

Perbandingan Berbagai Metode Klasifikasi Teks Untuk Sentimen Kebijakan Makan Gratis di Indonesia

Emi Yuspita¹, Ryan Randy Suryono²

emi_yuspita@teknokrat.ac.id¹, ryan@teknokrat.ac.id²

^{1,2} Sistem Informasi, Fakultas Teknik dan Ilmu Komputer, Universitas Teknokrat Indonesia

Informasi Artikel

Diterima : 15 Okt 2024
Direvisi : 23 Okt 2024
Disetujui : 30 Okt 2024

Kata Kunci

Kebijakan Makan Gratis, Analisis Sentimen, *Naïve Bayes*, *SVM*, *Decision Tree*.

Abstrak

Kebijakan makan gratis merupakan inisiatif penting untuk meningkatkan gizi anak balita dan ibu hamil serta mengurangi ketidaksetaraan sosial. Kebijakan ini mendukung keluarga rendah dengan menyediakan makanan dan susu gratis di sekolah dan pesantren. Di media sosial, terutama platform X (*Twitter*), kebijakan ini memicu diskusi publik. Tujuan dari penelitian ini adalah untuk menganalisis sentimen mengenai kebijakan makan gratis dengan menggunakan metode *Naïve Bayes*, *SVM*, dan *Decision Tree*, serta memberikan efektivitas algoritma klasifikasi dalam memahami opini publik. Dari 5.205 tweet yang dianalisis, terdapat 4.735 tweet positif dan 470 tweet negatif. Penerapan SMOTE pada analisis ini memberikan hasil yang signifikan. *SVM* mencapai akurasi 99%, *Decision Tree* juga menunjukkan performa baik dengan akurasi 98%. Sementara itu, *Naïve Bayes* mengalami peningkatan akurasi hingga 91%, meskipun masih kurang optimal dalam mendeteksi sentimen negatif dibandingkan dengan *SVM* dan *Decision Tree*.

Keywords

Free Meal Policy, Sentiment Analysis, *Naïve Bayes*, *SVM*, *Decision Tree*.

Abstract

The free meal policy is an important initiative to improve the nutrition of children under five and pregnant women and reduce social inequality. This policy supports low-income families by providing free food and milk in schools and Islamic boarding schools. On social media, especially platform X (*Twitter*), this policy sparked public discussion. This research aims to analyze sentiment regarding the free meal policy using *Naive Bayes*, *SVM*, and *Decision Tree* methods, as well as providing the effectiveness of classification algorithms in understanding public opinion. Of the 5,205 tweets analyzed, there were 4,735 positive tweets and 470 negative tweets. Applying SMOTE to this analysis provides significant results. *SVM* achieved 99% accuracy, *Decision Tree* also showed good performance with 98% accuracy. Meanwhile, *Naive Bayes* experienced an increase in accuracy of up to 91%, although it was still less than optimal in detecting negative sentiment compared to *SVM* and *Decision Tree*.

A. Pendahuluan

Kebijakan makan gratis adalah langkah penting yang bertujuan untuk meningkatkan status gizi anak balita dan ibu hamil, serta mengurangi kesenjangan sosial. Kebijakan ini mendukung keluarga rendah dengan menyediakan akses makanan dan susu gratis di sekolah dan pesantren [1]. Inisiatif ini sejalan dengan visi misi yang diusung oleh pasangan calon presiden dan wakil presiden, Prabowo dan Gibran. Mereka merencanakan agar kebijakan ini menjangkau sekitar 82 juta orang, termasuk balita dan ibu hamil. Oleh karena itu, diharapkan dapat memberikan manfaat yang signifikan dalam meningkatkan kesehatan masyarakat. Dengan adanya program ini, diharapkan kualitas gizi anak-anak dapat meningkat, yang pada gilirannya dapat berdampak positif pada perkembangan mereka di masa depan.

