

PERBANDINGAN METODE NAÏVE BAYES DAN K-NEAREST NEIGHBOR TERHADAP ANALISIS SENTIMEN ULASAN PROGRAM MAKAN SIANG GRATIS DI INDONESIA

Fathoni, Aulia Pinkan Maretta, Aisha Nuraini Kusuma,

Ruth Mei Sasmita, Anabel Fiorenza Rizkyllah, Ali Ibrahim

Fakultas Ilmu Komputer, Sistem Informasi, Universitas Sriwijaya

Jalan Palembang-Prabumulih KM 32 Indralaya, Kabupaten Ogan Ilir, Sumatera Selatan, Indonesia

fathoni@unsri.ac.id

ABSTRAK

Program Makan Siang Gratis di Indonesia sebagai kebijakan kontroversial memicu respons beragam di media sosial, sehingga analisis sentimen diperlukan untuk memahami persepsi publik secara komprehensif. Permasalahan utama terletak pada keterbatasan metode klasifikasi dalam menangani data teks informal, yang berpotensi membuat akurasi identifikasi menurun. Penelitian ini bertujuan membandingkan kinerja algoritma *Naive Bayes* dan *K-Nearest Neighbor* (K-NN) dalam mengklasifikasikan sentimen ulasan masyarakat terhadap program tersebut. Sebanyak 2.080 komentar dari X (Twitter) dan YouTube dikumpulkan melalui *web scraping*, kemudian diproses dengan tahapan *cleaning* (penghapusan *mention*, URL), penghapusan *stopwords*, tokenisasi, dan transformasi fitur menggunakan TF-IDF. Dataset dibagi dengan rasio 60:40 untuk *training* dan *testing*, lalu dievaluasi menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan *Naive Bayes* mencapai akurasi tertinggi (87,75%), lebih unggul dari K-NN (86,67%). Kedua algoritma mencatat *precision* sempurna (100%), namun memiliki kelemahan dalam *recall* (NB: 18,4%; K-NN: 11,2%) dan *F1-score* (NB: 31%; K-NN: 20,1%), yang mengindikasikan kesulitan dalam mengidentifikasi sentimen positif. Penelitian ini membuktikan keunggulan *Naive Bayes* dalam analisis sentimen kebijakan publik berbasis teks informal.

Kata kunci : Analisis Sentimen, Program Makan Siang Gratis, Data Mining, *Naive Bayes*, *K-Nearest Neighbor*

1. PENDAHULUAN

Semenjak pelantikan Presiden dan Wakil Presiden terpilih Indonesia periode 2024-2029, program makan siang gratis masih menjadi topik hangat yang diperbincangkan warga Indonesia. Pasangan Prabowo-Gibran mencanangkan program makan bergizi gratis yang ditargetkan mencapai 82,9 juta penerima manfaat dari golongan anak-anak sekolah dan ibu hamil. Tentunya pelaksanaan program ini menuai tanggapan dari masyarakat, dengan respon dan pandangan yang berbeda-beda. Terdapat pihak yang mendukung kebijakan ini, namun sebagian pihak lainnya menentang kebijakan ini dengan pertimbangan-pertimbangan yang ada seperti alokasi dana, proses distribusi yang merata, maupun keraguan terhadap klaim yang menilai program ini akan menurunkan angka stunting di Indonesia. Perdebatan pro dan kontra yang terjadi diantara masyarakat dapat dilihat melalui platform-platform media sosial. Linimasa berbagai sosial media seperti X dan YouTube berisi ulasan-ulasan yang secara langsung dinyatakan oleh masyarakat dan mencerminkan pandangan mereka terhadap program ini. Opini yang diberikan masyarakat dapat menjadi evaluasi bagi pelaksanaan program makan siang gratis. Dengan begitu, diperlukan analisis lebih lanjut pada komentar-komentar tersebut supaya bisa digunakan sebagai informasi yang komprehensif.

