yang dihasilkan. Untuk menilai tingkat akurasi model yang digunakan dalam analisis sentimen, dapat diterapkan metode evaluasi tertentu.[8] Beberapa metrik yang umum digunakan dalam proses ini antara lain adalah *accuracy*, *precision*, dan *recall*. Nilai-nilai tersebut diperoleh dengan menggunakan confusion matrix seperti pada tabel1

Tabel 1. Confussion Matrix

Predicted Class	True Class	
	Positif	Negatif
Positif	True Positif (TP)	False Positif (FP)
Negatif	False Positif (FP)	True Negatif (TN)

Keempat parameter diatas digunakan untuk menghitung 3 metode evaluasi yakni :

- Seberapa baik suatu sistem mampu menemukan kembali data yang relevan ditentukan oleh recall. Dalam klasifikasi, recall adalah ukuran proporsi data positif yang berhasil dikenali oleh model dibandingkan dengan total data positif.[10] Nilai recall diperoleh dengan membagi jumlah data TP dengan total data positif, atau hasil penjumlahan antara data TP dan NP .Perhitungan recall dapat dirumuskan sebagai berikut::

$$Recall = \frac{TP}{TP + FP} \times 100\% \quad (3)$$

- Kecocokan antara bagian data yang diambil dengan informasi yang dibutuhkan dikenal sebagai *precision*. *Precision* berfokus pada bagian data yang telah diprediksi sebagai positif, menilai seberapa banyak dari prediksi tersebut yang benar-benar merupakan nilai positif. *Precision* digunakan untuk memahami seberapa baik model dalam menghasilkan prediksi positif yang benar.[10]

Precision sering digunakan bersamaan dengan *recall* untuk mendapatkan gambaran yang lebih lengkap tentang kinerja model. *Precision* dapat dihitung melalui rumusan berikut ini :

$$Precision = \frac{TP}{TP + FN} \times 100\% \quad (4)$$

- c. *Accuracy* merupakan ukuran yang menunjukkan seberapa dekat hasil prediksi model dengan nilai sebenarnya dalam data. Metrik ini memberikan gambaran menyeluruh mengenai performa model dengan menunjukkan proporsi prediksi yang benar dari seluruh data yang diuji. Walaupun *accuracy* sering dijadikan tolak ukur utama dalam menilai kinerja model, penggunaannya tetap perlu disandingkan dengan metrik lain seperti *precision* dan *recall* untuk memperoleh evaluasi yang lebih menyeluruh dan akurat.[10] Nilai *accuracy* dapat dihitung menggunakan rumus berikut:

$$Accuracy = \frac{TN + TP}{TN + TP + FN + FP} \times 100\% \quad (5)$$

2.6. Pembobotan TF-IDF

Metode yang dikenal sebagai TF-IDF (Term Frequency-Inverse Document Frequency) digunakan untuk memberikan nilai pada kata-kata yang dihasilkan dari ekstraksi dengan menilai frekuensi kata-kata tertentu dalam suatu dokumen dibandingkan dengan frekuensi total kumpulan dokumen. Teknik ini umum digunakan dalam bidang *information retrieval* untuk mengidentifikasi kata-kata yang dianggap penting dalam konteks tertentu.. Tahapan ini digunakan untuk mengukur nilai bobot suatu kata dengan menimbang frekuensi kemunculan kata tersebut. Hasil pembobotan ini dapat digunakan untuk klasifikasi, *clustering* dan pengambilan informasi dokumen[3]. Persamaan untuk menghitung nilai TF-IDF dapat dilihat pada persamaan (6)

$$W = tf \times idf \quad (6)$$

Keterangan :

W = Bobot TF-IDF

tf = Hasil perhitungan nilai *term frequency*

idf = Hasil perhitungan nilai *Inverse Document Frequency*

2.7. Term Frequency (TF)

Nilai *Term Frekuensi* (TF) menunjukkan seberapa sering kata muncul dalam dokumen. Nilai TF yang lebih tinggi, atau jumlah kata yang muncul, akan menerima bobot yang lebih besar.[3] Rumus perhitungan *Term Frequency* dapat dilihat pada persamaan (7).

