



ANALISIS SENTIMEN ULASAN APLIKASI GOJEK MENGGUNAKAN SUPPORT VECTOR MACHINE DAN RANDOM FOREST

Azka Bima Aditya¹, Syafri Samsudin², Winahyu Pandu Rizki³, Mahir Mahendra⁴, Arif Setiawan⁵

^{1,2,3,4} Teknik Informatika, Universitas Muria Kudus

⁵ Sistem Informasi, Universitas Muria Kudus
Kudus, Jawa Tengah, Indonesia 59327

202251131@std.umk.ac.id, 202251124@std.umk.ac.id, 202251079@std.umk.ac.id, 202251059@std.umk.ac.id,
arif.setiawan@umk.ac.id

Abstract

The rapid development of digital transportation, such as Gojek, requires a deep understanding of user satisfaction. This study analyzes the sentiment of Gojek application reviews to evaluate public opinion and compare the performance of the Support Vector Machine (SVM) and Random Forest models. A quantitative experimental method was applied to 30,055 user reviews for versions "4" and "5" from the Google Play Store. The data underwent comprehensive text preprocessing, automatic sentiment labeling using VADER enriched with an Indonesian lexicon, and TF-IDF feature extraction. The training data imbalance was addressed using SMOTE before the data was split for training and testing. The results show that user sentiment was dominated by positive (38.9%) and neutral (38.2%) categories. In the performance evaluation, the SVM model demonstrated superior performance with 96% accuracy and an F1-score of 0.96, outperforming the Random Forest model, which achieved 93% accuracy and an F1-score of 0.93. In conclusion, SVM is a more effective model for sentiment classification of Gojek reviews. Future research is recommended to refine the lexicon and implement aspect-based analysis to obtain more detailed insights.

Keywords: Gojek, Machine Learning, Random Forest, Sentiment Analysis, Support Vector Machine

Abstrak

Perkembangan pesat transportasi digital seperti Gojek menuntut pemahaman yang mendalam tentang kepuasan pengguna. Penelitian ini menganalisis sentimen dari ulasan aplikasi Gojek untuk mengevaluasi opini publik dan membandingkan performa model *Support Vector Machine* (SVM) dan *Random Forest*. Metode eksperimen kuantitatif diterapkan pada 30.055 ulasan pengguna versi "4" dan "5" dari Google Play Store. Data melalui pra-pemrosesan teks yang komprehensif, pelabelan sentimen otomatis menggunakan VADER yang diperkaya dengan leksikon Bahasa Indonesia, dan ekstraksi fitur TF-IDF. Ketidakseimbangan data latih diatasi dengan SMOTE sebelum data dibagi untuk pelatihan dan pengujian. Hasil menunjukkan bahwa sentimen pengguna didominasi oleh kategori positif (38,9%) dan netral (38,2%). Dalam evaluasi kinerja, model SVM terbukti sangat unggul dengan akurasi 96% dan *F1-score* 0,96, melampaui *Random Forest* yang memperoleh akurasi 93% dan *F1-score* 0,93. Kesimpulannya, SVM merupakan model yang lebih efektif untuk klasifikasi sentimen pada ulasan Gojek. Penelitian selanjutnya direkomendasikan untuk menyempurnakan leksikon dan menerapkan analisis berbasis aspek guna mendapatkan wawasan yang lebih mendalam.

Kata kunci: Analisis Sentimen, Gojek, Machine Learning, Random Forest, Support Vector Machine

1. PENDAHULUAN

Indonesia memiliki perkembangan moda transportasi yang cukup pesat. Perkembangan ini dibarengi dengan melonjaknya jumlah penduduk yang mendesak kebutuhan akan transportasi yang efisien, adaptif, serta dapat dijangkau oleh semua lapisan masyarakat [1]. Sebelum adanya aplikasi transportasi, masyarakat Indonesia sudah mengenal jenis transportasi tradisional seperti angkutan umum, taksi, serta ojek pangkalan [2]. Namun, sarana transportasi ini

sering dianggap tidak efisien karena rutenya yang terbatas, waktu tunggu, dan kurangnya tarif yang jelas. Kondisi ini melahirkan transportasi berbasis teknologi, salah satu pelopornya adalah Gojek [3]. Gojek mula-mula diluncurkan pada 2010 sebagai inovasi layanan ojek di era digital, yang kini telah berevolusi menjadi sebuah platform multi-layanan yang digunakan jutaan masyarakat Indonesia [1], [4].

Gojek sangat mengutamakan pandangan serta kepuasan pengguna mengingat aplikasi ini digunakan jutaan orang setiap hari [5]. *Review* di Google Play Store atau App Store tidak hanya berupa *feedback*, tetapi juga memiliki potensi untuk memengaruhi pengguna berikutnya [4]. *Feedback* dari pengguna memiliki kalimat-kalimat penting yang harus diperhatikan untuk pengembangan produk. Sayangnya, analisis ulasan ini tidak dilakukan secara menyeluruh untuk memetakan sentimen. Dengan kata lain, perusahaan gagal untuk mengetahui kebutuhan, keluhan, dan ekspektasi pelanggan secara menyeluruh. Dalam hal ini, analisis sentimen terhadap ulasan yang diberikan pengguna menjadi isu pokok yang mendasar dalam penelitian ini.

Berdasarkan penelitian [6], [7], analisis sentimen menjadi metode paling efisien untuk memahami opini publik. Perusahaan dapat dengan jelas mengetahui tren, keluhan, atau fitur servis yang paling disukai menggunakan teknik analisis opini. Pengolahan data yang masif seperti ini sangat bergantung pada otomatisasi model *machine learning* dalam klasifikasi sentimen. Dalam riset yang dilakukan [8] mengenai sentimen pengguna terhadap aplikasi Picsart, dipergunakan algoritma *Random Forest* untuk melakukan prediksi dengan tingkat akurasi mencapai 95,17%. Sementara penelitian [9] algoritma SVM mendapat akurasi 98% pada analisis sentimen terkait pemakaian aplikasi Shopee. Di sisi lain, studi tentang analisis sentimen aplikasi Gojek yang dilakukan, menunjukkan bahwa algoritma SVM memberikan akurasi terbaik sebesar 87% dibandingkan dengan algoritma *Random Forest* dan *Decision Tree* yang masing-masing mendapatkan skor 84,50% dan 78,25% [4]. Namun demikian, sebagian besar riset terdahulu belum mengkaji secara sistematis integrasi NLP dengan TF-IDF, belum menekankan pada penggunaan VADER yang diperkaya leksikon Bahasa Indonesia, serta jarang membandingkan performa SVM dan *Random Forest* dengan strategi penanganan ketidakseimbangan data melalui SMOTE. Oleh karena itu, penelitian ini menawarkan kontribusi baru dengan menggabungkan *preprocessing* komprehensif, teknik *feature extraction* modern, serta evaluasi model berbasis metrik yang lebih detail, sehingga memperkuat literatur analisis sentimen pada konteks transportasi digital di Indonesia

Machine learning memungkinkan komputer berinteraksi serta memahami bahasa alami manusia melalui lingkup NLP [10], [11]. Melalui analisis sentimen, NLP secara spesifik memproses, memahami, serta mengevaluasi kemudian memberikan pemahaman terhadap teks ulasan yang ditulis oleh para pengguna. Salah satu teknik yang digunakan dalam representasi akan hal ini TF-IDF (*Term Frequency-Inverse Document Frequency*), dimana sistem akan menghitung sejauh mana kata tersebut penting dalam dokumen dibandingkan dengan seluruh dokumen [12]. Dengan menggabungkan pendekatan TF-IDF dan NLP, data teks dapat lebih mudah diolah dan diubah menjadi bentuk numerik sehingga memudahkan pemrosesan model seperti *Random Forest* dan SVM dalam klasifikasi ulasan.