Kebijakan ini, seperti halnya kebijakan lainnya, juga mendapat banyak tanggapan dari masyarakat di media sosial. Media sosial memainkan peran krusial dalam menyebarkan informasi mengenai kebijakan pemerintah dan berfungsi sebagai wadah diskusi publik di mana masyarakat dapat mengemukakan pendapat mereka [2]. Salah satu platform yang menonjol adalah X (Twitter), yang telah menjadi tempat utama bagi masyarakat untuk berbagi pandangan dan tanggapan, termasuk terkait kebijakan makan gratis ini. Dengan beragamnya respons yang muncul, penting untuk melakukan analisis sentimen masyarakat secara lebih terstruktur guna memahami sejauh mana kebijakan ini diterima oleh masyarakat.

Analisis sentimen adalah metode untuk mengukur dan mengklasifikasikan pendapat orang terhadap suatu topik tertentu. Analisis ini biasanya membagi opini menjadi dua kategori utama: sentimen positif dan sentimen negatif [3]. Dalam konteks ini, algoritma klasifikasi seperti Naive Bayes, SVM, dan Decision Tree banyak digunakan dalam klasifikasi teks dan analisis sentimen karena kinerja dan efisiensinya dalam memproses data dalam jumlah besar, termasuk teks media sosial. Berdasarkan penelitian sebelumnya, algoritme ini terbukti efektif dalam konteks analisis sentimen, namun performanya bervariasi bergantung pada data dan karakteristik kumpulan data. Misalnya, SVM terbukti lebih akurat dalam beberapa kasus, sedangkan Naive Bayes lebih baik dalam hal akurasi. Selain itu, ketiga algoritma ini diketahui beradaptasi dengan baik pada analisis data teks media sosial yang umumnya memiliki variasi bahasa, panjang teks, dan struktur kalimat tidak beraturan, yang relevan dengan penelitian ini.

Dalam penelitian yang dilakukan oleh Chindy Aulia Sari, Annisa Sukmawati, Rishani Putri Aprilli, Prais Sarah Kayaningtias, dan Novanto Yudistira dengan judul Perbandingan Metode *Naïve Bayes*, *Support Vector Machine*, dan *Decision Tree* dalam Klasifikasi Konsumsi Obat, hasil akurasi yang diperoleh masing-masing adalah 0,50 untuk metode *Naïve Bayes*, 0,41 untuk *Decision Tree*, dan 0,53 untuk *Support Vector Machine* (SVM). Berdasarkan hasil tersebut, algoritma SVM terbukti memiliki akurasi tertinggi dibandingkan kedua algoritma lainnya [4].

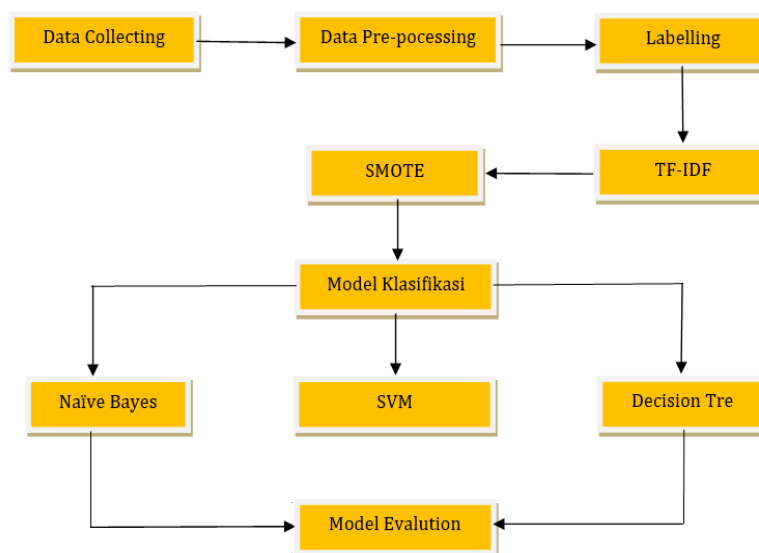
Penelitian yang dilakukan oleh Rima Tamara Aldisa dan Pandu Maulana tentang analisis sentimen opini terhadap masyarakat vaksinasi booster COVID-19 dengan menggunakan tiga metode, yaitu *Naive Bayes*, *Decision Tree*, dan *SVM*. Hasil penelitian menunjukkan bahwa SVM memiliki performa rata-rata terbaik dibandingkan dua model lainnya. Namun, untuk skor presisi, model Naive Bayes unggul dengan nilai 83,81%. Selain itu, dilakukan analisis terhadap jumlah