$$TF = 1 + \log(F_{t,d}), F_{t,d} > 0 \quad (7)$$

Keterangan :

TF = *Term Frequency*

$F_{t,d}$ = *Frequency Term* pada dokumen

2.8. Inverse Document Frequency (IDF)

Inverse Document Frequency digunakan untuk menentukan apakah kata yang dicari sama dengan kata kunci yang diinginkan [3]. Pada proses perhitungan IDF, semakin cocok term yang dicari dengan kata kunci, maka nilai IDF akan semakin rendah. Hal ini dikarenakan proses ini berfokus pada kata yang unik dan relevan dengan dokumen yang dianalisis dengan mengurangi pengaruh kata yang umum. Rumus untuk menghitung *Term Frequency* dapat dilihat pada persamaan (8).

$$IDF = \log\left(\frac{D}{df}\right) \quad (8)$$

Keterangan :

IDF = *Inverse Document Frequency*.

D = Jumlah keseluruhan dokumen.

Df = Jumlah dokumen yang mengandung *term*.

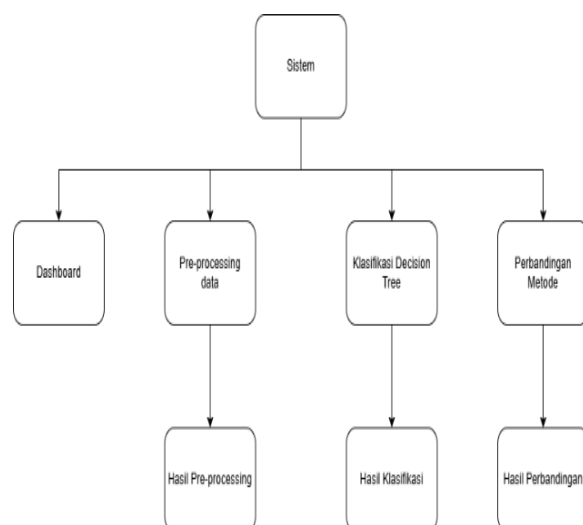
3. METODE PENELITIAN

3.1. Pengumpulan Data

Data dikumpulkan dari media sosial X (Twitter) menggunakan Twitter API, dengan kata kunci pencarian “Makan Bergizi Gratis”. Periode pengambilan data dilakukan antara bulan September 2024 hingga Februari 2025. Dataset terdiri dari 600 tweet berbahasa Indonesia, yang mencerminkan opini masyarakat terhadap program makan bergizi gratis yang dicanangkan pemerintah. Setiap tweet dianotasi secara manual menjadi tiga kategori sentimen: positif, negatif, dan netral.

3.2. Struktur Menu

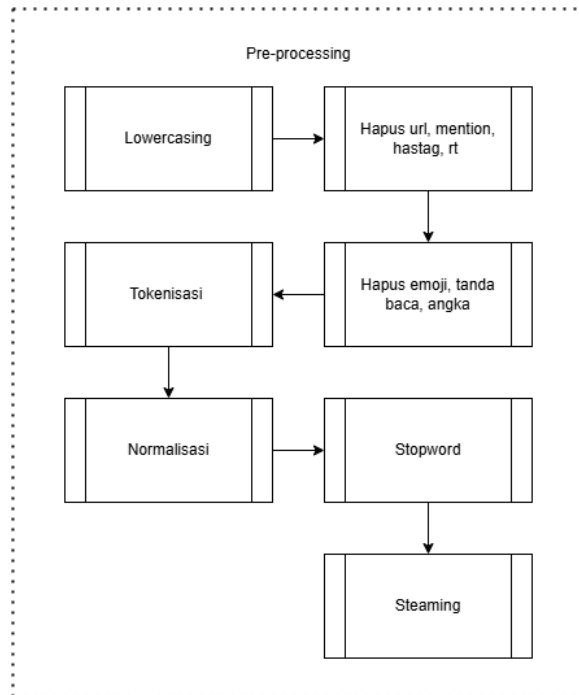
Perancangan struktur menu bertujuan untuk memberikan panduan yang jelas bagi pengguna aplikasi, sehingga mereka dapat dengan mudah mengakses fitur yang dibutuhkan terdapat Halaman Dashboard, Halaman Pre-processing, Halaman Klasifikasi Decision Tree, Halaman Perbandingan Klasifikasi