Penelitian ini menggunakan dua model *machine learning* populer dipakai untuk klasifikasi sentimen, *Support Vector Machine* (SVM) dan *Random Forest*. SVM dipilih karena algoritmanya sangat efektif untuk klasifikasi teks, efisien dalam hal komputasi, dan mudah diterapkan [13]. Sementara itu, *Random Forest* dipilih karena performanya bagus pada data dengan banyak fitur dan kemampuannya menangani *overfitting* [8]. Kedua model ini kemudian dibandingkan untuk mencari tahu mana yang paling akurat dalam mengklasifikasikan ulasan pengguna aplikasi Gojek.

Hasil dari penelitian ini diharapkan dapat memberikan gambaran jelas mengenai kepuasan dan keluhan pelanggan Gojek. Dengan memahami sentimen mayoritas pengguna, tim pengembang bisa membuat keputusan strategis untuk meningkatkan fitur, memperbaiki layanan, atau merancang kebijakan baru yang lebih sesuai dengan kebutuhan pelanggan. Lebih dari itu, analisis ini tak hanya berkontribusi pada pengembangan teknologi NLP dan *machine learning*, tapi juga menjadi alat penting bagi Gojek untuk meningkatkan kualitas layanan secara berkelanjutan.

2. METODE PENELITIAN

Penelitian ini menerapkan pendekatan eksperimen kuantitatif dengan menggunakan algoritma *Random Forest* dan *Support Vector Machine* untuk melakukan analisis sentimen. Pemilihan kedua metode ini didasarkan pada efektivitasnya dalam mendeteksi sentimen dari data berbasis teks, khususnya dalam mengevaluasi respons pengguna terhadap suatu aplikasi. Selain itu, pendekatan komputasional dimanfaatkan untuk menganalisis ulasan pengguna aplikasi Gojek yang diambil dari Google Play Store, dengan tujuan mengidentifikasi pola sentimen yang merefleksikan tingkat kepuasan dan kebahagiaan pengguna.

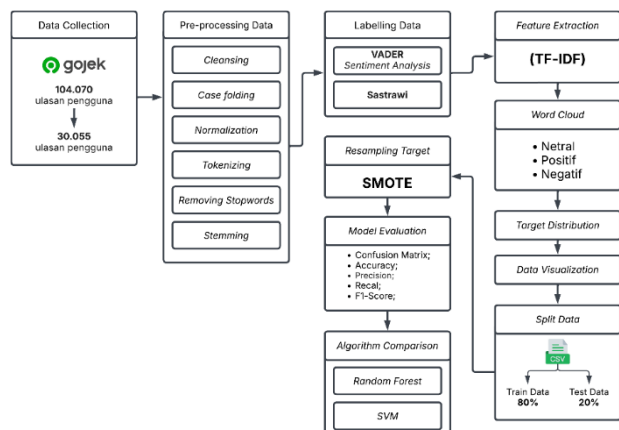
2.1 Metode Pengumpulan Data dan Analisis Data

Penelitian ini menggunakan data hasil *scraping* dari ulasan pada aplikasi Gojek yang terdapat dalam Google Play Store, data tersebut dimanfaatkan untuk mengkaji sentimen dalam ulasan pengguna aplikasi Gojek. Sumber data dalam penelitian ini berasal dari situs Kaggle melalui tautan <https://www.kaggle.com/datasets/dewanakretarta/gojek-playstore-reviews>. Informasi dalam data tersebut berupa ulasan, *rating*, versi aplikasi, jumlah *like*, dan waktu ulasan yang disimpan dengan format CSV.

Proses pengolahan data menggunakan teknik NLP (*Natural Language Processing*) dengan *preprocessing* teks yang meliputi *cleansing*, *tokenizing*, *removing stopwords*, dan *stemming*. Sentimen kemudian dilabeli menggunakan VADER, diikuti dengan ekstraksi fitur menggunakan TF-IDF dan visualisasi *Word Cloud*. Selanjutnya, distribusi sentimen dianalisis, data dipisahkan ke dalam dua kelompok, yaitu pelatihan dan pengujian, dan dilakukan *resampling* dengan SMOTE. Model dilatih menggunakan algoritma *Random Forest* dan SVM, kemudian dievaluasi berdasarkan *accuracy*, *precision*, *recall*, dan *F1-score*.

Seluruh proses ini diimplementasikan dengan bahasa pemrograman Python.

2.2 Tahapan penelitian



Gambar 1. Tahapan Penelitian

Beberapa tahapan penelitian ditunjukkan pada Gambar 1 sebagai berikut.

A. Data Collection

Pengumpulan data dari hasil *scraping* yang tersedia di situs Kaggle. Data mencakup informasi seperti isi ulasan (*content*), *rating* (*score*), versi aplikasi, jumlah *like*, dan waktu ulasan. Dari total 104.070 data ulasan yang tersedia peneliti menggunakan ulasan dari pengguna yang memakai versi aplikasi Gojek dengan awalan versi "4" dan "5". Setelah proses penyaringan versi tersebut, diperoleh sebanyak 30.055 data ulasan pengguna.

