

ANALISIS SENTIMEN PROGRAM MAKAN BERGIZI GRATIS MENGGUNAKAN SVM DAN KNN DENGAN TF-IDF DAN TF-ABS

Sevira Hukman Majidah^{1§}, Putriaji Hendikawati²

¹Program Studi Matematika, Universitas Negeri Semarang [Email: svrmajidah@students.unnes.ac.id]

²Program Studi Matematika, Universitas Negeri Semarang [Email: putriaji.mat@mail.unnes.ac.id]

[§]Corresponding Author

ABSTRACT

The Free Nutritious Meal Program, initiated by the government to improve child and maternal health, has generated varied public responses on social media. This study aims to classify public sentiment toward the program by applying machine learning models to social media text data. The dataset used consists of 9000 tweets from Platform X that have been manually labeled into positive, negative, and neutral categories. Classification was performed using Support Vector Machine (SVM) and K-Nearest Neighbor (KNN) algorithms combined with two term weighting techniques Term Frequency-Inverse Document Frequency (TF-IDF) and Term Frequency-Absolute (TF-ABS). The performance of each model was evaluated using accuracy, precision, recall, and F1-score. The results show that SVM model with TF-ABS achieved the best performance with 99.69% accuracy, precision 99.69%, recall 99.68%, and F1-score 99.68%. The KNN model with TF-ABS also performed well, reaching 95.55% accuracy. In contrast, models employing TF-IDF demonstrated noticeably lower performance, with the SVM achieving an accuracy of [75,53%] and the KNN reaching [70,89%], indicating a clear performance gap compared to the TF-ABS weighting scheme. This research provides insights into suitable machine learning models and term weighting methods for sentiment analysis of public opinion on government programs using social media data.

Keywords: sentiment analysis, SVM, KNN, TF-IDF, TF-ABS

1. PENDAHULUAN

Makanan bergizi merupakan salah satu kebutuhan pokok yang harus dipenuhi untuk mendukung tumbuh kembang dan kesehatan masyarakat, khususnya anak-anak. Apabila kebutuhan gizi tidak dipenuhi, kondisi ini akan berdampak pada kesehatan dan perkembangan generasi muda yang menjadi harapan masa depan bangsa. Menyadari persoalan ini, pasangan calon presiden Prabowo Subianto dan Gibran Rakabuming Raka pada Pemilu 2024 menggagas Program Makan Bergizi Gratis sebagai salah satu program prioritas. Program ini dirancang untuk menyediakan makanan bergizi secara gratis kepada sekitar 82,9 juta anak-anak sekolah, balita, dan ibu hamil, dengan harapan membangun generasi emas Indonesia 2045 yang sehat, cerdas, dan berdaya saing tinggi (Sitanggang et al., 2024).

Sejak awal diperkenalkan, program ini menjadi topik yang ramai diperbincangkan masyarakat. Mulai dari apresiasi dan dukungan

terhadap program, hingga kritik terhadap distribusi manfaat, keberlanjutan program, serta besarnya anggaran yang diperkirakan mencapai 460 triliun rupiah. Media sosial X kini menjadi sarana utama bagi masyarakat untuk menyuarakan pendapat mereka secara terbuka. Berdasarkan data Asosiasi Penyelenggara Jasa Internet Indonesia (APJII), pengguna internet Indonesia mencapai 221 juta pada 2024, dan lebih dari 60% di antaranya aktif menggunakan media sosial X (Nursinggh et al., 2024). Hal ini menjadikan media sosial sebagai sumber data yang potensial untuk mengukur sentimen publik terhadap kebijakan ini.

Seiring dengan perkembangan teknologi, analisis sentimen berbasis *machine learning* menjadi salah satu pendekatan yang banyak digunakan untuk memetakan opini publik dari data teks. Analisis sentimen bertujuan untuk mengklasifikasikan opini ke dalam kategori sentimen tertentu, seperti positif, negatif, atau

netral. Proses ini umumnya melibatkan beberapa tahapan, mulai dari *text pre-processing*, pembobotan kata, hingga penerapan algoritma klasifikasi (Prakash & Aloysius, 2021).

Salah satu algoritma yang banyak digunakan adalah *Support Vector Machine* (SVM) dan *K-Nearest Neighbor* (KNN). SVM bekerja dengan mencari *hyperplane* terbaik yang dapat memisahkan data dari kelas yang berbeda dengan *margin* maksimal. SVM dikenal efektif untuk menangani data dengan dimensi tinggi dan mampu menghasilkan klasifikasi yang akurat, terutama dalam kasus yang melibatkan data teks (Sari & Salsabila, 2024). Di sisi lain, KNN melakukan klasifikasi dengan menghitung jarak antara data uji dan data latih, kemudian menetapkan kelas berdasarkan mayoritas dari k tetangga terdekat (Başarslan & Kayaalp, 2024).

Sejumlah penelitian telah mengevaluasi kinerja kedua algoritma tersebut pada berbagai kasus analisis sentimen. Penelitian yang dilakukan oleh Sari & Salsabila (2024) membandingkan SVM dan KNN dalam mengklasifikasikan sentimen publik terhadap vaksinasi booster kedua COVID-19. Hasil penelitian tersebut menunjukkan bahwa KNN sedikit lebih unggul dengan akurasi 85,48% dan presisi 78,64%, sementara SVM mencatat akurasi 84,33% dan presisi 78,34%. Selain itu, Mustaqim et al. (2024) juga membandingkan kinerja algoritma SVM, Naive Bayes, dan KNN dalam klasifikasi sentimen. Hasil penelitian menunjukkan bahwa KNN memberikan performa terbaik dengan akurasi mencapai 99,21%.