Gambar 1. Struktur Menu

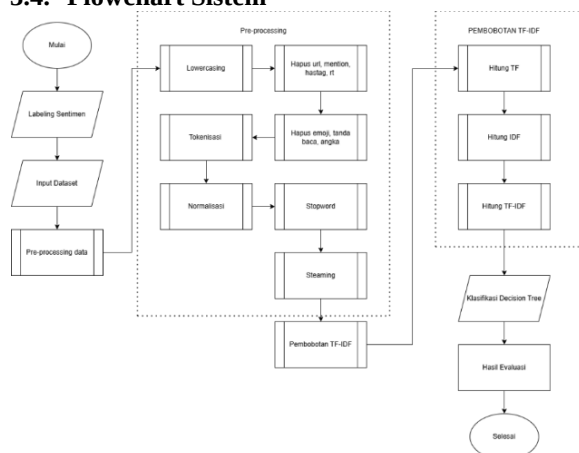
3.3. Pre-processing

Sebelum memasuki tahap klasifikasi, proses pertama yang dilakukan adalah preprocessing teks, yang terdiri dari beberapa langkah yang sangat penting. Case folding (mengubah seluruh huruf menjadi huruf kecil), cleansing (menghapus fitur seperti URL, emoji, angka, dan tanda baca), tokenisasi (memisahkan teks menjadi kata-kata individu), stopwords removal (menghilangkan kata-kata yang tidak penting), normalisasi (mengganti kata tidak baku menjadi bentuk baku), dan stemming (mengembalikan kata ke bentuk dasarnya).



Gambar 2. Pre-processing

3.4. Flowchart Sistem

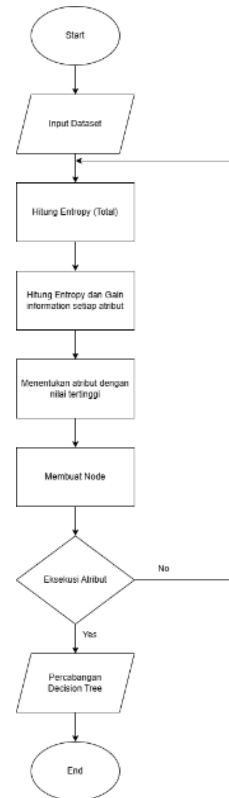


Gambar 3. Flowchart Sistem

3.5. Flowchart Metode

Alur algoritma Decision Tree dimulai dengan memasukkan dataset, kemudian sistem menghitung nilai entropi total untuk mengukur tingkat ketidakpastian data. Selanjutnya, dihitung entropi dan

gain informasi untuk setiap atribut guna menentukan atribut terbaik dengan nilai gain tertinggi. Atribut tersebut dipilih sebagai node dan dibuatkan percabangan berdasarkan nilai-nilainya. Proses ini dieksekusi secara berulang hingga semua atribut selesai digunakan, dan terbentuklah struktur pohon keputusan (decision tree) yang menjadi hasil akhir klasifikasi



Gambar 4. Flowchart Metode

Sistem memulai dengan pelabelan sentimen dan data dari media sosial X. Setelah itu, data diproses melalui proses pre-processing seperti lowercasing, penghapusan elemen tidak penting (seperti URL, emoji, dan tanda baca), tokenisasi, normalisasi, penghapusan stopwords, dan stemming. Selanjutnya, pembobotan dilakukan menggunakan TF-IDF. Data yang telah dibobot diklasifikasikan dengan metode Decision Tree, dan hasilnya dinilai dengan metrik akurasi, presisi, recall.