B. Preprocessing Data

Preprocessing mencakup serangkaian tahap awal dalam pengolahan data teks untuk membersihkan dan menyiapkannya sebelum dianalisis lebih lanjut [14]. Tahapan *preprocessing* dalam penelitian ini meliputi *cleansing*, *case folding*, *normalization*, *tokenizing*, *removing stopwords*, dan *stemming*. Setiap tahapan akan menjalani proses yang berbeda untuk memastikan data yang dihasilkan bersih dan terstruktur sesuai kebutuhan penelitian. Tahapan *preprocessing* dijelaskan sebagai berikut:

a) Cleansing

Guna meningkatkan kualitas data teks dan memastikan bahwa hanya informasi yang relevan yang dipertahankan, perlu dilakukan proses pembersihan dengan menghapus elemen-elemen yang tidak diperlukan, seperti emoji, angka, tautan (URL), serta tanda baca. Proses pembersihan ini bertujuan untuk mengurangi gangguan dalam analisis, sehingga model dapat lebih fokus mengutamakan kata-kata yang memiliki nilai informasi saat menginterpretasi sentimen atau isi teks. Dengan menghilangkan komponen-komponen tersebut, data menjadi lebih terstruktur dan siap untuk tahap pemrosesan selanjutnya [15].

b) Case Folding

Pada tahap *case folding*, seluruh teks ulasan dikonversi kedalam huruf kecil. Ini dilakukan untuk menghindari perbedaan antara kata yang sama yang ditulis dengan huruf kapital yang berbeda. Tujuan dari proses ini adalah untuk meningkatkan konsistensi, menyederhanakan analisis teks, dan mempertahankan struktur data supaya lebih teratur untuk tahap pemrosesan berikutnya.

c) Normalization

Normalization adalah proses mengganti kata tidak baku dan singkatan diganti dengan bentuk yang sesuai dengan kaidah bahasa yang berlaku. Tujuan dari langkah ini adalah untuk menyamakan bentuk kata, sehingga analisis teks menjadi lebih akurat dan lebih mewakili makna yang sebenarnya.

d) Tokenizing

Proses *tokenizing* dimulai dengan memecah teks menjadi elemen linguistik yang lebih kecil, seperti kata, frasa, atau simbol, sehingga dapat dianalisis lebih lanjut oleh algoritma. Setiap kalimat atau pernyataan diuraikan menjadi unit-unit dasar yang umumnya berbentuk kata, guna mempermudah tahapan pemrosesan selanjutnya.[16].

e) Removing Stopwords

Pada tahap *removing stopwords*, akan menghapus kata-kata umum seperti "yang", "dan", "oh", "di", serta kata-kata lain yang dianggap kurang memiliki bobot informasi serta berkontribusi penting terhadap analisis. Proses ini bertujuan untuk menyaring informasi yang kurang relevan, agar data yang diolah menjadi lebih efisien serta terpusat pada kata-kata yang mengandung makna penting.

f) Stemming

Tahap *stemming* adalah proses mengonversi kata-kata ke dalam bentuk dasar atau akar katanya. Contohnya, kata "orderan" diubah menjadi "order". Dengan cara ini, dapat mengurangi variasi bentuk kata dalam teks. Hal ini membuat analisis menjadi lebih mudah, dan model dapat lebih cepat mengenali pola atau sentimen, karena fokusnya ada pada inti makna kata tanpa terganggu oleh imbuhan atau perubahan bentuk lainnya [17].

g) Labeling

Pelabelan sentimen pada penelitian ini dilakukan menggunakan metode *VADER Sentiment Analysis*. *VADER* merupakan *library* analisis sentimen populer dalam pemrosesan bahasa alami, yang memang dirancang khusus untuk mendeteksi sentimen positif, negatif, atau netral dalam teks [18]. Untuk meningkatkan akurasi, *VADER* dikombinasikan dengan kamus leksikon Bahasa Indonesia, yang membantu memahami nuansa bahasa lokal yang mungkin tidak terdeteksi oleh kamus standar. Setelah analisis, sentimen dikategorikan menjadi tiga kelas utama: positif, netral, dan negatif. Penentuan kategori ini

didasarkan pada nilai *compound* yang dihasilkan VADER, yang merepresentasikan skor gabungan sentimen keseluruhan teks.

C. Feature Extraction (TF-IDF)

Setelah pra-pemrosesan, data teks diubah menjadi fitur numerik dengan metode TF-IDF (*Term Frequency – Inverse Document Frequency*). Metode ini berfungsi untuk menunjukkan sejauh mana suatu kata memiliki bobot informasi dalam dokumen tertentu dibandingkan seluruh dokumen yang dianalisis [7].

D. Word Cloud

Word Cloud adalah metode yang digunakan dalam penelitian ini dalam *text mining* untuk menganalisis data teks, dan digunakan untuk memvisualisasikan frekuensi kata dominan di setiap kelas sentimen. Dalam visualisasi tersebut, kata yang berukuran besar menandakan frekuensi kemunculan yang tinggi, begitu pun sebaliknya [19]. Ini membantu mengidentifikasi pola kata yang dominan muncul pada masing-masing kelompok sentimen.

E. Target Distribution

Distribusi target merupakan langkah penting dalam eksplorasi data yang membantu memahami bagaimana data tersebar di setiap kategori dalam variabel target. Dalam penelitian ini, variabel target yang dibahas adalah sentimen, yang terbagi menjadi tiga kategori yaitu positif, netral, dan negatif. Analisis ini bertujuan untuk memahami komposisi *dataset* secara kuantitatif dan mengidentifikasi potensi adanya ketidakseimbangan kelas. Ketidakseimbangan kelas terjadi ketika jumlah sampel di setiap kategori target tidak terdistribusi secara merata. Situasi ini dapat berdampak besar pada proses pelatihan dan evaluasi model, karena model cenderung lebih condong kepada kelas mayoritas, yaitu kelas yang memiliki jumlah sampel terbanyak.

F. Visualisasi

Visualisasi distribusi target dilakukan untuk memahami bagaimana data tersebar di setiap kategori sentimen yang berfungsi sebagai variabel target. Visualisasi dengan menggunakan *bar chart* dan *funnel chart* untuk memberikan gambaran yang lebih jelas tentang proporsi masing-masing kelas. *Bar chart* menampilkan jumlah data di setiap kategori sentimen dalam bentuk *bar chart*, sementara *funnel chart* menggambarkan urutan dan proporsi data secara hierarkis. Kedua jenis visualisasi ini sangat membantu peneliti dalam mengidentifikasi potensi ketidakseimbangan kelas yang bisa memengaruhi analisis lebih lanjut, terutama dalam proses pemodelan.

G. Split Data

Dalam pembagian *dataset*, data dipisahkan menjadi dua bagian, dengan 80% dialokasikan untuk pelatihan dan 20% untuk pengujian model. Pembagian ini penting untuk

mengevaluasi seberapa baik performa model saat dihadapkan pada data yang sama sekali baru.

H. Resampling Target

Resampling yang digunakan pada penelitian ini adalah SMOTE (*Synthetic Minority Over-sampling Technique*). Metode ini diterapkan untuk mengatasi ketidakseimbangan data yang ada, SMOTE bekerja dengan menghasilkan sampel data sintetis untuk kelas minoritas [20]. Tujuannya adalah menyeimbangkan jumlah data di setiap kelas, sehingga model *machine learning* tidak bias dan dapat belajar mengenali pola dari kelas minoritas dengan lebih baik.