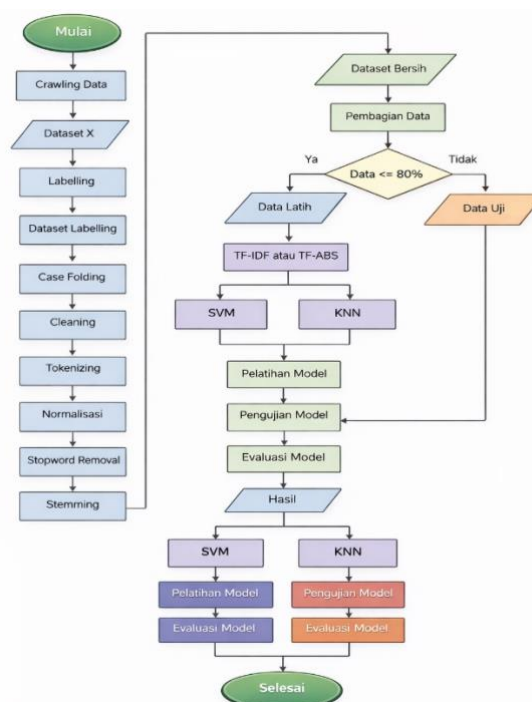
Hasil analisis sentimen yang optimal tidak hanya bergantung pada algoritma klasifikasi, melainkan juga pada strategi pembobotan kata yang digunakan. Pembobotan kata bertujuan untuk mengubah data teks menjadi representasi numerik agar dapat diproses oleh algoritma *machine learning*. Salah satu teknik yang paling umum adalah *Term Frequency - Inverse Document Frequency* (TF-IDF). TF-IDF menilai pentingnya sebuah kata berdasarkan seberapa sering kata tersebut muncul dalam dokumen tertentu dibandingkan dengan keseluruhan dokumen di korpus. Selain TF-IDF, terdapat teknik alternatif bernama *Term Frequency - Absolute* (TF-ABS) yang dikembangkan untuk memperbaiki kelemahan TF-IDF. TF-ABS memanfaatkan transformasi logit dari *odds ratio* untuk mengukur hubungan antara kata dan kelas kategori secara lebih seimbang.

Penelitian Supono & Suprayogi (2021) mengkaji penerapan kedua teknik pembobotan tersebut pada teks layanan *helpdesk* Direktorat Jenderal Kekayaan Negara (DJKN) dengan menggunakan algoritma KNN. Hasil penelitian menunjukkan bahwa kinerja TF-ABS lebih baik dibandingkan TF-IDF dengan nilai akurasi sebesar 90,04%.

Berdasarkan temuan sebelumnya, penelitian ini bertujuan untuk mengombinasikan algoritma SVM dan KNN dengan dua teknik pembobotan kata TF-IDF dan TF-ABS dalam menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis. Tujuan utama penelitian ini adalah untuk mengukur dan membandingkan efektivitas masing-masing kombinasi metode berdasarkan metrik evaluasi seperti akurasi, presisi, *recall*, dan *F1-score*, sehingga dapat memberikan rekomendasi metode yang paling tepat dalam memahami opini publik berbasis data teks dari media sosial.

2. METODE PENELITIAN

Penelitian ini dilakukan melalui sejumlah tahapan penting yang alurnya dapat dilihat pada Gambar 1. Proses dimulai dengan pengumpulan data, kemudian dilanjutkan dengan tahap *labelling*, *preprocessing*, pembobotan kata, pelatihan dan pengujian model klasifikasi, hingga evaluasi model untuk mengukur kinerja klasifikasi yang dihasilkan.



Gambar 1. Alur Penelitian

2.1 Pengumpulan Data

Penelitian ini memanfaatkan data berupa *tweet* dari Platform X yang memuat kata kunci “Makan Bergizi Gratis” dengan periode pengambilan data dari 1 hingga 31 Januari 2025. Data tersebut dikumpulkan melalui teknik *crawling*. Dari hasil *crawling*, diperoleh 9.000 *tweet* yang terdiri dari 15 atribut.

2.2 Labelling

Tahap pelabelan dilakukan untuk memberikan informasi awal mengenai jenis sentimen yang terkandung dalam setiap *tweet*. Sentimen dibagi ke dalam tiga kategori utama, yaitu positif, negatif, dan netral. Pada penelitian ini, pelabelan dilakukan secara manual dengan menelaah secara saksama isi setiap *tweet*, memperhatikan konteks dan makna kalimat agar label yang diberikan sesuai dengan maksud sebenarnya dari opini yang disampaikan.

2.3 Preprocessing

Tahap *preprocessing* bertujuan untuk membersihkan dan merapikan data teks yang semula tidak terstruktur menjadi lebih rapi, terstruktur, dan siap digunakan untuk tahap analisis berikutnya (Saputra & Hasan, 2024). Proses ini mencakup beberapa langkah penting, seperti *case folding*, *cleaning*, *tokenization*, *normalization*, *stopwords removal*, dan *stemming*.