4. HASIL DAN PEMBAHASAN

Dataset terdiri dari tweet hasil crawling, yang kemudian dilabeli secara manual ke dalam tiga kategori sentimen: positif, negatif, dan netral. Setelah proses *pre-processing* dan pembobotan TF-IDF, data dibagi ke dalam rasio 80:20 untuk pelatihan dan pengujian model

4.1. Implementasi Pre-processing

Tentang data terdapat 643 data yang telah di crawling pada sosial media X terdapat data positif 486, data negative 78, data netral 69.

Table 2. Tentang Data

Full text	Sentimen
sejujurnya klo anggaran dipangkas untuk efisiensi yg dimana dananya untuk keberlanjutan makan bergizi gratis mnurutku mnding dihapus aja sih dripda pendidikan dan kesahtan juga dikorbankan dimana seharusnya ini jdi pondasi utama.	POSITIF
Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. https://t.co/5rkkip15n6	NEGATIF

Table 3. Menghapus Mention

Sebelum	Sesudah
*Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. https://t.co/5rkkip15n6 https://t.co/GcvpK4itv8 @Dusi #Majubersama Prabowo	*Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. * Program Makan Bergizi https://t.co/GcvpK4itv8 #Majubersama Prabowo Gratis di 8

Table 4. Menghapus link

Sebelum	Sesudah
*Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. * https://t.co/GcvpK4itv8 @Dusi #Majubersama Prabowo	*Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. * #Majubersama Prabowo

Table 5. Tanda Baca, Angka, Emoji

Sebelum	Sesudah
Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. https://t.co/GcvpK4itv8 @Dusi #Majubersama Prabowo	Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. #MajubersamaPrabowo

Table 6. Casefolding

Sebelum	Sesudah
Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl.	ratusan pelajar di papua menolak makan gratis. tdk makan bergizi karena dibungkus pake plastik dl.

Table 7. Stopword dan steaming

Sebelum	Sesudah
ratusan pelajar di papua menolak makan gratis. tdk makan bergizi karena dibungkus pake plastik dl. https://t.co/GcvpK4itv8 @Dusi #Majubersama Prabowo	Ratus pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl.

Table 8 Hapus Hastag

Sebelum	Sesudah
Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl.	Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl.

Table 9 Tokenisasi

Sebelum	Sesudah
Ratusan pelajar di Papua menolak makan gratis. Tdk makan bergizi karena dibungkus pake plastik dl. https://t.co/GcvpK4itv8 @Dusi #Majubersama Prabowo	'Ratusan', 'Pelajar', 'menolak', 'makan', 'gratis', 'tidak', 'makan', 'bergizi', 'karena', 'dibungkus', 'pake', 'plastik', 'dl',

4.2. Pengujian Metode

Dalam penelitian ini, dataset dibagi ke dalam dua bagian: data pelatihan dan data pengujian. Data pelatihan memiliki proporsi 80% dan data pengujian 20%. 127 data uji dievaluasi. Algoritma Decision Tree digunakan untuk menguji model klasifikasi. Ini menilai performa model berdasarkan confusion matrix dan metrik evaluasi seperti ketepatan, recall, dan skor f1. Data dikategorikan ke dalam tiga kategori sentimen, menurut model: positif, negatif, dan netral.