sentimen dalam dataset, di mana sentimen positif mencapai 2248 dan sentimen negatif berjumlah 751. Berdasarkan konfusi matriks, model Naive Bayes menunjukkan kinerja yang baik dengan nilai *True Positive* (TP) sebesar 1224 dan *True Negative* (TN) sebesar 252, yang lebih besar dibandingkan dengan *False Positive* (FP) sebesar 114 dan *False Negatif* 210[5].

Penelitian ini bertujuan untuk menganalisis terkait kebijakan makan gratis di Indonesia serta mengevaluasi efektivitas berbagai algoritma klasifikasi dalam mengidentifikasi opini publik.

B. Metode Penelitian

Pada penelitian ini penulis menggunakan beberapa tahapan, sehingga penelitian ini memiliki alur kerja dari awal sampai akhir. Tahapan penelitian ini dapat dilihat pada gambar 1. berikut.



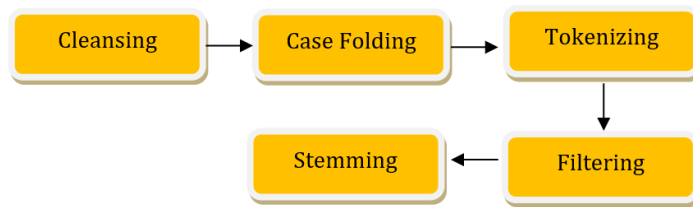
Gambar 1. Tahapan Penelitian

1. Data Collecting

Data collecting adalah proses yang dilakukan untuk mengumpulkan data dan informasi [6]. Pada tahap ini data diambil melalui sosial media X dengan menggunakan metode *crawling*, yang memanfaatkan bahasa pemrograman python. Dengan menganakan kata kunci “makan siang gratis”. Sebanyak 5205 tweet . Dataset ini berisi teks berbahasa indonesia yang yang selanjunya disimpan dalam bentuk csv.

2. Data Preprocessing

Data Preprocessing adalah serangkaian langkah atau prosedur yang dilakukan untuk mempersiapkan dan membersihkan data mentah sebelum diaplikasikan ke model atau algoritma analisis data [7]. Ada beberapa tahapan yang dilakukan dalam preprocessing. Tahapan ini dapat dilihat pada gambar 2 berikut.



Gambar 2. Tahapan Pada Data Pre-Pocessing

2.1 Cleansing

Tahap pertama pada data *preprocessing* adalah *cleansing*. *Cleansing* adalah proses membersihkan data mentah dari kesalahan, ketidakakuratan, dan ketidak konsistenan [8].

Tabel 1. Hasil *Cleansing*

Tweet	Cleansing
@DS_yantie @NenkMonica Dan kalian masih percaya	DSyantie NenkMonica Dan kalian masih percaya

2.2 Case Folding

Case Folding adalah perubahan karakter dalam format teks, biasanya mengubah semua huruf besar menjadi kecil[9]. tujuannya adalah untuk menyamakan representasi teks agar analisis yang dilakukan tidak dipengaruhi oleh perbedaan penggunaan huruf besar dan kecil.

Tabel 2. Hasil *Case Folding*

Tweet	Case Folding
@DS_yantie @NenkMonica Dan kalian masih percaya	dsyantie nenkmonica dan kalian masih percaya

2.3 Tokenizing

Tokenizing adalah proses dalam pemrosesan bahasa alami (NLP) yang bertujuan untuk memecah teks atau kalimat menjadi unit-unit yang lebih kecil.