I. Model Evaluation

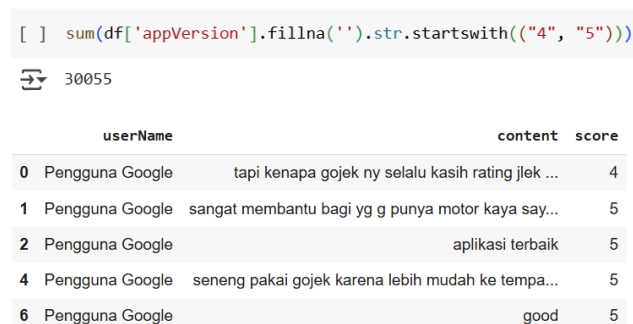
Dalam studi ini, proses pemodelan dilakukan dengan menerapkan dua algoritma klasifikasi, yaitu *Random Forest* dan *Support Vector Machine* (SVM). Setelah model dilatih, evaluasi performa dilakukan dengan mengukur sejumlah metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score* untuk menilai tingkat efektivitas model dalam melakukan klasifikasi sentimen.

Selain itu, bagian ini juga memaparkan secara mendetail setiap tahapan dalam pelaksanaan penelitian, lengkap dengan hasil yang diperoleh pada masing-masing tahap. Penyajian tersebut bertujuan untuk memberikan gambaran komprehensif mengenai alur proses penelitian secara sistematis.

3. HASIL DAN PEMBAHASAN

3.1 Data Collection

Data ulasan pengguna Gojek yang telah dianalisis berjumlah 104.070, yang diambil dari Kaggle. Data ini mencakup berbagai informasi seperti isi ulasan, *rating*, versi aplikasi, jumlah *like*, dan waktu ulasan. Penelitian ini difokuskan pada ulasan dari pengguna Gojek versi "4" dan "5", karena kedua versi ini dianggap paling mewakili pengalaman pengguna saat ini. Setelah melakukan penyaringan, 30.055 data ulasan yang relevan berhasil dikumpulkan, dan jumlah ini dianggap cukup untuk analisis sentimen yang lebih mendalam serta pelatihan model yang dapat dilihat pada Gambar 2.



```
[ ] sum(df['appVersion'].fillna('').str.startswith(("4", "5")))
```

30055

	userName	content	score
0	Pengguna Google	tapi kenapa gojek ny selalu kasih rating jlek ...	4
1	Pengguna Google	sangat membantu bagi yg g punya motor kaya say...	5
2	Pengguna Google	aplikasi terbaik	5
4	Pengguna Google	seneng pakai gojek karena lebih mudah ke tempa...	5
6	Pengguna Google	good	5

Gambar 2. Data Collection

3.2 Preprocessing Data

Langkah pertama dalam penerapan *text mining* dimulai dengan pra-pemrosesan. Di sini, informasi yang dianggap relevan dari masing-masing dokumen diseleksi untuk analisis lebih lanjut. Tujuan utamanya adalah membersihkan elemen-elemen yang kurang penting agar data yang dihasilkan berkualitas tinggi dan siap untuk dianalisis. Dalam penelitian ini, tahap pra-pemrosesan mencakup beberapa langkah penting, seperti *cleansing*, *case folding*, normalisasi, tokenisasi, penghapusan *stopword*, *stemming*, dan *labeling*. Semua langkah ini bertujuan untuk menormalkan data dengan menghilangkan kata-kata yang kurang bermakna, sehingga hasil pemrosesan menjadi lebih optimal dan dapat disajikan dengan lebih terstruktur. Tahap *preprocessing* data ditunjukkan pada Tabel 1.

Tabel 1. Preprocessing Data

Tahap	Deskripsi	Ulasan
<i>Cleansing</i>	Menghapus karakter khusus, angka, tanda baca, dan spasi berlebih	Driver nya sekarang pada malas malas bgt orderan lama pick up nya apa lagi ambil fitur hemat sering teman kejadian di mintai tip tolong segera di tegur
<i>Case Folding</i>	Mengubah seluruh huruf menjadi huruf kecil	driver nya sekarang pada malas malas bgt orderan lama pick up nya apa lagi ambil fitur hemat sering teman kejadian di mintai tip tolong segera di tegur
<i>Normalization</i>	Mengganti kata tidak baku atau slang menjadi bentuk baku	drivernya sekarang pada malas malas banget orderan lama pick upnya apalagi ambil fitur hemat sering teman kejadian dimintai tip tolong segera ditegur
<i>Tokenizing</i>	Memisahkan kalimat menjadi kata-kata	[drivernya, sekarang, pada, malas, malas, banget, orderan, lama, pick, upnya, apalagi, ambil, fitur, hemat, sering, teman, kejadian, dimintai, tip, tolong, segera, ditegur]
<i>Removing Stopwords</i>	Menghapus <i>stopwords</i> atau kata-kata umum yang tidak penting.	['drivernya', 'malas', 'malas', 'orderan', 'lama', 'pickupnya', 'fitur', 'hemat', 'tip', 'ditegur']
<i>Stemming</i>	Mengubah kata ke bentuk dasar menggunakan pustaka Sastrawi	['driver', 'malas', 'malas', 'order', 'lama', 'pickup', 'fitur', 'hemat', 'tip', 'tegur']

3.3 Labeling

Pelabelan sentimen pada penelitian ini menggunakan metode *VADER Sentiment Analysis*, yang diperkaya dengan kamus leksikon Bahasa Indonesia. Metode ini berhasil mengklasifikasikan 21.698 data ulasan ke dalam

tiga kategori positif, netral, dan negatif. Distribusi hasilnya adalah 38,9% positif, 38,2% netral dan 22,9% negatif disajikan pada Tabel 2. Ketimpangan distribusi ini, terutama dominasi sentimen netral, menjadi alasan utama diterapkannya metode *resampling* pada tahap pelatihan model.

Tabel 2. Distribusi Sentimen

Positif	Netral	Negatif
38,9% (8430)	38,2% (8298)	22,9% (4970)

3.4 Feature TF-IDF

Transformasi data teks menjadi fitur numerik setelah melalui tahap pra-pemrosesan dilakukan dengan metode TF-IDF (*Term Frequency – Inverse Document Frequency*). Pemilihan TF-IDF ini sangat penting dan krusial dalam menyoroti seberapa signifikan sebuah kata dalam dokumen tertentu dibandingkan dengan seluruh kumpulan dokumen [7]. Dengan representasi numerik ini, model klasifikasi mendapatkan *input* yang memungkinkan 'pemahaman' terhadap konteks dan makna setiap kata, yang pada akhirnya membantu meningkatkan akurasi dalam mengidentifikasi sentimen dari ulasan pengguna Gojek. Implementasi TF-IDF dapat dilihat pada Gambar 3.