Analisis sentimen merupakan metode untuk mengkategorikan teks menjadi sentimen positif, negatif, atau netral. Teknik *text mining* tersebut memungkinkan peneliti untuk memahami dan mengidentifikasi informasi untuk membantu dalam

menentukan penanganan yang tepat terkait suatu keputusan atau isu yang terjadi [1]. Algoritma *text mining* yang populer diantaranya *Naive Bayes Classifier* dan *K-Nearest Neighbor*, kedua algoritma tersebut memiliki kelebihan masing-masing. Algoritma K-NN dapat mencapai akurasi sebesar 0.80-0.82, lebih tinggi dari *Naive Bayes* dengan akurasi 0.74-0.76 dan membuktikan bagaimana algoritma K-NN lebih mudah dan efektif untuk diimplementasikan [2]. Dalam kasus lainnya, *Naive Bayes* juga dapat mencapai tingkat akurasi yang lebih unggul jika dari K-NN jika menggunakan normalisasi teks dan dataset yang relatif kecil yaitu skor akurasi 87.14% [3]. Penelitian ini bertujuan untuk membandingkan kinerja dari dua algoritma *text mining*, *Naive Bayes Classifier* dan *K-Nearest Neighbor*, dalam mengklasifikasi sentimen masyarakat Indonesia terhadap program makan siang gratis. Diharapkan melalui penelitian ini bisa diperoleh metode mana yang paling akurat dalam mengidentifikasi persepsi publik tentang kebijakan pemerintah.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis sentimen mengumpulkan pendapat masyarakat melalui jejaring sosial, yang mencakup berbagai layanan umum dan masalah terbaru. Analisis sentimen juga dapat digunakan untuk menilai kinerja, pelayanan, dan hal lainnya. Analisis sentimen dalam bentuk dokumen atau kalimat adalah jenis yang paling sering digunakan oleh para peneliti [4].

Analisis sentimen adalah cara untuk menilai emosi, sikap, dan pendapat yang disampaikan oleh

orang-orang tentang barang, merek, layanan, politik, atau organisasi di berbagai media. Data dapat dikategorikan ke dalam kategori seperti sangat positif, positif, netral, negatif, dan sangat negatif berkat penggunaan *machine learning* dan *lexicon-based* [5].

2.2. Data Mining

Data mining adalah suatu proses analisis data yang digunakan untuk menemukan pola, tren, atau informasi penting yang tersembunyi dalam kumpulan data yang besar. Proses ini tidak hanya berfokus pada pengumpulan data, tetapi juga mencakup langkah-langkah penting seperti pemilihan data (atau pemilihan data), pembersihan data (atau preprocessing), transformasi, proses eksplorasi menggunakan algoritma tertentu, dan akhirnya evaluasi hasil. *Data mining* bertujuan untuk menghasilkan informasi baru yang dapat digunakan untuk membuat keputusan [6].

2.3. Naïve Bayes

Berdasarkan prinsip teorema Bayes, algoritma Naïve Bayes menentukan kemungkinan setiap kelas pada data yang dianalisis dengan menggunakan probabilitas bersyarat. Metode ini menganggap bahwa setiap variabel tidak bergantung satu sama lain, tetapi dalam kenyataannya, hal ini tidak selalu terjadi. Meskipun demikian, metode klasifikasi Naïve Bayes masih berhasil dalam banyak kasus klasifikasi. Menurut Fadhlurohman (2025), algoritma ini dapat dengan akurat memperkirakan kemungkinan bahwa suatu objek akan termasuk dalam suatu kategori [7].

2.4. K- Nearest Neighbor

Algoritma pembelajaran terawasi *K-Nearest Neighbor* (K-NN) digunakan untuk tujuan mengklasifikasikan data baru berdasarkan seberapa dekat mereka dengan data pelatihan yang telah ada. Algoritma ini bekerja dengan membandingkan objek uji dengan jumlah K data terdekat yang ada dalam dataset. Selanjutnya, algoritma ini menentukan kelas objek berdasarkan mayoritas label dari objek tetangga terdekatnya. Melainkan membangun model secara eksplisit, K-NN menyimpan seluruh data pelatihan dan melakukan perhitungan selama proses klasifikasi [8]. Rumus jarak *Euclidean* biasanya digunakan untuk menentukan kedekatan antar data, sehingga semakin dekat data uji dengan data latih, semakin besar kemungkinan data tersebut memiliki label yang sama.

2.5. Penelitian Terdahulu

Kajian terhadap penelitian sebelumnya diperlukan untuk memahami metode, hasil, dan keterbatasan yang telah ada. Dalam konteks analisis sentimen program makan siang gratis, studi terdahulu menjadi acuan penting untuk melihat tren pendekatan yang digunakan serta menemukan celah yang dapat dijadikan dasar dalam penelitian ini.