Tabel 3. Hasil *Tokenizing*

Tweet	Tokenizing
@DS_yantie @NenkMonica Dan kalian masih percaya	[ds, yantie, nenkmonica, dan, kalian, masih, percaya]

2.4 Filtering

Filtering digunakan untuk menyeleksi teks yang tidak relevan dari dataset[10]. Proses ini dapat mencakup penghapusan dokumen atau bagian teks yang tidak berhubungan dengan topik atau tujuan analisis.

Tabel 4. Hasil *Filtering*

Tweet	Filtering
@DS_yantie @NenkMonica Dan kalian masih percaya	[ds, yantie, nenkmonica, kalian, percaya]

2.5 Stemming

Stemming adalah tahap akhir dalam preprocessing yang bertujuan untuk menghilangkan semua imbuhan yang terdapat di awal, akhir, dan kombinasi keduanya pada suatu kata atau kalimat[11].

Tabel 5. Hasil Stemming

Tweet	Stemming
@DS_yantie @NenkMonica Dan kalian masih percaya	[ds, yantie, nenkmonica, kalian, percaya]

3. Labelling

Pada tahap pelabelan, dilakukan pemberian label pada data teks yang telah diproses [12]. biasanya dengan kategori seperti negatif dan positif. Proses ini penting dalam analisis sentimen karena membantu dalam mengklasifikasikan opini.

Tabel 6. Hasil Labelling

Tweet	Labelling
ds yantie nenkmonica dan kalian masih percaya	Positif
ini yang paling saya takutkan kalau makan siang gratis tanpa prosedur	Negatif

4. TF-IDF

TF-IDF membantu mengevaluasi pentingnya kata dalam sebuah dokumen dengan mempertimbangkan tidak hanya jumlah dokumen dalam koleksi, tetapi juga jumlah kemunculan kata dalam dokumen.

5. SMOTE

SMOTE adalah metode yang digunakan untuk mengatasi masalah ketidakseimbangan[13]. Di mana jumlah data dalam satu kelompok jauh lebih banyak dibandingkan dengan kelompok lainnya.

6. Model Klasifikasi

Penelitian ini berusaha menerapkan tiga algoritma klasifikasi yang berbeda untuk menganalisis sentimen pada data tweet. Algoritma yang diterapkan dalam penelitian ini adalah naïve bayes, Support Vector Machine (svm), decision tree. yang bertujuan membandingkan dan mengevaluasi kinerja dari masing-masing metode klasifikasi dalam analisis sentimen terhadap data tweet.

7. Model Evaluation

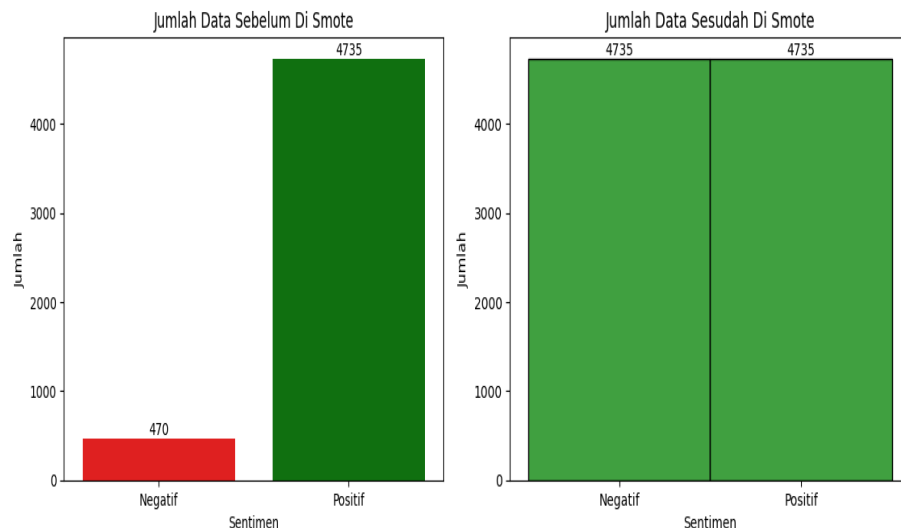
Evaluasi model merupakan langkah penting dalam analisis sentimen, tujuan utamanya adalah menganalisis hasil prediksi sentimen yang dihasilkan model berdasarkan rating pada dataset. Proses evaluasi mencakup beberapa metrik seperti akurasi, presisi, recall, dan F1-score, yang digunakan untuk mengevaluasi seberapa efektif model dalam mengklasifikasikan sentimen [14].