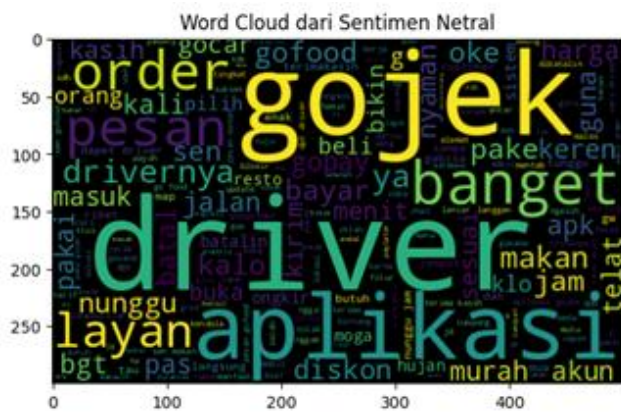
```
# TF-IDF
from sklearn.feature_extraction.text import TfidfVectorizer

tfidf_vectorizer = TfidfVectorizer()
X_tfidf = tfidf_vectorizer.fit_transform(df['content'])
```

Gambar 3. TF-IDF

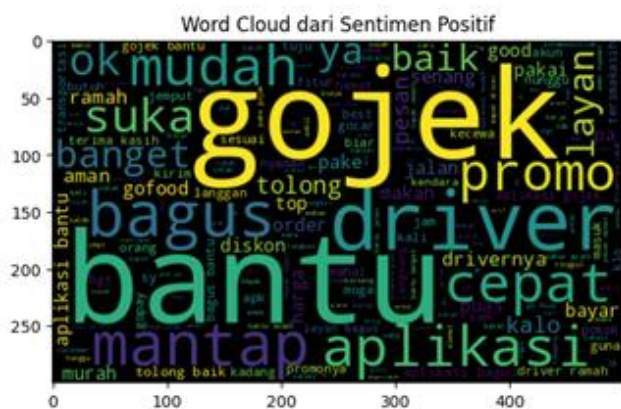
3.5 Word Cloud

Visualisasi *word cloud* digunakan untuk mengidentifikasi kata-kata yang sering muncul dalam masing-masing kategori sentimen ulasan pengguna Gojek. Visualisasi ini dibedakan berdasarkan klasifikasi sentimen yang dilakukan, yaitu positif, netral, dan negatif. Ulasan sentimen netral pada aplikasi Gojek cenderung berfokus pada deskripsi layanan umum, kerap menyebut "gojek", "aplikasi", dan "driver" tanpa muatan emosional kuat. Kata-kata seperti "cepat", "mudah", "layanan", dan "pesan" juga menonjol, merefleksikan pengalaman pengguna terkait kepraktisan. Namun, konteks netralnya menunjukkan penggunaan kata-kata ini lebih sebagai penjelasan objektif, bukan ekspresi sentimen positif atau negatif ditampilkan pada Gambar 4.



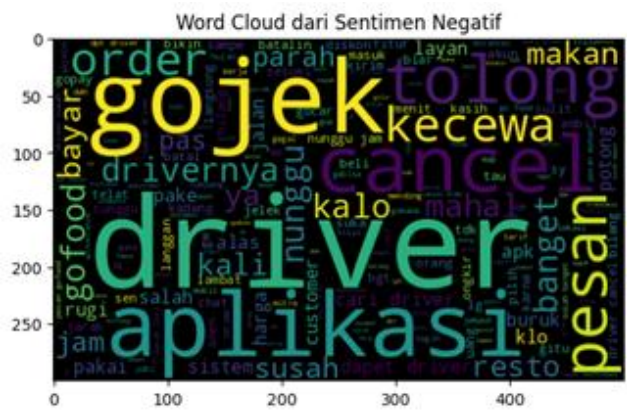
Gambar 4. Visualisasi Word Cloud Netral

Sementara analisis sentimen positif menunjukkan kecenderungan pengguna Gojek menggunakan kata-kata seperti "bagus", "baik", "mantap", "promo", dan "diskon". Ini merefleksikan kepuasan mereka terhadap layanan, kinerja aplikasi, dan penawaran menarik. Kemunculan dominan kata "good", "driver", "puas", dan "layanan" juga mengindikasikan apresiasi terhadap mitra pengemudi serta kenyamanan aplikasi, menyoroti umpan balik positif pengguna terhadap keseluruhan pengalaman Gojek divisualisasikan pada Gambar 5.



Gambar 5. Visualisasi Word Cloud Positif

Disisi lain, ulasan negatif pada aplikasi Gojek didominasi oleh kata-kata seperti "driver", "susah", "kecewa", dan "parah", merefleksikan ketidakpuasan terhadap kinerja pengemudi dan layanan umum. Kata "cancel", "mahal", dan "error" juga menonjol, menandakan keluhan pengguna terkait pembatalan pesanan, biaya layanan, dan gangguan teknis aplikasi dapat dilihat pada Gambar 6.



Gambar 6. Visualisasi Word Cloud Negatif

3.6 Target Distribution

Pada tahap analisis distribusi target, dilakukan perhitungan jumlah data di setiap kategori sentimen untuk memahami proporsi dalam *dataset*. Dari Gambar 7, kita bisa melihat bahwa *dataset* ini terdiri dari 8.430 data dengan sentimen positif, 8.298 data netral, dan 4.970 data negatif. Distribusi ini menunjukkan adanya ketidakseimbangan kelas, karena jumlah data dengan sentimen negatif jauh lebih sedikit dibandingkan dengan dua kategori lainnya. Ketidakseimbangan ini bisa memengaruhi performa model klasifikasi yang kita kembangkan, karena model cenderung lebih terlatih pada kelas mayoritas, yaitu sentimen positif dan netral, sehingga ada risiko kurang akurat dalam mengenali sentimen negatif.

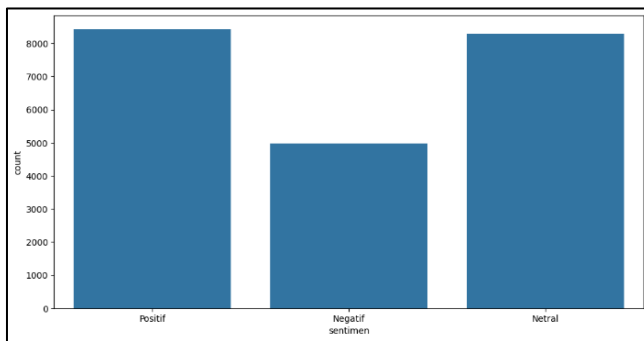
```
[ ] temp = df.groupby('sentimen').count()['content'].reset_index().sort_values(by='content', ascending=False)
temp.style.background_gradient(cmap='inferno_r')
```

sentimen	content
2	Positif 8430
1	Netral 8298
0	Negatif 4970

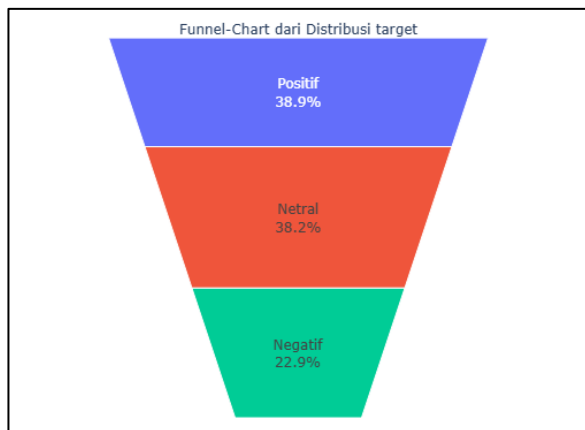
Gambar 7. Target Distribution

3.7 Visualisasi

Tahapan visualisasi bertujuan untuk menggambarkan sebaran data target secara grafis. Pada Gambar 8 (*Bar Chart*), tampak bahwa jumlah data dengan sentimen 'Negatif' jauh lebih rendah dibandingkan dengan kategori 'Positif' dan 'Netral'. Hal ini diperjelas melalui Gambar 9 (*Funnel Chart*), yang menyajikan proporsi persentase masing-masing sentimen, di mana sentimen negatif hanya mencakup 22,9% dari total data, sementara sentimen positif dan netral masing-masing mencapai 38,9% dan 38,2%. Kedua visualisasi ini secara efektif mengonfirmasi adanya ketimpangan distribusi kelas yang cukup mencolok dalam *dataset* yang digunakan.