2.4 Pembobotan TF-IDF dan TF-ABS

Pembobotan merupakan proses untuk menentukan seberapa penting atau seberapa besar kontribusi setiap kata dalam dataset terhadap model klasifikasi (Prasetyo et al., 2023). Sebelum dilakukan proses pembobotan, dataset terlebih dahulu dibagi menjadi dua bagian, yakni data latih dan data uji, dengan proporsi 80% untuk pelatihan model dan 20% sisanya untuk pengujian performa model. Dalam penelitian ini digunakan dua teknik pembobotan, yaitu *Term Frequency-Inverse Document Frequency* (TF-IDF) dan *Term Frequency-Absolute* (TF-ABS).

TF-IDF bekerja dengan menghitung bobot setiap kata berdasarkan seberapa sering kata tersebut muncul dalam sebuah dokumen, dikombinasikan dengan seberapa jarang kata tersebut muncul di seluruh dokumen (Rasyid et al., 2024). Semakin sering suatu kata muncul dalam satu dokumen, bobot TF-nya semakin tinggi. Sementara itu, IDF akan memberikan bobot lebih besar pada kata-kata yang jarang

muncul di keseluruhan dokumen. TF-IDF dirumuskan sebagai berikut.

$$TF_{(t,d)} = f_{(t,d)} \quad (1)$$

$$IDF_t = \log\left(\frac{N}{df_t}\right) \quad (2)$$

$$TF - IDF_{t,d} = TF_{(t,d)} \times IDF_t \quad (3)$$

Di mana t merupakan *term* atau kata, d adalah dokumen, $f_{(t,d)}$ adalah jumlah kemunculan kata t dalam dokumen d , N adalah jumlah dokumen, dan df_t adalah jumlah dokumen yang memuat kata t .

Sementara itu, TF-ABS dikembangkan sebagai alternatif TF-IDF untuk meningkatkan akurasi klasifikasi dengan memperhitungkan keterkaitan kata dengan kategori sentimen. Berbeda dari TF-IDF, TF-ABS menggantikan komponen IDF dengan nilai absolut dari logit *odds ratio* (Matsunaga & Ebecken, 2008). Bobot TF-ABS dihitung dengan mengalikan frekuensi kata dengan nilai absolut logit dari perbandingan kemunculan kata pada satu kategori dan di luar kategori tersebut (Supono & Suprayogi, 2021). Perhitungan TF-ABS dirumuskan sebagai berikut.

$$ABS(t_j, c_k) = \left| \ln \frac{(n_{kj}+0,5)(n_{\bar{k}j}+0,5)}{(n_{\bar{k}j}+0,5)(n_{kj}+0,5)} \right| \quad (4)$$

$$TF - ABS(t_j, c_k) = TF_{(t,d)} \times ABS(t_j, c_k) \quad (5)$$

Dengan n_{kj} adalah jumlah dokumen pada kategori c_k yang memuat *term* t_j , $n_{\bar{k}j}$ adalah jumlah dokumen tidak pada kategori c_k yang memuat *term* t_j , n_{kj} adalah jumlah dokumen pada kategori c_k yang tidak memuat *term* t_j , dan $n_{\bar{k}j}$ adalah jumlah dokumen tidak pada kategori c_k yang tidak memuat *term* t_j .

2.5 Klasifikasi SVM dan KNN

Setelah proses pembobotan dilakukan, langkah berikutnya adalah membangun model klasifikasi untuk mengelompokkan *tweet* ke dalam kategori sentimen. Penelitian ini menggunakan dua algoritma *machine learning*, yakni *Support Vector Machine* (SVM) dan *K-Nearest Neighbor* (KNN).

SVM merupakan algoritma klasifikasi yang bekerja berdasarkan konsep sebuah permukaan yang disebut *hyperplane*. *Hyperplane* berfungsi sebagai batas pemisah yang memaksimalkan jarak antar kelas (Pratiwi et al., 2021). Tujuan SVM adalah mencari *hyperplane* terbaik untuk memisahkan data ke dalam kelas-kelas yang berbeda dengan jarak pemisah (margin) maksimal. Secara matematis, *hyperplane* pada SVM dirumuskan sebagai berikut.

$$w \cdot x_i + b = 0 \quad (6)$$

Dengan w adalah vektor bobot, x adalah vektor fitur, dan b adalah bias. Klasifikasi dilakukan dengan menentukan apakah data memenuhi:

$$w \cdot x_i + b \geq +1 \quad \text{untuk } y_i = +1 \quad (7)$$

$$w \cdot x_i + b \leq -1 \quad \text{untuk } y_i = -1 \quad (8)$$

Hyperplane yang dipilih adalah yang memiliki margin terbesar terhadap data dari kedua kelas. Untuk kasus data non-linear, SVM memanfaatkan fungsi *kernel* guna memetakan data ke ruang berdimensi lebih tinggi sehingga dapat dipisahkan secara linear.

Sementara itu, KNN adalah algoritma klasifikasi yang mengkategorikan data uji berdasarkan mayoritas label dari Presisi mengukur seberapa tepat prediksi positif yang dilakukan model, yaitu rasio antara jumlah *true positive* dengan total prediksi positif, yang dirumuskan tetangga terdekatnya (Wisnu et al., 2020). Jarak antara data uji dan data latih dihitung menggunakan *Euclidean distance* dengan rumus berikut.

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (9)$$

Kelas data uji kemudian ditentukan berdasarkan kelas yang paling banyak muncul di antara tetangga terdekat tersebut. Pemilihan nilai K sangat berpengaruh terhadap hasil klasifikasi; nilai K kecil bisa membuat model sensitif terhadap *noise*, sedangkan nilai K besar dapat mengaburkan batas antar kelas.