C. Hasil dan Pembahasan

1. Dataset

Dalam penelitian yang telah dilakukan, jumlah seluruh data yang dianalisis adalah 5205 tweets. Proses pelabelan menunjukkan bahwa jumlah

data dengan sentimen positif mencapai 4735, sementara sentimen negatif hanya berjumlah 470. Perbedaan yang signifikan antara jumlah data sentimen positif dan negatif ini menunjukkan adanya ketidakseimbangan yang perlu ditangani. Oleh karena itu, dilakukan optimasi menggunakan SMOTE (Synthetic Minority Over-sampling Technique) untuk menyeimbangkan dataset. Dengan penerapan SMOTE, diharapkan model klasifikasi seperti *Naive Bayes*, *SVM*, dan *Decision Tree* tidak hanya mendominasi satu sentimen, tetapi mampu mengenali dan mengklasifikasikan kedua sentimen secara lebih seimbang. Hasil perbandingan kinerja model sebelum dan sesudah penerapan SMOTE dapat dilihat dalam Gambar 3 di bawah ini.



Gambar 3. Perbandingan SMOTE

Berdasarkan Gambar 3 di atas, terlihat perbedaan yang signifikan antara jumlah data sebelum dan sesudah penerapan teknik SMOTE. sebelum di SMOTE, jumlah data untuk kategori negatif adalah 470, sedangkan untuk kategori positif mencapai 4.735. Ketidakseimbangan ini menunjukkan adanya dominasi data positif, yang dapat menyebabkan model algoritma lebih sulit untuk belajar dari data minoritas (negatif). setelah penerapan SMOTE, jumlah data untuk kedua kategori menjadi seimbang, dengan masing-masing kategori berjumlah 4.735 tweet. Dengan cara ini, model dapat mempelajari sentimen dengan lebih baik, tanpa adanya bias terhadap kategori mayoritas.

Berdasarkan model algoritma yang diterapkan, persentase data yang digunakan adalah 80% dan latih dan 20% data uji. Model ini menggunakan data yang telah dioptimalkan dengan teknik SMOTE, serta membandingkannya dengan data sebelum dioptimalkan. Dengan pendekatan ini, hasil penerapan algoritma dapat langsung diamati. Hasil dari model algoritma dapat dilihat pada Tabel 8 berikut ini:

Tabel 7. Modeling Sebelum Menggunakan SMOTE

<i>Modelling</i>	<i>Sentimen</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recal</i>	<i>F1-Score</i>
<i>Naive</i>	<i>Negatif</i>	0.92	0.81	0.34	0.48
<i>Bayes</i>	<i>Positif</i>		0.93	0.99	0.96
<i>SVM</i>	<i>Negatif</i>	0.98	0.98	0.82	0.89

Decision Tree	Positif	0.98	0.98	1.00	0.99
	Negatif		0.95	0.90	0.92
	Positif		0.99	0.99	0.99