Gambar 8. Visualisasi Bar Chart



Gambar 9. Visualisasi Funnel Chart

3.8 Split Data

Pada tahap ini, untuk keperluan pelatihan dan evaluasi model, *dataset* dipisahkan menjadi data latih dan data uji menggunakan metode *train_test_split* dari *Scikit-learn* untuk evaluasi model. Sebagaimana ditampilkan pada Gambar 10, proses pembagian data dilakukan dengan proporsi 80% sebagai data latih dan 20% sebagai data uji, yang diatur menggunakan parameter *test_size=0.2*. Selain itu, penggunaan *random_state=42* bertujuan untuk menjaga konsistensi hasil pemisahan data pada setiap proses eksekusi. Proses ini menghasilkan 17.358 sampel data untuk pelatihan (*X_train*) dan 4.340 sampel data untuk pengujian (*X_test*), yang keduanya siap untuk tahap pemodelan dan evaluasi selanjutnya.

```
[ ] # splitting
from sklearn.model_selection import train_test_split

X_train, X_test, y_train, y_test = train_test_split(X_tfidf, df['sentimen'], test_size=0.2, random_state=42)
X_train.shape, X_test.shape

((17358, 11677), (4340, 11677))
```

Gambar 10. Split Data

3.9 Resampling Target

Penanganan masalah ketidakseimbangan kelas yang telah diidentifikasi sebelumnya dilakukan dengan menerapkan teknik *resampling* pada data latih menggunakan metode SMOTE (*Synthetic Minority Over-sampling Technique*), seperti yang terlihat pada Gambar 11. Metode ini berfungsi dengan menciptakan sampel sintesis untuk kelas minoritas

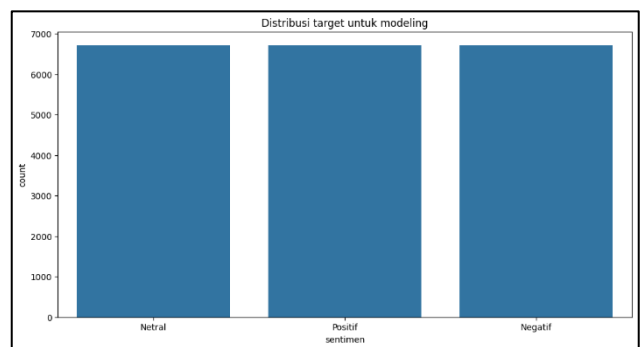
('Negatif') hingga jumlahnya sebanding dengan kelas mayoritas. Hasil dari proses ini, yang divisualisasikan pada Gambar 12, menunjukkan distribusi data latih yang kini seimbang sempurna, setiap kelas sentimen (Positif, Negatif, dan Netral) memiliki jumlah sampel yang sama. Langkah ini sangat krusial untuk memastikan bahwa model yang akan dilatih tidak mengalami bias dan dapat mengenali semua kelas dengan baik.

```
from imblearn.over_sampling import SMOTE

smote = SMOTE(random_state=42)
X_train, y_train = smote.fit_resample(X_train, y_train)

plt.figure(figsize=(12, 6))
sns.countplot(x=y_train)
plt.title('Distribusi target untuk modeling')
plt.show()
```

Gambar 11. Resampling Target dengan SMOTE



Gambar 12. Hasil Resampling Target

3.10 Model Evaluation

Penelitian ini melakukan perbandingan terhadap performa dua algoritma klasifikasi, yaitu *Support Vector Machine* (SVM) dan *Random Forest*. *Dataset* yang digunakan dalam penelitian ini dibagi ke dalam dua *subset*, 80% dialokasikan sebagai data latih (*training set*) dan 20% sebagai data uji (*testing set*). Pembagian ini dilakukan untuk memastikan bahwa model dapat dievaluasi secara objektif terhadap data yang belum pernah "dilihat" sebelumnya. Mengingat adanya ketidakseimbangan pada distribusi data sentimen, teknik *resampling* menggunakan SMOTE diterapkan pada data latih untuk menyeimbangkan jumlah sampel di setiap kelas sentimen.

Setelah proses pelatihan, model *Support Vector Machine* (SVM) pada Gambar 13 menunjukkan performa terbaik dengan akurasi mencapai 96%. Untuk metrik evaluasi lainnya, nilai *precision*, *recall*, dan *F1-score* rata-rata tertimbang masing-masing adalah 0.96, 0.96, dan 0.96. Rincian *F1-score* untuk setiap kelas adalah 0,94 untuk sentimen negatif, 0,97 untuk netral, dan 0,97 untuk positif. Sementara, model *Random Forest* pada Gambar 14, berhasil meraih akurasi sebesar 93%, dengan nilai *precision*, *recall*, dan *F1-score* rata-rata tertimbang masing-masing 0.94,

0.93, dan 0.93. Rincian *F1-score* menunjukkan 0,88 untuk sentimen negatif, 0,95 untuk netral, dan 0,94 untuk positif.

Classification Report for SVM (Tuned):				
	precision	recall	f1-score	support
Negatif	0.95	0.94	0.94	956
Netral	0.95	0.98	0.97	1673
Positif	0.98	0.95	0.97	1711
accuracy			0.96	4340
macro avg	0.96	0.96	0.96	4340
weighted avg	0.96	0.96	0.96	4340

Gambar 13. Evaluasi Model *Support Vector Machine (SVM)*

Classification Report for Random Forest (Tuned):				
	precision	recall	f1-score	support
Negatif	0.81	0.96	0.88	956
Netral	0.96	0.94	0.95	1673
Positif	0.99	0.91	0.94	1711
accuracy			0.93	4340
macro avg	0.92	0.94	0.93	4340
weighted avg	0.94	0.93	0.93	4340