Salah satu penelitian melakukan analisis sentimen pengguna aplikasi X terhadap program makan siang gratis menggunakan algoritma Naïve

Bayes dengan pendekatan SEMMA [5]. Data sejumlah 4.038 *tweet* jangka waktu antara Februari hingga Maret 2024 diambil untuk penelitian tersebut dengan 53,04% *tweet* mengandung sentimen positif dan 44,95% negatif, dengan akurasi model mencapai 80,31%, serta presisi dan *recall* di atas 79% untuk kedua kategori sentimen. Walaupun begitu, pada penelitian ini metode klasifikasi (Naïve Bayes) tidak dibandingkan dengan metode lain, maka dari itu belum didapatkan seberapa besar efektivitas metode lainnya seperti *K-Nearest Neighbor* (K-NN) dalam topik dan lingkup data yang sama.

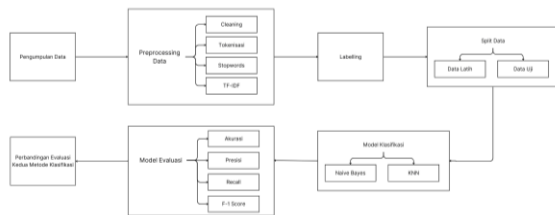
Adapun penelitian lainnya juga menggunakan metode Naïve Bayes untuk analisis opini publik terhadap program makan siang gratis di platform X [9]. Dari total 400 komentar yang diperoleh, jumlah terbanyak sentimen yang ditemukan bersifat negatif (75,31%), sedangkan sentimen positif hanya 9,06%. Model menghasilkan akurasi rendah, yaitu hanya 20% pada data uji, akibat dominasi data negatif dan tidak dilakukannya penyeimbangan data. Dalam studi ini pun belum membandingkan metode alternatif lainnya yang mungkin memiliki ketahanan lebih baik terhadap distribusi data yang tidak seimbang, sehingga memberi peluang untuk menguji reliabilitas metode lain.

Selanjutnya adapun penelitian yang menggunakan performa tiga algoritma klasifikasi antara lain Naïve Bayes, *Support Vector Machine* (SVM), dan *Decision Tree* untuk dibandingkan pada penelitian analisis *tweet* kebijakan makan siang gratis sebanyak 5.205 *tweet* [10]. SMOTE digunakan dalam penelitian ini untuk mengatasi ketimpangan data (4.735 positif dan 470 negatif). Hasil penelitian ini menunjukkan bahwa model SVM mencapai akurasi tertinggi yaitu 99%, diikuti *Decision Tree* dengan 98%, dan Naïve Bayes dengan 91%. Sebelum diterapkan SMOTE, Naïve Bayes hanya memperoleh *recall* 34% pada sentimen negatif, kendati demikian meningkat menjadi 97% setelah SMOTE diterapkan. Dengan hasil penelitian ini, penting untuk melakukan teknik penyeimbangan data, tetapi tidak melibatkan K-NN dalam perbandingan, di sisi lain metode ini cukup potensial untuk pengolahan data berbasis fitur jarak yang relevan dalam klasifikasi sentimen media sosial. Oleh karena itu, terdapat peluang untuk mengevaluasi efektivitas metode K-NN dalam skenario serupa dan membandingkannya langsung dengan Naïve Bayes

3. METODE PENELITIAN

Penelitian ini menggunakan Naïve Bayes dan *K-Nearest Neighbor* (KNN), dua algoritma yang biasa digunakan dalam analisis sentimen berbasis pembelajaran mesin dan menggunakan metode berbeda untuk mengklasifikasikan teks ulasan ke dalam kategori sentimen (positif dan negatif). Untuk mengevaluasi kinerja dalam mengklasifikasikan sentimen masyarakat terhadap kebijakan tersebut, kedua teknik ini akan diterapkan pada data ulasan Program Makan Siang Gratis di Indonesia. Sebelum proses klasifikasi dimulai, data ulasan harus melalui

tahap *preprocessing*. Ini dilakukan untuk meningkatkan kualitas data dan memastikan bahwa teks siap untuk diproses oleh algoritma pembelajaran mesin.



Gambar 1. Tahap metode penelitian

4.1. Data Collecting

Data collecting adalah proses yang dilakukan untuk mengumpulkan data dan informasi [10]. Proses ini dapat digunakan dalam penelitian atau pengambilan keputusan.