2.6 Evaluasi

Evaluasi model merupakan langkah penting untuk mengukur kinerja model klasifikasi dalam memetakan data ke dalam kategori yang tepat sesuai tujuan penelitian (Purnamasari et al., 2023). Salah satu teknik yang banyak digunakan dalam evaluasi performa klasifikasi adalah *confusion matrix*. *Confusion matrix* memuat informasi mengenai jumlah prediksi benar dan salah dari setiap kelas yang diklasifikasikan (Musfiroh et al., 2021).

Tabel 1. *Confusion Matrix*

<i>Actual</i>	<i>Precision</i>	
	<i>Positive</i>	<i>Negative</i>
<i>Positive</i>	<i>True Positive (TP)</i>	<i>False Negative (FN)</i>
<i>Negative</i>	<i>False Positive (FP)</i>	<i>True Negative (TN)</i>

Berdasarkan komponen *confusion matrix*, sejumlah metrik evaluasi seperti akurasi, *precision*, *recall*, dan *F1-score* dapat dihitung untuk mengukur kinerja model. Akurasi mengukur proporsi prediksi yang benar terhadap seluruh prediksi yang dilakukan. Nilai akurasi didefinisikan oleh Persamaan (10).

$$Accuracy = \frac{TP+TN}{TP+FP+FN+TN} \times 100\% \quad (10)$$

Precision mengukur seberapa tepat prediksi positif yang dilakukan model, yaitu rasio antara jumlah *true positive* dengan total prediksi positif. *Precision* didefinisikan oleh Persamaan (11).

$$Precision = \frac{TP}{TP+FP} \times 100\% \quad (11)$$

Recall digunakan untuk mengetahui kemampuan model dalam mengenali seluruh data positif yang sebenarnya. *Recall* didefinisikan oleh Persamaan (12).

$$Recall = \frac{TP}{TP+FN} \quad (12)$$

Sementara itu, *F1-Score* merupakan ukuran keseimbangan antara *precision* dan *recall*. Nilai *F1-score* didefinisikan oleh Persamaan (13).

$$F1 - score = \frac{2 \times recall \times precision}{recall + precision} \quad (13)$$

3. HASIL DAN PEMBAHASAN

3.1 Pengumpulan Data

Data yang digunakan dalam penelitian ini merupakan data *tweet* dari Platform X yang memuat kata kunci “Makan Bergizi Gratis”. Data dikumpulkan menggunakan teknik *crawling* dengan *twitter_auth_token*. Dari proses *crawling* yang dilakukan, berhasil dikumpulkan sebanyak 9.000 *tweet* dengan 15 atribut. Untuk mempermudah analisis, hanya kolom *full_text* yang digunakan. Data yang telah disaring kemudian disimpan dalam format *csv*. Contoh data hasil *crawling* ditampilkan pada Tabel 2.

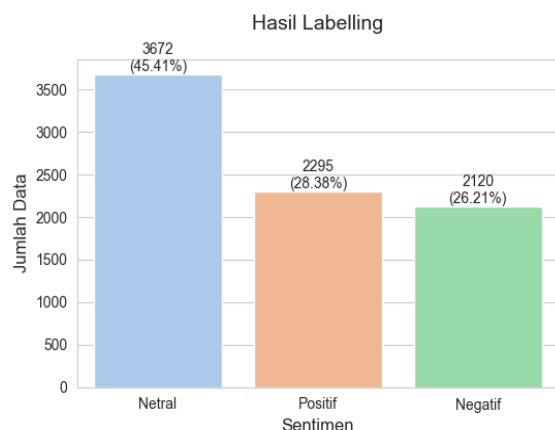
Tabel 2. Contoh Data Hasil *Crawling*

No	<i>full_text</i>
1	makan bergizi gratis ngapain sih anjir so asik bngt jadi prioritas utama kayak yg bergizi beneran aja tu makanan
2	Makan Bergizi Gratis dapat Sambutan Hangat dari Warganet Global https://t.co/J1tSNoVJXs

3.2 Labelling

Pelabelan terhadap 9000 data dilakukan secara manual untuk mengklasifikasikan sentimen setiap *tweet* ke dalam tiga kategori,

yaitu positif, negatif, dan netral. Distribusi *tweet* hasil pelabelan menurut kategori sentimen ditampilkan pada Gambar 2.



Gambar 2. Hasil Labelling

Grafik tersebut menunjukkan bahwa sentimen netral mendominasi dengan 3.672 data, diikuti positif sebanyak 2.295 dan negatif 2.120 *tweet*.

3.3 Preprocessing

Sebelum data digunakan dalam analisis sentimen, dilakukan tahap *preprocessing* untuk membersihkan dan menyiapkan teks agar dapat diproses secara optimal oleh model klasifikasi. Tahap ini mencakup beberapa langkah utama, yakni *case folding*, *cleaning*, *tokenization*, *normalization*, *stopword removal*, dan *stemming*.

1. Case Folding

Case folding merupakan proses mengubah seluruh huruf dalam teks menjadi huruf kecil. Tujuannya untuk mengurangi jumlah token unik dan dalam data sehingga dimensi fitur menjadi lebih sederhana (Prasetyo et al., 2023).

2. Cleaning

Cleaning bertujuan untuk membersihkan data dari elemen yang tidak relevan seperti tanda baca, *hashtag*, *mention*, *URL*, symbol, serta karakter khusus yang tidak mendukung analisis.