Gambar 14. Evaluasi Model *Random Forest*

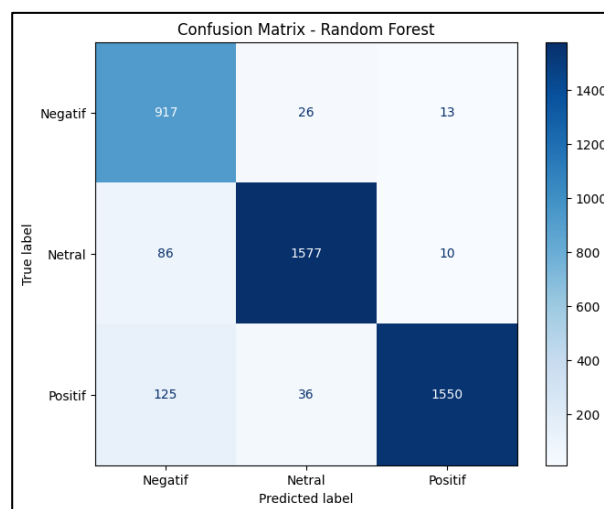
Hasil penelitian ini sejalan dengan temuan sebelumnya yang menunjukkan bahwa algoritma SVM sangat efektif dalam mengklasifikasikan sentimen ulasan aplikasi Gojek, dengan tingkat akurasi yang cukup tinggi [4]. Penelitian serupa juga menegaskan posisi SVM sebagai algoritma unggulan, mencatat akurasi hingga 98% dalam analisis sentimen ulasan aplikasi Shopee [9]. Di sisi lain, penelitian yang menggunakan *Random Forest* pada aplikasi Picsart menunjukkan akurasi yang tinggi, sekitar 95,17%. Ini membuktikan bahwa kedua algoritma memiliki potensi besar untuk tugas klasifikasi sentimen.

Perbedaan nilai akurasi dalam penelitian ini terutama dipengaruhi oleh penerapan teknik *preprocessing* yang lebih menyeluruh dan penggunaan metode SMOTE untuk mengatasi ketidakseimbangan kelas data. Ini memberikan peningkatan performa model, terutama pada *dataset* yang besar dan beragam, seperti ulasan aplikasi Gojek versi 4 dan 5 yang telah dianalisis. Penelitian sebelumnya cenderung hanya menggunakan satu algoritma atau tidak mengintegrasikan kombinasi NLP dan ekstraksi fitur TF-IDF secara menyeluruh, sehingga kurang memberikan gambaran lengkap tentang performa masing-masing algoritma dalam konteks tersebut.

Dengan pendekatan yang lebih sistematis dan menyeluruh ini, studi ini berkontribusi pada pengembangan teori dan praktik analisis sentimen, serta mengisi kesenjangan dalam literatur yang berkaitan dengan perbandingan performa SVM dan *Random Forest* secara komprehensif pada layanan transportasi *online* di Indonesia.

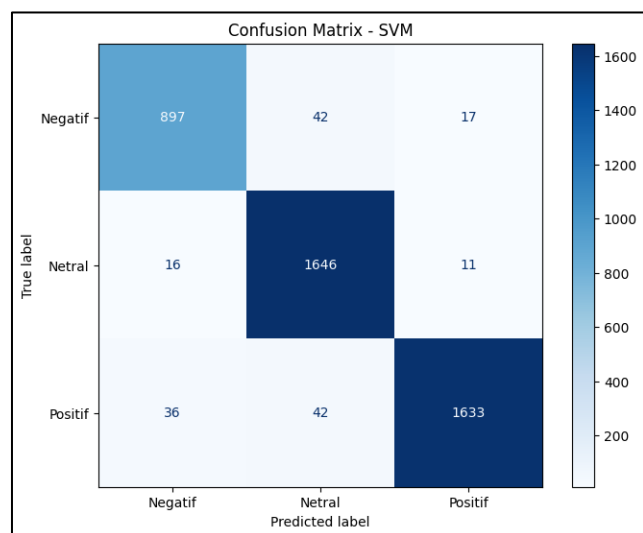
3.11 Confusion Matrix

Model ini tercatat berhasil mengklasifikasikan 1.577 data netral dengan benar, serta 1.550 data positif secara tepat. Meskipun demikian, model ini menghadapi tantangan signifikan pada kelas negatif, terbukti dari 86 data netral dan 125 data positif yang keliru diklasifikasikan sebagai negatif. Fenomena ini menjelaskan mengapa nilai *F1-score* untuk kelas negatif berada di angka terendah (0,88), karena model cenderung lebih sering keliru dalam membedakan sentimen negatif dari yang lainnya. Keseluruhan performa ini, model mampu melakukan klasifikasi terhadap mayoritas data secara akurat terutama untuk kelas netral dan positif, dapat dilihat pada *confusion matrix* di Gambar 15.



Gambar 15. *Confusion Matrix Random Forest*

Model SVM berhasil mengklasifikasikan 1.646 data netral dan 1.633 data positif dengan tepat, serta hanya melakukan sedikit kesalahan di masing-masing kelas. Walaupun kelas negatif sedikit lebih menantang bagi model ini, dengan 42 data negatif yang keliru diklasifikasikan sebagai netral, hasil keseluruhannya tetap menunjukkan konsistensi yang tinggi. Ketepatan model dalam membedakan sentimen netral dan positif sangat mencolok, terbukti dari nilai *F1-score* yang tinggi untuk kedua kelas tersebut, yakni 0,97, menjadikan SVM sebagai model dengan performa terbaik dalam penelitian ini. Seluruh hasil klasifikasi performa yang lebih akurat di semua kelas ditunjukkan pada Gambar 16.



Gambar 16. Confusion Matrix SVM

4. KESIMPULAN

Dari rangkaian proses penelitian yang telah dilaksanakan, dapat ditarik kesimpulan bahwa analisis sentimen ulasan pengguna aplikasi Gojek versi awalan “4” dan “5” dari Google Play Store sejumlah 30.055 data berbahasa Indonesia berhasil dilaksanakan. Melalui proses *preprocessing* yang komprehensif, data teks berhasil disiapkan untuk diolah oleh *machine learning*. Pelabelan sentimen dengan VADER, didukung leksikon Bahasa Indonesia, sukses mengklasifikasikan ulasan menjadi positif, netral, dan negatif, dengan mayoritas ulasan didominasi oleh sentimen positif. Ekstraksi fitur menggunakan TF-IDF efektif mengubah teks menjadi representasi numerik untuk *input* model klasifikasi. Dari perbandingan dua algoritma, *Support Vector Machine* terbukti unggul dengan akurasi 96%, mengalahkan *Random Forest* yang mencapai 93%, menunjukkan efektivitas SVM dalam mengklasifikasikan sentimen ulasan Gojek pada studi ini.

Berdasarkan temuan dan keterbatasan penelitian ini, beberapa saran untuk riset di masa mendatang mencakup:

1. Penyempurnaan Leksikon: Dengan memanfaatkan leksikon Bahasa Indonesia yang lebih lengkap dan sesuai dengan konteks ulasan aplikasi, akurasi dalam pelabelan sentimen otomatis dapat ditingkatkan.
2. Eksplorasi Algoritma: Melakukan pengujian dengan algoritma *machine learning* lain seperti *XGBoost*, *Logistic Regression*, atau metode *deep learning* untuk memperoleh perbandingan kinerja model yang lebih beragam.
3. Analisis Berbasis Aspek: Menambahkan fitur analisis sentimen berbasis aspek (*aspect-based sentiment analysis*) untuk memberikan wawasan lebih mendalam mengenai elemen layanan Gojek yang paling disukai atau dikeluhkan pengguna.