4.2. Preprocessing Data

Proses ini mencakup pembersihan data dengan menghilangkan informasi yang tidak diperlukan seperti URL, mention, dan hashtag, serta menghapus karakter yang bukan huruf alfabet [11]. Tahapan yang dibahas dalam penelitian ini adalah tokenisasi, *stopwords*, dan pembersihan.

4.3. Labelling

Pelabelan dalam analisis sentimen adalah proses menetapkan label pada suatu teks sesuai dengan emosi atau sentimen yang terkandung di dalamnya [12].

Dataset dilabeli berdasarkan karakteristik atau kelas yang relevan untuk memungkinkan model untuk mengidentifikasi pola dalam dataset dan membuat prediksi berdasarkan pola tersebut.

4.4. Split Data

Split data yaitu membagi dataset menjadi data Latih (Training) Dan Data Uji (Test). Tujuan utama dari pembagian ini adalah untuk melatih model, menyetel parameter, dan mengevaluasi kinerja model secara objektif, sehingga dapat menghindari masalah seperti *overfitting* dan memastikan bahwa model berfungsi dengan baik pada data yang belum pernah dilihat sebelumnya.

4.5. Model Klasifikasi

Dalam penelitian ini, digunakan dua metode klasifikasi utama, yaitu Naïve Bayes dan *K-Nearest Neighbor* (KNN), untuk menganalisis sentimen ulasan masyarakat terhadap program makan siang gratis di Indonesia.

Naïve Bayes merupakan metode klasifikasi berbasis probabilitistik yang menggunakan Teorema Bayes dengan asumsi bahwa setiap fitur bersifat independen satu sama lain (*independent feature assumption*). Naïve Bayes sering digunakan dalam klasifikasi teks karena kesederhanaannya, efisiensi komputasi, dan kemampuannya bekerja dengan dataset berukuran besar [9].

K-Nearest Neighbor (KNN) adalah metode klasifikasi berbasis jarak yang menentukan kategori suatu data berdasarkan kedekatannya dengan data lain dalam ruang fitur. KNN bekerja dengan mencari *k* tetangga terdekat (*nearest neighbors*) dari suatu data baru dan menentukan kelasnya berdasarkan mayoritas kelas dari tetangga-tetangga tersebut. Metode KNN sering digunakan dalam analisis sentimen karena kemampuannya dalam menangani data non-linear serta tidak memerlukan asumsi distribusi data tertentu [13].

4.6. Model Evaluasi

Untuk mengukur performa model dalam melakukan klasifikasi sentimen dilakukan evaluasi menggunakan empat metrik utama: akurasi, presisi, *recall*, dan skor *f-1*. Tujuan evaluasi ini adalah untuk membandingkan kinerja metode Naïve Bayes dan *K-Nearest Neighbor* (KNN) dalam menganalisis sentimen ulasan terhadap program makan siang gratis di Indonesia.

4.7. Akurasi

Sebuah metrik yang dikenal sebagai akurasi digunakan untuk mengukur kemampuan model untuk mengklasifikasikan data dengan benar.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Akurasi digunakan sebagai indikator awal performa model, tetapi tidak selalu mencerminkan keefektifan model jika dataset memiliki distribusi kelas yang tidak seimbang.

4.8. Presisi

Presisi adalah rasio antara jumlah prediksi positif yang benar terhadap seluruh prediksi positif yang dihasilkan model.

$$Precision = \frac{TP}{TP + FP} \times 100\% \quad (2)$$

Presisi tinggi menunjukkan bahwa model jarang mengklasifikasikan data negatif sebagai positif, sehingga cocok digunakan pada aplikasi di mana kesalahan klasifikasi positif harus dikurangi, seperti dalam analisis opini publik yang sensitif.

4.9. Recall

Recall mengukur seberapa baik model dalam menemukan semua data positif dalam dataset. Nilai *recall* tinggi menunjukkan bahwa model jarang melewatkan data positif yang sebenarnya ada dalam dataset.

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

Recall sangat penting dalam skenario di mana kegagalan dalam mendeteksi data positif bisa berakibat fatal, misalnya dalam analisis tren kebijakan publik di mana semua opini negatif atau positif harus diidentifikasi dengan akurat.