3. Tokenization

Tokenization merupakan tahap memisahkan teks menjadi bagian-bagian kecil yang disebut token. Token ini dapat berupa kata, frasa, maupun simbol.

4. Normalization

Normalization bertujuan untuk mengubah teks ke dalam bentuk yang sesuai standar bahasa baku berdasarkan Kamus Besar Bahasa Indonesia (KBBI).

5. Stopword Removal

Proses ini bertujuan untuk menghilangkan kata-kata umum yang sering muncul tetapi tidak memiliki kontribusi signifikan dalam analisis, seperti “yang”, “dan”, atau “di”.

6. Stemming

Stemming adalah proses mengembalikan kata ke bentuk dasar dengan menghapus awalan, akhiran, maupun imbuhan lain yang melekat pada kata.

3.4 Pembobotan

Dataset pada penelitian ini dibagi menjadi dua bagian dengan rasio 80:20, di mana 80% untuk data latih dan 20% untuk data uji. Hasil pembagian dataset ditampilkan pada Tabel 3.

Tabel 3. Hasil Pembagian Dataset

	Data Train	Data Test
Presentase	80%	20%
Jumlah	6468	1618

Setelah dilakukan pembagian dataset, proses selanjutnya adalah pembobotan kata menggunakan TF-IDF dan TF-Abs. TF-IDF menghitung bobot kata berdasarkan frekuensi kemunculannya dalam dokumen (TF) dan seberapa jarang kata tersebut muncul di seluruh dokumen (IDF). Semakin sering sebuah kata muncul pada dokumen tertentu namun jarang pada dokumen lainnya, bobotnya akan semakin tinggi. Penerapan TF-IDF pada penelitian ini dilakukan menggunakan *TfidfVectorizer* dari pustaka *scikit-learn*. Berdasarkan hasil perhitungan menggunakan Python, diperoleh 10 kata dengan nilai bobot TF-IDF tertinggi yang disajikan pada Tabel 4.

Tabel 4. 10 Kata Bobot TF-IDF Tertinggi

No	Kata	TF-IDF
1	makan	0.0742
2	gizi	0.0703
3	gratis	0.0678
4	program	0.0487
5	anak	0.0244
6	dukung	0.0228
7	indonesia	0.0194
8	mbak	0.0189
9	prabowo	0.0175
10	sehat	0.0175

Selanjutnya, dilakukan pembobotan kedua, yaitu TF-ABS, yang menghitung bobot dengan mempertimbangkan distribusi kata di dalam dan di luar kelas tertentu. Proses pembobotan TF-ABS diawali dengan membentuk representasi *Term Frequency* (TF) menggunakan pustaka *CountVectorizer*. Nilai TF dihitung dengan mengukur seberapa sering suatu kata muncul dalam suatu dokumen. Setelah mendapatkan nilai TF, selanjutnya dihitung distribusi dokumen yang memuat atau tidak memuat kata tersebut, baik pada kelas yang dimaksud maupun di luar kelas. Hasil perhitungan ini selanjutnya digunakan untuk mendapatkan nilai ABS. Bobot akhir TF-ABS diperoleh dari hasil perkalian nilai TF dan ABS untuk setiap kata. Sepuluh kata dengan bobot TF-ABS tertinggi disajikan pada Tabel 4.5.

Tabel 5. 10 Kata Bobot TF-ABS Tertinggi

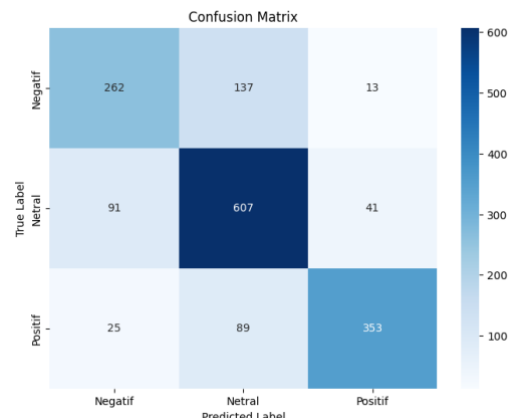
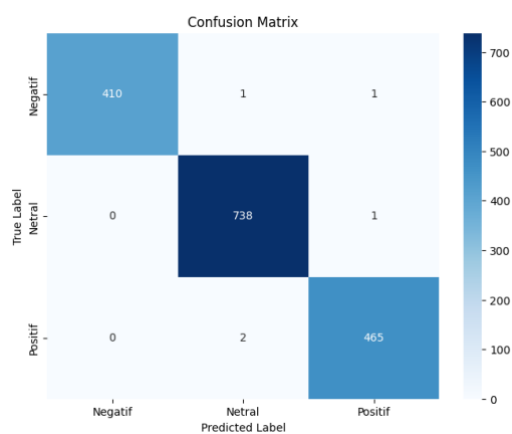
No	Kata	TF-ABS
1	makan	1.2073
2	gizi	0.7923
3	gratis	0.6110
4	dukung	0.2610
5	sinergiindonesiajepang	0.2391
6	program	0.1795
7	jepang	0.1460
8	indonesia	0.1376
9	sehat	0.1051
10	pm	0.0960

Kata “*sinergiindonesiajepang*” memiliki bobot TF-ABS yang relatif tinggi karena berkaitan dengan isu atau pemberitaan yang sedang ramai dibahas pada periode pengambilan data. Istilah tersebut diduga muncul secara masif dalam konteks pemberitaan atau percakapan publik mengenai kerja sama antara Indonesia dan Jepang, khususnya yang dikaitkan dengan dukungan terhadap Program Makan Bergizi Gratis.