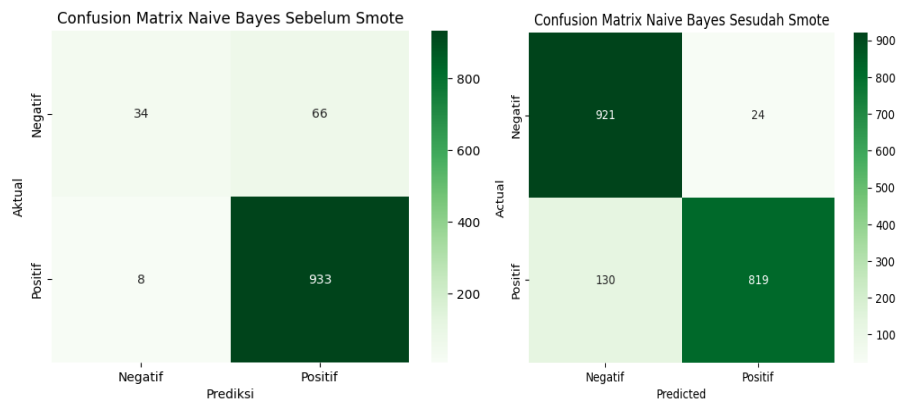
Tabel 8. Modelling Sesudah Menggunakan SMOTE

Modelling	Sentimen	Accuracy	Precision	Recal	F1-Score
Naïve Bayes	Negatif	0.91	0.88	0.97	0.92
	Positif		0.97	0.86	0.91
SVM	Negatif	0.99	0.99	1.00	1.00
	Positif		1.00	0.99	1.00
Decision Tre	Negatif	0.98	0.98	0.99	0.99
	Positif		0.99	0.98	0.99

Perbandingan antara Tabel 7 dan Tabel 8 menunjukkan kinerja algoritma klasifikasi *Naive Bayes*, *SVM (Support Vector Machine)*, dan *Decision Tree* sebelum dan sesudah penerapan SMOTE. Sebelum penerapan SMOTE, Naive Bayes memiliki akurasi 92%, dengan recal negatif rendah 34%, menunjukkan kesulitan dalam mendeteksi sentimen negatif. Untuk SVM dan Decision Tree menunjukkan performa yang baik, masing-masing dengan akurasi 98% dan recal negatif 82% untuk SVM. Setelah penerapan SMOTE Naive Bayes mengalami penurunan sedikit pada akurasi 91%, tetapi recall meningkat signifikan menjadi 97%. SVM mencapai kinerja optimal dengan akurasi 99% dan nilai maksimum untuk *precision* dan *recall* pada kategori negatif. Decision Tree tetap stabil dengan akurasi 98% dan recal negatif 99%. Secara keseluruhan, penerapan SMOTE berhasil meningkatkan kemampuan deteksi sentimen negatif, terutama pada model SVM.

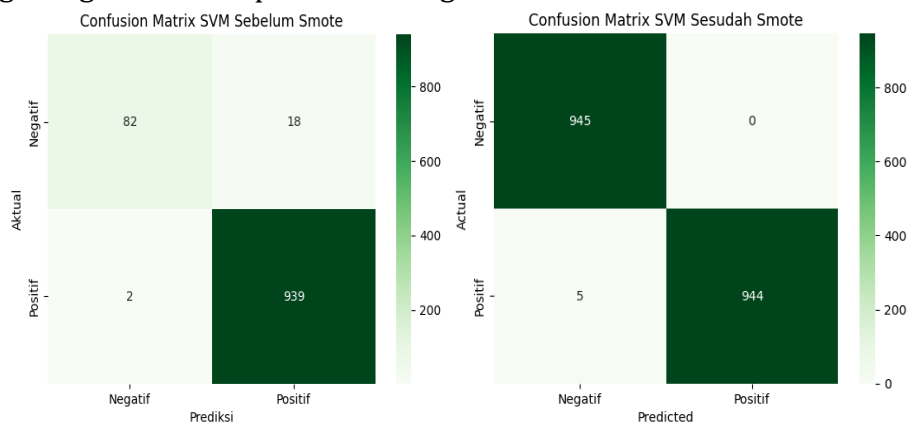
2. Confusion Matrix

Tujuan melakukan perbandingan antar model sebelum dan sesudah penerapan SMOTE sangat penting untuk mengevaluasi efektivitas teknik penyeimbangan data. *Confusion Matrix* memberikan gambaran yang jelas tentang kinerja model dalam hal prediksi kelas yang benar dan salah. Hasil confusion matrix sebelum dan sesudah SMOTE dapat dilihat pada gambar 4 dibawah.