4.10. F-1 Score

F1-Score merupakan salah satu metrik evaluasi dalam sistem informasi yang menggabungkan nilai

recall dan presisi, dengan menunjukkan rata-rata harmonik dari keduanya yang telah diberi bobot secara seimbang [14].

$$F1 = 2x \frac{\text{precision} \times \text{recall}}{\text{precision} + \text{recall}} \times 100\% \quad (4)$$

Secara keseluruhan, *F1-score* adalah metrik yang penting dalam evaluasi model klasifikasi, terutama ketika keseimbangan antara presisi dan *recall* menjadi pertimbangan utama dalam analisis data.

4. HASIL DAN PEMBAHASAN

4.1. Data Collecting

Penelitian ini mengumpulkan data komentar dari dua platform media sosial, yaitu X (Twitter) dan YouTube, melalui teknik web *scraping*. Komentar-komentar yang diambil secara spesifik berkaitan dengan opini masyarakat terhadap program makan siang gratis di Indonesia. Proses pengumpulan data dilakukan dengan mempertimbangkan validitas, relevansi, dan keaslian opini, guna menghindari duplikasi maupun data yang tidak relevan.

Total data yang berhasil dikumpulkan sebanyak 2.080 komentar, terdiri dari 1.063 komentar dari X (Twitter) dan 1.017 komentar dari YouTube. Data yang diperoleh masih belum terstruktur dan mengandung *noise* seperti simbol, emoji, serta kata-kata tidak baku, sehingga perlu dilakukan tahap *preprocessing* sebelum masuk ke proses klasifikasi. Komentar-komentar tersebut kemudian dianalisis dan dikategorikan ke dalam dua kelas sentimen, yaitu positif dan negatif.

4.2. Preprocessing Data

Tahap *pre-processing* merupakan langkah penting dalam pengolahan data teks untuk memastikan kualitas input bagi model klasifikasi. Data komentar yang diperoleh dari platform X (Twitter) dan YouTube masih mengandung berbagai elemen yang tidak relevan, sehingga dilakukan beberapa tahap pembersihan sebagai berikut:

a. Cleaning

Pada tahap ini dilakukan penghapusan elemen-elemen tidak penting seperti mention (@username), URL, tanda baca, angka, dan karakter khusus.

Tabel 1. Hasil Cleaning

Text	Clean_Text
@OposisiCerdas Ya di gaji sama makan siang gratis lah	oposisi cerdas ya di gaji sama makan siang gratis lah

Dari tabel tersebut, dapat dilihat bahwa setelah tahap *cleaning*, elemen-elemen tidak penting seperti mention (@OposisiCerdas) telah dihapus, sehingga hanya tersisa teks yang relevan yaitu "ya di gaji sama makan siang gratis lah".

b. Stopwords

Proses ini menghilangkan kata-kata umum yang tidak memiliki kontribusi besar terhadap makna (*stopwords*).

Tabel 2. Hasil Stopwords

Clean_text	Stopwords
oposisicerdas ya di gaji sama makan siang gratis lah	oposisicerdas ya gaji makan siang gratis

Dari tabel tersebut, dapat dilihat bahwa setelah tahap *stopwords removal*, kata-kata umum seperti "di", "sama", "lah" telah dihapus.

c. Tokenization

Setelah dibersihkan, kalimat dipecah menjadi satuan kata (*token*) menggunakan metode *tokenizing*.

Tabel 3. Hasil Tokenization

Clean_text	Tokenization
oposisi cerdas ya gaji makan siang gratis	[oposisi cerdas, ya, gaji, makan, siang, gratis]

Dari tabel tersebut, dapat dilihat bahwa setelah proses *tokenisasi*, kalimat yang sudah dibersihkan dipisahkan menjadi kata-kata individu (*token*).

d. TF IDF

Setelah teks diproses, setiap dokumen komentar direpresentasikan dalam bentuk numerik menggunakan metode TF-IDF.

	aamiin	abab	abil	abis	acara	accountability	action	ad	adab
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
4	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

Gambar 2. Hasil TF-IDF

Gambar di atas menunjukkan hasil dari representasi TF-IDF, di mana setiap kata dalam korpus diberi nilai numerik berdasarkan frekuensinya.

4.3. Labelling

Pada tahap *labelling*, setiap teks yang telah diproses akan diberi label berdasarkan sentimen yang terkandung dalam teks tersebut.