3.5 Klasifikasi

Klasifikasi terhadap data hasil preprocessing dan pembobotan TF-IDF serta TF-ABS dilakukan menggunakan algoritma SVM dan KNN. Pada bagian berikut, akan dibahas terlebih dahulu penerapan SVM dalam klasifikasi sentimen, kemudian dilanjutkan dengan pembahasan KNN.

Pada model SVM, klasifikasi dilakukan menggunakan kernel linear dari pustaka *sklearn.svm*. Kernel linear dipilih karena cocok untuk data teks yang memiliki dimensi fitur tinggi dan mendukung pemisahan linear. Model dilatih dengan parameter *default* $C = 1$. Pelatihan model SVM dilakukan untuk dua jenis pembobotan, yaitu TF-IDF dan TF-ABS. Setelah model SVM dilatih, dilakukan prediksi pada data uji untuk mengukur performa model. Hasil *confusion matrix* model SVM terhadap data uji untuk setiap pembobotan ditampilkan pada Gambar 3 dan Gambar 4.

Gambar 3. Hasil *Confusion Matrix* SVM dengan TF-IDFGambar 4. Hasil *Confusion Matrix* SVM dengan TF-ABS

Berdasarkan hasil prediksi, metrik evaluasi dihitung untuk menilai performa model. *Precision*, *recall*, dan *F1-score* dihitung dengan pendekatan *weighted average* agar sesuai dengan distribusi data sentimen yang tidak seimbang. Ringkasan hasil evaluasi disajikan pada Tabel 6.

Tabel 6. Evaluasi Performa SVM

Metrik	TF-IDF	TF-ABS
Akurasi	75,53%	99,69%
<i>Precision</i>	75,96%	99,69%
<i>Recall</i>	75,52%	99,68%
<i>F1-score</i>	75,48%	99,68%

Tabel 6 menunjukkan bahwa SVM dengan pembobotan TF-ABS memberikan hasil yang unggul dibandingkan TF-IDF. Model dengan TF-ABS mencapai akurasi 99,69%, dengan nilai *precision*, *recall*, dan *F1-score* di atas 99%. Sebaliknya, SVM dengan TF-IDF hanya memperoleh akurasi 75,53%, dengan *precision*, *recall*, dan *F1-score* yang juga lebih rendah.

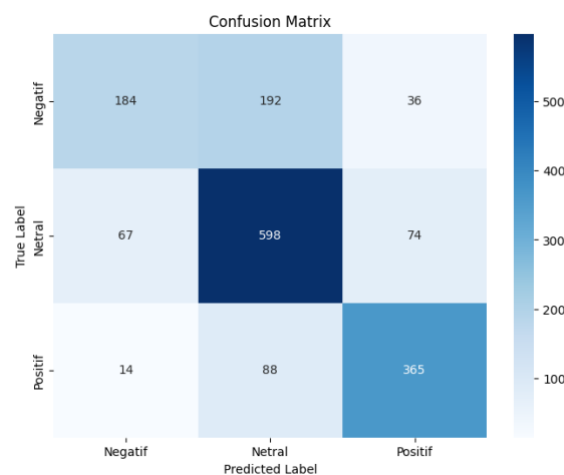
Selanjutnya, klasifikasi dilakukan menggunakan algoritma KNN dengan metode pembobotan TF-IDF dan TF-ABS. KNN merupakan algoritma klasifikasi berbasis jarak yang menentukan kelas suatu data uji dengan melihat mayoritas kelas dari K tetangga terdekatnya. Salah satu faktor penting dalam KNN adalah pemilihan nilai K , yang menentukan jumlah tetangga terdekat yang digunakan untuk mengklasifikasikan data uji. Pemilihan nilai K dilakukan melalui proses *tuning* parameter dengan *5-fold cross validation* yang menguji nilai K ganjil dari 1 hingga 20 untuk menghindari hasil imbang. Hasil *tuning* parameter pada berbagai nilai K disajikan pada Tabel 7 berikut.

Tabel 7. Hasil Evaluasi pada Berbagai Nilai K

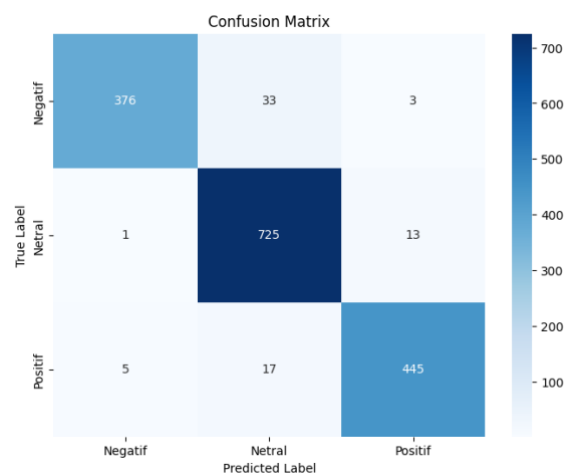
Nilai K	Akurasi	
	KNN dengan TF-IDF	KNN dengan TF-ABS
$K=1$	65,54%	95,73%
$K=3$	67,32%	95,78%
$K=5$	68,34%	96,01%
$K=7$	69,05%	96,01%
$K=9$	69,05%	96,01%
$K=11$	69,45%	95,92%
$K=13$	69,16%	95,92%
$K=15$	69,62%	96,07%
$K=17$	69,67%	96,40%
$K=19$	70,10%	96,51%