Tabel 10. Confusion Matrix

True Class	Predicted Class			
	Negatif	Netral	Positif	Total
Negatif	4	2	9	15
Netral	5	4	8	17
Positif	7	3	85	95
Total	16	9	102	127

Tabel 11. Evaluasi Performa

Sentimen	Precision	Recall	F1-Score	Support
Negatif	0.25	0.267	0.258	15
Netral	0.444	0.235	0.307	17
Positif	0.833	0.894	0.862	95

Tabel 10 dan 11 menampilkan matrix confusion yang menunjukkan hasil klasifikasi pada data uji. Dari 17 data aktual negatif, hanya 4 yang berhasil dianggap benar, sementara sisanya salah dianggap netral dan positif. Dari 17 data aktual netral, hanya 4 yang dianggap tepat. Namun, model berhasil mengklasifikasikan 85 dari 95 data positif dengan benar; secara keseluruhan, model mampu mengklasifikasikan 93 dari 127 data dengan tepat, menghasilkan akurasi sebesar 76,38%

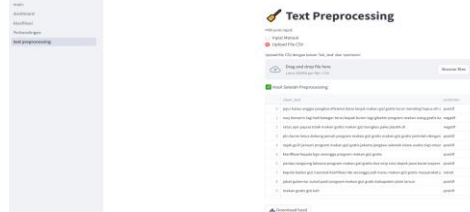
4.3. Implementasi Dashboard



Gambar 5. Halaman Dashboard

Dashboard ini adalah aplikasi analisis sentimen yang dirancang untuk topik "Makan Bergizi Gratis" (MBG), yang memungkinkan pengguna untuk mengolah data tweet terkait topik tersebut. Menu utama mencakup tahap Text Preprocessing, di mana data tweet dibersihkan dan diproses agar siap dianalisis, dan Training Model, yang digunakan untuk melatih model

4.4. Implementasi Halaman Text Pre-processing



Gambar 6. Halaman Text Pre-processing

Pada menu Text Preprocessing ini, pengguna dapat memilih untuk mengimpor data melalui dua opsi: Input Manual atau Upload File CSV. Pada bagian ini, file CSV yang diunggah harus berisi dua kolom utama: 'full_text' yang berisi teks (seperti tweet atau komentar), dan 'sentimen' yang menunjukkan label sentimen (positif, negatif, atau netral) untuk setiap teks

4.5. Implementasi Visualisasi Data Cloudword



Gambar 7. Visualisasi Data Cloudword

Menu Visualisasi Data ini menampilkan dua wordcloud yang menggambarkan distribusi kata-kata berdasarkan sentimen dalam dataset yang telah diproses di sebelah kiri positif sebelah kanan negatif

4.6. Implementasi Klasifikasi dan Evaluasi Decision Tree

Klasifikasi Sentimen - Decision Tree

Hasil Evaluasi

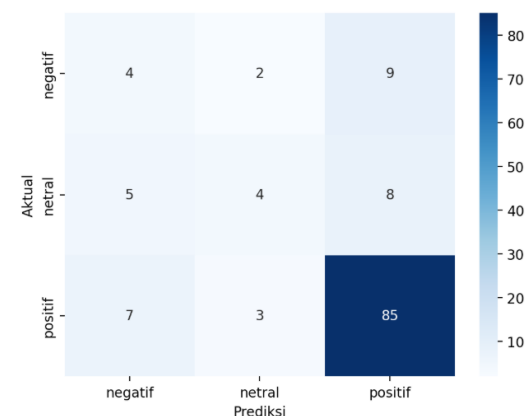
Akurasi: 0.73

	precision	recall	f1-score	support
negatif	0.25	0.2667	0.2581	15
netral	0.4444	0.2353	0.3077	17
positif	0.8333	0.8947	0.8629	95
accuracy	0.7323	0.7323	0.7323	0.7323
macro avg	0.5093	0.4656	0.4762	127
weighted avg	0.7124	0.7323	0.7172	127

Gambar 8. Klasifikasi Decision Tree

Halaman ini menampilkan hasil evaluasi model klasifikasi sentimen menggunakan algoritma Decision Tree dengan akurasi sebesar 76,38%, di mana performa model dievaluasi berdasarkan tiga kelas sentimen yaitu negatif, netral, dan positif, dengan hasil metrik menunjukkan bahwa precision, recall, dan f1-score tertinggi terdapat pada kelas positif, yang mengindikasikan bahwa model lebih unggul dalam mengklasifikasikan sentimen positif dibandingkan sentimen lainnya