Tabel 4. Hasil Tokenization

Text	Label
bank indonesia dukung program makan siang gratis	Positif

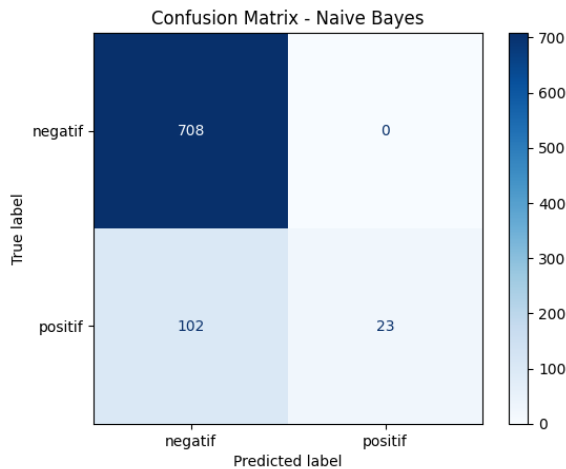
Tabel tersebut menunjukkan contoh hasil *tokenization* pada teks "bank indonesia dukung program makan siang gratis" yang diberikan label "Positif." Dalam hal ini, label tersebut menunjukkan bahwa teks tersebut mengandung sentimen positif terhadap program makan siang gratis.

4.4. Split Data

Dalam penelitian ini, dataset yang telah dikumpulkan dibagi menjadi dua bagian utama, yaitu data pelatihan (*training data*) dan data pengujian (*testing data*). Pembagian data dilakukan dengan rasio 60:40, di mana 60% dari data digunakan untuk pelatihan dan 40% sisanya digunakan untuk pengujian.

Data pelatihan terdiri dari 1249 sampel, sedangkan data uji coba terdiri dari 833 sampel.

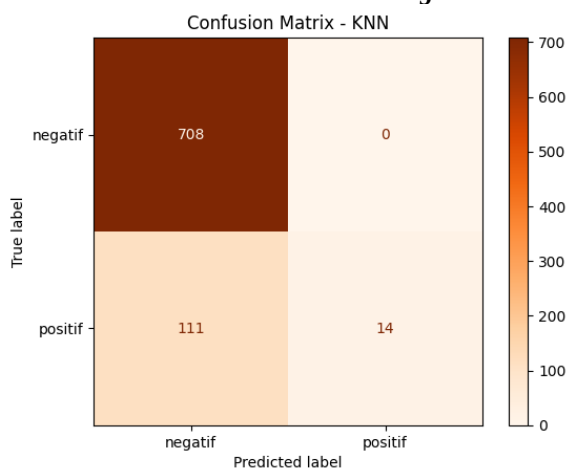
4.5. Model Klasifikasi Naïve Bayes



Gambar 3. Hasil Naïve Bayes

Berdasarkan hasil evaluasi model Naïve Bayes yang ditunjukkan melalui *confusion matrix*, terlihat bahwa model ini memiliki kinerja yang sangat baik dalam mengklasifikasikan kelas negatif, tetapi masih mengalami kesulitan dalam mengidentifikasi kelas positif. Model berhasil memprediksi dengan benar 708 *instance* sebagai kelas negatif (*True Negatives/TN*) dan tidak ada *instance* negatif yang salah diklasifikasikan sebagai positif (*False Positives/FP*). Hal ini menunjukkan tingkat akurasi yang tinggi dalam mengenali kelas negatif. Namun, di sisi lain, model hanya berhasil memprediksi 23 *instance* dengan benar sebagai kelas positif (*True Positives/TP*), sementara 102 *instance* positif lainnya salah diprediksi sebagai negatif (*False Negatives/FN*).

4.6. Model Klasifikasi K-Nearest Neighbor



Gambar 4. Hasil K-Nearest Neighbor

Berdasarkan *confusion matrix* yang diperoleh dari evaluasi model K-Nearest Neighbors (KNN), dapat diamati bahwa model ini menunjukkan kinerja yang baik dalam mengklasifikasikan kelas negatif

namun masih memiliki keterbatasan dalam memprediksi kelas positif. Secara spesifik, model berhasil mengidentifikasi dengan benar 708 *instance* sebagai kelas negatif (*True Negatives/TN*) tanpa adanya kesalahan prediksi kelas negatif sebagai positif (*False Positives/FP*). Hal ini mengindikasikan tingkat akurasi yang tinggi untuk klasifikasi kelas negatif. Di sisi lain, model hanya mampu memprediksi 14 *instance* dengan benar sebagai kelas positif (*True Positives/TP*), sementara 111 *instance* positif lainnya salah diklasifikasikan sebagai negatif (*False Negatives/FN*).