Gambar 4. *Confusion Matrix* Naive Bayes Sebelum Dan Sesudah SMOTE

Gambar di atas menampilkan dua *Confusion Matrix* dari model Naive Bayes yang digunakan untuk klasifikasi sebelum dan sesudah penerapan SMOTE. Pada *Confusion Matrix* sebelum SMOTE, terdapat 34 prediksi *True Negative*, 66 *False Positive*, 8 *False Negative*, dan 933 *True Positive*. Sementara itu, pada *Confusion Matrix* setelah diterapkan SMOTE, terdapat 921 prediksi *True Negative*, 24 *False Positive*, 130 *False Negative*, dan 891 *True Positive*. Hasil ini menunjukkan bahwa model memiliki banyak prediksi yang benar pada kelas positif (933), tetapi mengalami kesalahan yang cukup tinggi pada kelas negatif, terutama sebelum penerapan SMOTE. Penerapan SMOTE tampak memberikan perbaikan yang signifikan dalam jumlah prediksi *True Negative* dan mengurangi kesalahan pada kelas negatif.



Gambar 5. *Confusion Matrix* SVM Sebelum Dan Sesudah SMOTE

Berdasarkan *Confusion Matrix* SVM Sebelum Dan Sesudah SMOTE, menunjukkan model SVM sebelum SMOTE menghasilkan 82 *True Negative*, dan 939 *True Positive*. Namun terdapat 18 *False Positive* dan 2 *False Negative*. Setelah penerapan SMOTE, performa model SVM meningkat secara signifikan dengan 945 *True Negative* dan 944 *True Positive*. Kesalahan klasifikasi juga berkurang drastis, dengan hanya 0 *False Positive* dan 5 *False Negative*.



3. Visualisasi Data

Wordcloud with stopwords



Berdasarkan penelitian mengenai sentimen kebijakan makan gratis di Indonesia yang menggunakan metode *Naive Bayes*, *SVM*, dan *Decision Tree* dengan jumlah tweet 5.205, serta memanfaatkan bahasa pemrograman Python di Google

Colab, dapat disimpulkan bahwa SVM memberikan akurasi tertinggi baik sebelum maupun sesudah penerapan SMOTE, dengan akurasi masing-masing 98% dan 99%. Setelah *SVM*, *Decision Tree* menunjukkan kinerja yang baik dengan akurasi 98%, sementara Naive Bayes memiliki akurasi terendah, yaitu 92% dan 91%. Hasil penelitian ini menunjukkan bahwa SVM merupakan metode paling efisien untuk menganalisis sentimen dalam konteks kebijakan ini. Untuk penelitian selanjutnya, disarankan agar peneliti mengeksplorasi penggunaan algoritma lain, seperti Random Forest atau teknik deep learning, yang berpotensi memberikan kinerja lebih baik dalam klasifikasi sentimen. Selain itu, penerapan teknik penyeimbangan data lain, seperti ADASYN atau Borderline-SMOTE, dapat dipertimbangkan untuk meningkatkan efektivitas deteksi sentimen negatif.

E. Ucapan Terima Kasih

Penulis mengucapkan terima kasih yang sebesar-besarnya kepada dosen pembimbing yang telah membimbing dan arahan selama proses penelitian ini. Selain itu, Penulis juga mengucapkan terima kasih kepada orang tua yang senantiasa mendukung dan memberikan dukungan.