Berdasarkan Tabel 7, nilai K terbaik untuk kedua pembobotan diperoleh pada $K = 19$, dengan akurasi sebesar 70,10% untuk TF-IDF dan 96,51% untuk TF-ABS. Temuan ini menunjukkan bahwa TF-ABS mampu memberikan bobot kata yang lebih relevan

terhadap sentimen sehingga mendukung prediksi yang lebih akurat dibandingkan TF-IDF. Nilai K optimal yang diperoleh dari proses *tuning* kemudian digunakan untuk pelatihan akhir model KNN. Setelah model dilatih, dilakukan prediksi terhadap data uji dan evaluasi untuk mengetahui performa model. Hasil *confusion matrix* model KNN untuk setiap pembobotan ditampilkan pada Gambar 5 dan Gambar 6.



Gambar 5. Hasil *Confusion Matrix* KNN dengan TF-IDF



Gambar 6. Hasil *Confusion Matrix* KNN dengan TF-ABS

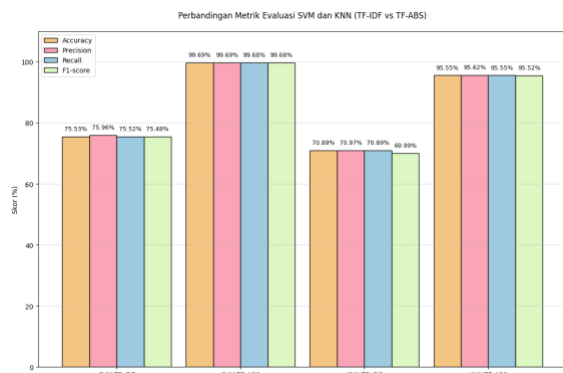
Berdasarkan *confusion matrix* di atas, dilakukan perhitungan metrik evaluasi untuk kedua model. Hasil evaluasi performa KNN dengan pembobotan TF-IDF dan TF-ABS ditampilkan pada Tabel 8 berikut.

Tabel 8. Evaluasi Performa KNN

Metrik	TF-IDF	TF-ABS
Akurasi	70,89%	95,55%
<i>Precision</i>	70,97%	95,62%
<i>Recall</i>	70,89%	95,55%
<i>F1-score</i>	69,99%	95,52%

Berdasarkan Tabel 8, terlihat bahwa KNN dengan TF-ABS memberikan performa yang jauh lebih baik dibandingkan TF-IDF di seluruh metrik evaluasi dengan akurasi sebesar 95,55%, *precision* sebesar 95,62%, *recall* sebesar 95,55%, dan *F1-score* sebesar 95,52%. Sementara itu, KNN dengan TF-IDF menghasilkan akurasi 70,89%, *precision* 70,97%, *recall* 70,89%, dan *F1-score* 69,99%. Hasil ini menunjukkan bahwa TF-ABS lebih efektif dalam mendukung KNN untuk membedakan kategori sentimen pada data.

Untuk mengevaluasi kinerja model, dilakukan perbandingan terhadap empat kombinasi metode klasifikasi dan pembobotan fitur. Perbandingan dilakukan menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*.



Gambar 7. Perbandingan Performa Kombinasi SVM dan KNN dengan TF-IDF dan TF-ABS

Berdasarkan grafik pada Gambar 7, kombinasi SVM dengan TF-ABS menunjukkan kinerja terbaik dibandingkan kombinasi lainnya, dengan akurasi mencapai 99,69%, *precision* 96,69%, *recall* 99,68%, dan *F1-score* 99,68%. Model KNN dengan TF-ABS menempati peringkat kedua dengan akurasi sebesar 95,55%, *precision* sebesar 95,62%, *recall* sebesar 95,55%, dan *F1-score* sebesar 95,52%. Sebaliknya, kedua model dengan pembobotan TF-IDF menunjukkan kinerja yang lebih rendah. SVM dengan TF-IDF hanya mencapai akurasi 75,53%, sementara KNN dengan TF-IDF

memperoleh akurasi terendah, yakni 70,89%. Perbedaan kinerja ini menegaskan bahwa pemilihan metode pembobotan kata sangat berpengaruh terhadap kualitas klasifikasi sentimen. SVM terbukti lebih unggul dibandingkan KNN karena karakteristiknya yang lebih tangguh dalam menangani data teks berdimensi tinggi dan bersifat jarang. SVM mampu membentuk *hyperplane* optimal dengan margin maksimum untuk memisahkan kelas sentimen secara lebih efektif, sedangkan KNN yang berbasis perhitungan jarak cenderung kurang stabil pada ruang fitur berdimensi besar. Kombinasi antara kemampuan SVM dan pembobotan TF-ABS menghasilkan peningkatan akurasi model secara signifikan.

Distribusi hasil klasifikasi terbaik menunjukkan bahwa sentimen netral mendominasi dengan jumlah 738 data, diikuti sentimen positif sebanyak 465 data, dan negatif sebanyak 410 data. Temuan ini mengindikasikan bahwa opini publik terhadap Program Makan Bergizi Gratis bersifat variatif, dengan adanya kombinasi dukungan, masukan, serta pengawasan terhadap implementasi kebijakan di lapangan.