Confusion Matrix



Gambar 9. Confusin Matrix

Hasil confusion matrix ini menunjukkan performa model Decision Tree dalam mengklasifikasikan sentimen dengan total 127 data, di mana dari 15 data aktual berlabel negatif, hanya 4 yang diklasifikasikan dengan benar sebagai negatif, sementara 2 diklasifikasikan sebagai netral dan 9 justru salah diklasifikasikan sebagai positif; dari 17 data aktual netral, hanya 4 yang tepat diklasifikasikan, sementara 5 salah menjadi negatif dan 8 menjadi positif; sedangkan dari 95 data aktual positif, model berhasil mengklasifikasikan 85 dengan benar sebagai positif, sementara 7 salah menjadi negatif dan 3 menjadi netral, sehingga dapat disimpulkan bahwa model memiliki kinerja klasifikasi terbaik pada sentimen positif, namun masih lemah dalam membedakan antara sentimen negatif dan netral

4.7. Implementasi Perbandingan Klasifikasi

Evaluasi Model Klasifikasi

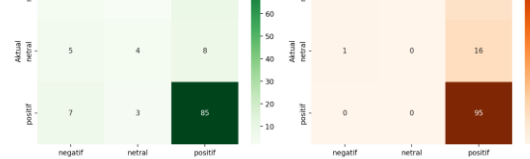
Akurasi Model

Decision Tree: 0.73

Naive Bayes: 0.76

Decision Tree - Confusion Matrix

Naive Bayes - Confusion Matrix



Gambar 10. Perbandingan Confusion Matrix

Gambar confusion matrix menunjukkan bahwa model Decision Tree berhasil mengklasifikasikan

sebagian besar data positif dengan benar (85 dari 95), namun kesulitan dalam mengklasifikasikan sentimen negatif dan netral. Terdapat banyak kesalahan klasifikasi pada dua kelas tersebut, terutama sentimen negatif yang sebagian besar diklasifikasikan sebagai positif.

Sementara itu, model Naïve Bayes menunjukkan performa yang sangat baik dalam mengenali sentimen positif (95 dari 95 data diklasifikasikan dengan benar), namun gagal hampir sepenuhnya dalam mengenali sentimen negatif dan netral. Data pada kedua kelas ini cenderung diklasifikasikan ke dalam kelas positif, yang menandakan terjadinya *bias klasifikasi* terhadap mayoritas kelas.

 Decision Tree

	precision	recall	f1-score	support
negatif	0.25	0.2667	0.2581	15
netral	0.4444	0.2353	0.3077	17
positif	0.8333	0.8947	0.8629	95
accuracy	0.7323	0.7323	0.7323	0.7323
macro avg	0.5093	0.4656	0.4762	127
weighted avg	0.7124	0.7323	0.7172	127

 Naive Bayes

	precision	recall	f1-score	support
negatif	0.5	0.0667	0.1176	15
netral	0	0	0	17
positif	0.76	1	0.8636	95
accuracy	0.7559	0.7559	0.7559	0.7559
macro avg	0.42	0.3556	0.3271	127
weighted avg	0.6276	0.7559	0.6599	127

Gambar 11. Tabel Evaluasi Klasifikasi

Meskipun Naïve Bayes memiliki akurasi keseluruhan yang lebih tinggi, ini disebabkan oleh fakta bahwa model hanya dapat mengidentifikasi sentimen positif, yang merupakan kelas dominan dalam dataset. Sebaliknya, Decision Tree menunjukkan kinerja yang lebih proporsional antar kelas, meskipun memiliki akurasi keseluruhan yang sedikit lebih rendah, seperti yang ditunjukkan oleh nilai macro average F1 yang lebih tinggi pada Decision Tree, yang menunjukkan bahwa model ini lebih mampu mengidentifikasi.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil membangun aplikasi analisis sentimen terhadap topik *Makan Bergizi Gratis* di media sosial X menggunakan metode Decision Tree dalam platform Streamlit, yang mencakup fitur dashboard, preprocessing teks, dan klasifikasi sentimen. Proses preprocessing seperti case folding, cleansing, tokenizing, stopword removal, dan normalisasi kata berperan penting dalam meningkatkan kualitas data sebelum pelatihan model. Model Decision Tree mampu mengklasifikasikan sentimen negatif, netral, dan positif dengan akurasi sebesar 76,38%. Hasil evaluasi menunjukkan model memiliki performa terbaik pada kelas positif dengan f1-score 0.8947, namun kurang optimal pada kelas negatif dan netral yang memiliki f1-score lebih rendah. Ketimpangan distribusi data menyebabkan model