F. Referensi

- [1] A. Sitanggang, Y. Umaidah, Y. Umaidah, R. I. Adam, and R. I. Adam, "Analisis Sentimen Masyarakat Terhadap Program Makan Siang Gratis Pada Media Sosial X Menggunakan Algoritma Naive Bayes," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, Aug. 2024.
- [2] D. P. Sitompul, Y. Sitorus, E. G. B. Sibuea, and S. D. Elsi, "Peran Media Sosial Dalam Mempengaruhi Perilaku Pemilih Pemula," *J. Law, Adm. Soc. Sci.*, vol. 4, no. 5, pp. 767–775, 2024.
- [3] T. Program, M. Siang, G. Tundo, and D. N. Rachmawati, "Implementasi Algoritma Naive Bayes untuk Analisis Sentimen," 2024.
- [4] C. Aulia Sari, A. Sukmawati, R. Putri Aprilli, P. Sarah Kayaningtias, and N. Yudistira, "Perbandingan Metode Naive Bayes, support vector machine, Dan Decision Tree Dalam Klasifikasi Konsumsi Obat," vol. 3, pp. 33–41, 2022.
- [5] R. T. Aldisa and P. Maulana, "Analisis Sentimen Opini Masyarakat Terhadap Vaksinasi Booster COVID-19 Dengan Perbandingan Metode Naive Bayes, Decision Tree dan SVM," *Build. Informatics, Technol. Sci.*, vol. 4, no. 1, pp. 106–109, Jun. 2022.
- [6] A. Supian, B. Tri Revaldo, N. Marhadi, L. Efrizoni, and R. Rahmadden, "Perbandingan Kinerja Naive Bayes Dan Svm Pada Analisis Sentimen Twitter Ibukota Nusantara," *J. Ilm. Inform.*, vol. 12, no. 01, pp. 15–21, 2024.
- [7] V. D. Yunanda and N. Hendrastuty, "Perbandingan Kernel Polynomial dan RBF Pada Algoritma SVM Untuk Analisis Sentimen Skincare di Indonesia," *J. Media Inform. Budidarma*, vol. 8, no. 2, p. 726, 2024.
- [8] P. Studi *et al.*, "Analisis Sentimen Review Aplikasi Mobile Legends Di Google Playstore Menggunakan Algoritma Support Vector Machine Dan TF-IDF," 2024.
- [9] A. Setiawan and R. R. Suryono, "Analisis Sentimen Ibu Kota Nusantara menggunakan Algoritma Support Vector Machine dan Naive Bayes," *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 183–192, 2024.

- [10] E. R. Lidinillah, T. Rohana, and A. R. Juwita, "Analisis sentimen twitter terhadap steam menggunakan algoritma logistic regression dan support vector machine," *TEKNOSAINS J. Sains, Teknol. dan Inform.*, vol. 10, no. 2, pp. 154–164, 2023.
- [11] S. Saepudin, S. Widiastuti, and C. Irawan, "Sentiment Analysis of Social Media Platform Reviews Using the Naïve Bayes Classifier Algorithm," *J. Sisfokom (Sistem Inf. dan Komputer)*, vol. 12, no. 2, pp. 236–243, 2023.
- [12] A. Ilham, "Analisis Sentimen Masyarakat Terhadap Kesehatan Mental Pada Twitter Menggunakan Algoritme K-Nearest Neighbor," *Pros. Semin. Nas. Mhs. Fak. ...*, vol. 2, no. September, pp. 539–547, 2023.
- [13] H. Hidayatullah, P. Purwantoro, and Y. Umaidah, "Penerapan Naïve Bayes Dengan Optimasi Information Gain Dan SMOTE Untuk Analisis Sentimen Pengguna Aplikasi Chatgpt," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 3, pp. 1546–1553, 2023.
- [14] I. D. A. N. I. Pada, "Perbandingan Kinerja Pre-Trained Tokopedia Seller Center," vol. 11, no. 2, pp. 13–20, 2024.