4. KESIMPULAN

Berdasarkan hasil penelitian, dapat disimpulkan bahwa kombinasi model SVM dengan pembobotan TF-ABS memberikan hasil klasifikasi sentimen terbaik dibandingkan kombinasi metode lainnya. Model ini mencapai tingkat akurasi, *precision*, *recall*, dan *F1-score* yang sangat tinggi, serta secara konsisten menunjukkan performa yang lebih baik dibandingkan pembobotan TF-IDF pada kedua algoritma yang diuji.

Hasil klasifikasi sentimen menunjukkan bahwa opini publik terhadap Program Makan Bergizi Gratis didominasi oleh sentimen netral, diikuti oleh sentimen positif dan negatif. Temuan ini mencerminkan bahwa respons masyarakat terhadap program tersebut bersifat beragam, mencakup dukungan, masukan, maupun kritik terhadap pelaksanaan kebijakan di lapangan.

DAFTAR PUSTAKA

- Başarslan, M. S., & Kayaalp, F. (2024). Sentiment analysis of coronavirus data with ensemble and machine learning methods. *Turkish Journal of Engineering*, 8(2), 175–185.

- <https://doi.org/10.31127/tuje.1352481>
- Matsunaga, L. A., & Ebecken, N. F. F. (2008). Term weighting approaches for text categorization improving. *Proceedings - 8th International Conference on Intelligent Systems Design and Applications, ISDA 2008*, 1, 409–414.
<https://doi.org/10.1109/ISDA.2008.21>
- Musfiroh, D., Khaira, U., Utomo, P. P. E., & Suratno, T. (2021). Sentiment Analysis of Online Lectures in Indonesia from Twitter Dataset Using InSet Lexicon. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 1(1), 24–33.
- Mustaqim, I. Z., Puspasari, H. M., Utami, A. T., Syalevi, R., & Ruldeviyani, Y. (2024). Assessing public satisfaction of public service application using supervised machine learning. *IAES International Journal of Artificial Intelligence*, 13(2), 1608–1618.
<https://doi.org/10.11591/ijai.v13.i2.pp1608-1618>
- Nursingah, L., Ruuhwan, R., & Mufizar, T. (2024). ANALISIS SENTIMEN PENGGUNA APLIKASI X TERHADAP PROGRAM MAKAN SIANG GRATIS DENGAN METODE NAÏVE BAYES CLASSIFIER. *Jurnal Informatika Dan Teknik Elektro Terapan*, 12(3).
<https://doi.org/10.23960/jitet.v12i3.4336>
- Prakash, T. N., & Aloysius, A. (2021). *Textual Sentiment Analysis using Lexicon Based Approaches* (Vol.25).
<http://annalsofrscb.ro>
- Prasetyo, D. S., Hilabi, S. S., & Nurapriani, F. (2023). Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naïve Bayes dan KNN. *Jurnal KomtekInfo*, 1–7.
<https://doi.org/10.35134/komtekinfo.v10i1.330>
- Pratiwi, R. W., Febbi, H. S., Dairoh, Af'idah, D. I., Rifa, A. Q., & Gian, F. A. (2021). Analisis Sentimen Pada Review Skincare Female Daily Menggunakan Metode Support Vector Machine(SVM). *Journal of Informatics, Information System, Software Engineering and Applications*, 1(1), 040–046.
- Purnamasari, D., Bayu, A., Desy, A., Fanka, W. A. P., Reza, A., Safrila, M., Yanda, O. N., & Hidayati, U. (2023). *Pengantar Metode Analisis Sentimen*.
- Rasyid, R. Al, Handayani, D., & Ningsih, U. (2024). Penerapan Algoritma TF-IDF dan Cosine Similarity untuk Query Pencarian Pada Dataset Destinasi Wisata. *Jurnal Teknologi Informasi Dan Komunikasi*, 8(1), 2024. <https://doi.org/10.35870/jti>
- Saputra, R., & Hasan, F. N. (2024). Analisis Sentimen Terhadap Program Makan Siang & Susu Gratis Menggunakan Algoritma Naive Bayes. *Jurnal Teknologi Dan Sistem Informasi Bisnis*, 6(3), 411–419.
<https://doi.org/10.47233/jteksis.v6i3.1378>
- Sari, S. F., & Salsabila, I. (2024). Comparison Of K-Nearest Neighbor and Support Vector Machine for Sentiment Analysis of The Second Covid-19 Booster Vaccination. *Journal of Applied Statistics and Data Science*, 1(1), 58–70.
- Sitanggang, A., Umaidah, Y., & Adam, R. I. (2024). Analisis Sentimen Masyarakat Terhadap Program Makan Siang Gratis Pada Media Sosial X Menggunakan Algoritma Naïve Bayes. *Jurnal Informatika Dan Teknik Elektro Terapan*, 12(3).
<https://doi.org/10.23960/jitet.v12i3.4902>
- Supono, R. A., & Suprayogi, M. A. (2021). Perbandingan Metode TF-ABS dan TF-IDF Pada Klasifikasi Teks Helpdesk Menggunakan K-Nearest Neighbor. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(5), 911–918.
<https://doi.org/10.29207/resti.v5i5.3403>
- Wisnu, H., Afif, M., & Ruldevyani, Y. (2020). Sentiment analysis on customer satisfaction of digital payment in Indonesia: A comparative study using KNN and Naïve Bayes. *Journal of Physics: Conference Series*, 1444(1).
<https://doi.org/10.1088/1742-6596/1444/1/012034>