lebih cenderung mengenali sentimen positif, dan juga decision Tree dinilai lebih layak digunakan dalam konteks klasifikasi multi-kelas yang seimbang, terutama saat interpretabilitas dan klasifikasi yang adil untuk semua kategori menjadi fokus utama

DAFTAR PUSTAKA

- [1] I. Zulfahmi, "Analisis Sentimen Aplikasi PLN Mobile Menggunakan Metode Decision Tree," *Jurnal Penelitian Rumpun Ilmu Teknik*, vol. 3, no. 1, pp. 11–21, Dec. 2023.
- [2] Ucu Agustini. (2025). Efektivitas dan Tantangan Kebijakan Program Makan Bergizi Gratis sebagai Intervensi Pendidikan di Indonesia. *Jurnal Kiprah Pendidikan*, 4(3),
- [3] P. Wahyuni and M. A. Romli, "Comparison of Naive Bayes Classifier and Decision Tree Algorithms for Sentiment Analysis on the House of Representatives' Right of Inquiry on Twitter," *Journal of Applied Informatics and Computing*, vol. 8, no. 2, pp. 523–530, Nov. 2024.
- [4] N. C. Ramadani, I. Tahyudin, and A. S. Barkah, "Perbandingan Algoritma Support Vector Machine, Decision Tree, dan Logistic Regression pada Analisis Sentimen Ulasan Aplikasi Netflix," *Jurnal Nasional Teknologi dan Sistem Informasi*, vol. 10, no. 2, pp. 110–117, 2024.
- [5] *View of Analysis of Community Sentiment Towards Free Nutrition Meal Programs on Twitter Using Naïve Bayes, Support Vector Machine, K-Nearest Neighbors, and Ensemble Methods.* (2025).
- [6] Anggy Dwi Anggreny, & Suhardi Suhardi. (2025). Sentiment Analysis of Free Meal Program For School Students Using Algoritma Naive Bayes. *International Journal of Science and Environment (IJSE)*, 5(3), 484–492. <https://doi.org/10.51601/ijse.v5i3.215>
- [7] E. Aprilianto and Y. S. Rahayu, "Comparison of Sentiment Analysis from Twitter Data Collection with Naïve Bayes, Decision Tree, and k-Nearest Neighbor Methods," *Jurnal Ilmiah SINUS*, vol. 22, no. 2, pp. 1–10, 2024.
- [8] Analisis Sentimen Peran Artificial Intelligence Terhadap Kreativitas Dan Efektivitas Mahasiswa Dalam Penyelesaian Tugas Akhir Menggunakan Decision Tree Berbasis Smote. (2025)
- [9] Analisis Sentimen Ulasan Aplikasi Satu Sehat Pada Google Play Store Menggunakan Naïve Bayes Classifier. (2025).
- [10] Jurnal Analisis Klasifikasi Sentimen Pengguna Mypertamina Menggunakan Metode Evaluasi Precision, Recall, Dan F1-Score. (2025).
- [11] Itn.ac.id, "Penerapan Metode Naïve Bayes Classifier pada Analisis Sentimen Aplikasi GoPay," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 1, 2025.
- [12] Analisis Sentimen Komentar Politik Di Media Sosial X Dengan Pendekatan Deep Learning. (2025). Itn.Ac.Id