4.7. Model Evaluasi Naïve Bayes

- Accuracy

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} = \frac{23 + 708}{23 + 708 + 0 + 102} = \frac{731}{833} = 87.75\% \quad (5)$$

- Precision

$$Precision = \frac{TP}{TP + FP} \times 100\% = \frac{23}{23 + 0} \times 100\% = 100\% \quad (6)$$

- Recall

$$Recall = \frac{TP}{TP + FN} \times 100\% = \frac{23}{23 + 102} \times 100\% = 18.4\% \quad (7)$$

- F-1 Score

$$F-1 \text{ Score} = 2 \times \frac{Precision \times Recall}{Precision + Recall} = 2 \times \frac{1.0 \times 0.184}{1.0 + 0.184} = 31\% \quad (8)$$

4.8. Model Evaluasi K-Nearest Neighbor

- Accuracy

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} = \frac{14 + 708}{14 + 708 + 0 + 111} = \frac{722}{733} = 86.67\% \quad (9)$$

- Precision

$$Precision = \frac{TP}{TP + FP} \times 100\% = \frac{14}{14 + 0} \times 100\% = 100\% \quad (10)$$

- Recall

$$Recall = \frac{TP}{TP + FN} \times 100\% = \frac{14}{14 + 111} \times 100\% = 11.2\% \quad (11)$$

- F-1 Score

$$F-1 \text{ Score} = 2 \times \frac{Precision \times Recall}{Precision + Recall} = 2 \times \frac{1.0 \times 0.112}{1.0 + 0.112} = 20.1\% \quad (12)$$

4.9. Perbandingan Model Evaluasi

Setelah melakukan perhitungan terhadap metrik evaluasi seperti *Accuracy*, *Precision*, *Recall*, dan *F-1 Score* pada kedua model, berikut adalah perbandingan hasil evaluasi antara Naïve Bayes dan K-Nearest Neighbors (KNN) berdasarkan perhitungan yang telah dilakukan.

Tabel 5. Perbandingan Model Evaluasi

Model	Accuracy	Precision	Recall	F1-score
Naïve Bayes	87,75%	100%	18,4%	31%
KNN	86,67%	100%	11,2%	20,1%

Tabel ini menunjukkan hasil evaluasi dari dua model yang digunakan dalam analisis sentimen, yaitu Naïve Bayes dan *K-Nearest Neighbor* (KNN). Berdasarkan hasil yang ditampilkan, model Naïve Bayes memiliki nilai akurasi tertinggi sebesar 87,75%, namun nilai *recall* dan *F1-score*-nya lebih rendah dibandingkan dengan akurasi. Sementara itu, model KNN menunjukkan nilai akurasi yang sedikit lebih rendah (86,67%), dengan *precision* yang sama yaitu 100%, tetapi *recall* dan *F1-score*-nya juga lebih rendah dibandingkan Naïve Bayes.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian terhadap 2.080 komentar dari X (Twitter) dan YouTube mengenai program makan siang gratis di Indonesia, metode Naïve Bayes menunjukkan performa yang lebih baik dibandingkan *K-Nearest Neighbor* (KNN) dalam analisis sentimen, dengan akurasi tertinggi sebesar 87,75%. Meskipun kedua metode memiliki *precision* yang sama, yaitu 100%, model Naïve Bayes lebih unggul secara keseluruhan karena nilai akurasi yang lebih tinggi. Proses *preprocessing* yang melibatkan *cleaning*, *tokenisasi*, *stopwords removal*, dan TF-IDF membantu meningkatkan kualitas data sebelum diklasifikasikan. Untuk penelitian selanjutnya, disarankan untuk mencoba algoritma lain seperti *Support Vector Machine* (SVM) atau *Random Forest* guna memperkaya perbandingan performa, serta memperluas sumber data dan mempertimbangkan pendekatan *ensemble* agar hasil klasifikasi sentimen semakin akurat dan aplikatif dalam pengambilan keputusan kebijakan publik.

DAFTAR PUSTAKA

- [1] J. R. Jim, M. A. R. Talukder, P. Malakar, M. M. Kabir, K. Nur, and M. F. Mridha, "Recent advancements and challenges of NLP-based sentiment analysis: A state-of-the-art review," *Nat. Lang. Process. J.*, vol. 6, p. 100059, 2024, doi: 10.1016/j.nlp.2024.100059.
- [2] H. Bojda and D. Gala, "Analysis the performance of Naive Bayes and K-Nearest Neighbor Classifiers," *CEUR Workshop Proc.*, vol. 3885, pp. 366–373, 2024.
- [3] N. Charibaldi, A. Harfiani, and O. Samuel Simanjuntak, "Comparison of the Effect of Word Normalization on Naïve Bayes Classifier and K-Nearest Neighbor Methods for Sentiment Analysis," *Inf. J. Ilm. Bid. Teknol. Inf. dan Komun.*, vol. 9, no. 1, pp. 25–31, 2023, doi: 10.25139/inform.v9i1.7111.
- [4] M. Syarifuddin, "Analisis Sentimen Opini Publik Mengenai Covid-19 Pada Twitter Menggunakan Metode Naïve Bayes Dan Knn," *INTI Nusa Mandiri*, vol. 15, no. 1, pp. 23–28, 2020, doi: 10.33480/inti.v15i1.1347.
- [5] L. Nursinggah, R. Ruuhwan, and T. Mufizar, "Analisis Sentimen Pengguna Aplikasi X Terhadap Program Makan Siang Gratis Dengan Metode Naïve Bayes Classifier," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, 2024, doi: 10.23960/jitet.v12i3.4336.
- [6] M. Sutra Safira, N. Rahaningsih, and R. Danar Dana, "Penerapan Data Mining Untuk Klasifikasi Penjualan Obat Menggunakan Metode K-Nearest Neighbor," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 380–385, 2024, doi: 10.36040/jati.v8i1.8325.
- [7] N. A. Fadhlurrohman, A. Primajaya, A. N. Dimiyati, U. S. Karawang, T. Timur, and J. Barat, "ANALISIS SENTIMEN TERHADAP SKEMA STUDENT LOAN UNTUK BIAYA PERGURUAN TINGGI PADA TWITTER," vol. 9, no. 2, pp. 2115–2123, 2025.
- [8] F. Farahdinna, M. N. Shofy, U. Nasional, and B. Iris, "IMPLEMENTASI K-NEAREST NEIGHBOR UNTUK KLASIFIKASI BUNGA IRIS," vol. 9, no. 2, pp. 3510–3514, 2025.
- [9] W. Anggriyani and M. Fakhriza, "Analisis Sentimen Program Makan Gratis Pada Media Sosial X Menggunakan Metode NLP," vol. 5, no. 4, pp. 1033–1042, 2024, doi: 10.47065/josyc.v5i4.5826.
- [10] E. Yuspita and R. R. Suryono, "Indonesian Journal of Computer Science," *Indones. J. Comput. Sci.*, vol. 13, no. 5, pp. 8447–8457, 2024, [Online]. Available: <http://ijcs.stmikindonesia.ac.id/ijcs/index.php/ijcs/article/view/3135>
- [11] M. R. Firdaus, N. Rahaningsih, and R. D. Dana, "Analisis Sentimen Aplikasi Shopee di Goole Play Store Menggunakan Klasifikasi Algoritma Naïve Bayes," *J. Inform. dan Rekayasa Perangkat Lunak*, vol. 6, no. 1, pp. 228–237, 2024.
- [12] O. N. Julianti, N. Suarna, and W. Prihartono, "Penerapan Natural Language Processing Pada Analisis Sentimen Judi Online Di Media Sosial Twitter," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 3, pp. 2936–2941, 2024, doi: 10.36040/jati.v8i3.9613.
- [13] Syahril Dwi Prasetyo, Shofa Shofiah Hilabi, and Fitri Nurapriani, "Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naïve Bayes dan KNN," *J. KomtekInfo*, vol. 10, pp. 1–7, 2023, doi: 10.35134/komtekinfo.v10i1.330.
- [14] T. Arifqi, N. Suarna, and W. Prihartono, "Penggunaan Naive Bayes Dalam Menganalisis Sentimen Ulasan Aplikasi Mcdonald'S Di Indonesia," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 2, pp. 1949–1956, 2024, doi: 10.36040/jati.v8i2.8740.