DAFTAR PUSTAKA

- [1] B. Astuti and M. R. Daud, “Kepastian Hukum Pengaturan Transportasi Online,” *Al-Qisth Law Rev.*, vol. 6, no. 2, p. 205, Feb. 2023, doi: 10.24853/al-qisth.6.2.205-244.
- [2] V. Rhesy Modompit, J. Bintang Kalangi, and J. I. Sumual, “Analisis Permintaan Transportasi Gojek Online di Kota Manado,” *J. Berk. Ilm. Efisiensi*, vol. 20, no. 3, pp. 1–12, 2020.
- [3] P. Fakhriyah, “Pengaruh Layanan Transportasi Online (Gojek) Terhadap Perluasan Lapangan Kerja Bagi Masyarakat Di Kota Cimahi,” *Comm-Edu (Community Educ. Journal)*, vol. 3, no. 1, p. 34, Jan. 2020, doi: 10.22460/comm-edu.v3i1.3719.
- [4] G. Kanugrahan, V. Hafizh, C. Putra, and Y. Ramdhani, “Analisis Sentimen Aplikasi Gojek Menggunakan SVM, Random Forest dan Decision Tree,” *J. Infortech*, vol. 6, no. 2, pp. 171–178, 2024.
- [5] Harun Raudhatul Na’im and Wiyadi, “Analisis Pengaruh Kualitas Pelayanan, Harga, Citra Merek Dan Word Of Mouth Terhadap Kepuasan Pelanggan Transportasi Online (Gojek),” *J. LENTERA BISNIS*, vol. 13, no. 3, pp. 1789–1805, Sep. 2024, doi: 10.34127/jrlab.v13i3.1223.
- [6] A. Khusnul Khotimah, “Analisis Sentimen Terhadap Kualitas Pelayanan,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 3, pp. 3044–3048, May 2024, doi: 10.36040/jati.v8i3.9520.
- [7] S. Azzahra, Z. Kusuma, D. E. Ratnawati, and N. Y. Setiawan, “Analisis Sentimen Pengguna Sosial Media Twitter / X Terhadap Acara Clash Of Champions Menggunakan Metode Multinomial Naïve Bayes,” *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 9, no. 3, pp. 1–10, 2025.
- [8] S. Farkhatul Jannah, R. Astuti, and F. Muhamad Basysyar, “Implementasi Algoritma Random Forest Pada Aplikasi Picsart Berdasarkan Respon Pengguna,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 274–283, Feb. 2024, doi: 10.36040/jati.v8i1.8329.
- [9] I. S. K. Idris, Y. A. Mustofa, and I. A. Salihi, “Analisis Sentimen Terhadap Penggunaan Aplikasi Shopee Menggunakan Algoritma Support Vector Machine (SVM),” *Jambura J. Electr. Electron. Eng.*, vol. 5, no. 1, pp. 32–35, Jan. 2023, doi: 10.37905/jjee.v5i1.16830.
- [10] I. Huda, “Implementasi Natural Language Processing (Nlp) Untuk Aplikasi Pencarian Lokasi,” *J. Nas. Teknol. Terap.*, vol. 3, no. 2, p. 15, Oct. 2021, doi: 10.22146/jntt.35036.
- [11] F. N. Zaman, M. A. Fadhilah, M. A. Ulinuha, and K. Umam, “Menganalisis Respons Netizen Twitter Terhadap Program Makan Siang Gratis Menerapkan Nlp Metode Naïve Bayes,” *Just IT J.*

- Sist. Informasi, Teknol. Inf. dan Komput.*, vol. 14, no. 3, pp. 150–233, 2024, [Online]. Available: <https://jurnal.umj.ac.id/index.php/just-it/index>
- [12] F. D. Adhiatma and A. Qoiriah, “Penerapan Metode TF-IDF dan Deep Neural Network untuk Analisa Sentimen pada Data Ulasan Hotel,” *J. Informatics Comput. Sci.*, vol. xx, pp. 183–193, Nov. 2022, doi: 10.26740/jinacs.v4n02.p183-193.
- [13] M. R. Adrian, M. P. Putra, M. H. Rafialdy, and N. A. Rakhmawati, “Perbandingan Metode Klasifikasi Random Forest dan SVM Pada Analisis Sentimen PSBB,” *J. Inform. Upgris*, vol. 7, no. 1, pp. 36–40, Jun. 2021, doi: 10.26877/jiu.v7i1.7099.
- [14] U. Khairani, V. Mutiawani, and H. Ahmadian, “Pengaruh Tahapan Preprocessing Terhadap Model Indobert Dan Indobertweet Untuk Mendeteksi Emosi Pada Komentar Akun Berita Instagram,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 11, no. 4, pp. 887–894, Aug. 2024, doi: 10.25126/jtiik.1148315.
- [15] L. K. Sukiman, A. R. D. Saribu, and A. Wijaya, “Analisis Sentimen Ulasan Aplikasi LinkedIn Dalam Google Play Store Dengan Model Naïve Bayes,” *Djtechno J. Teknol. Inf.*, vol. 4, no. 2, pp. 374–385, 2023.
- [16] N. Z. Putri, M. Martanto, A. R. Dikananda, and A. Rifa’i, “Analisis Sentimen Aplikasi SeaBank dengan Algoritma Naive Bayes untuk Optimalisasi Pelayanan,” *J. Inform. Terpadu*, vol. 11, no. 1, pp. 55–62, Apr. 2025, doi: 10.54914/jit.v11i1.1721.
- [17] D. S. Nurrochmah, N. Rahaningsih, R. D. Dana, and C. L. Rohmat, “Penerapan Algoritma Naive Bayes dalam Analisis Sentimen Ulasan Aplikasi KitaLulus di Google Play Store,” *J. Inform. Terpadu*, vol. 11, no. 1, pp. 1–11, 2025.
- [18] S. Ernawati and R. Wati, “Evaluasi Performa Kernel SVM dalam Analisis Sentimen Review Aplikasi ChatGPT Menggunakan Hyperparameter dan VADER Lexicon,” *J. Buana Inform.*, vol. 15, no. 01, pp. 40–49, Apr. 2024, doi: 10.24002/jbi.v15i1.7925.
- [19] Y. Matira, Junaidi, and I. Setiawan, “Pemodelan Topik pada Judul Berita Online Detikcom Menggunakan Latent Dirichlet Allocation,” *Estimasi J. Stat. Its Appl.*, vol. 4, no. 1, pp. 53–63, 2023, [Online]. Available: <http://journal.unhas.ac.id/index.php/ESTIMASI>
- [20] K. Pramayasa, I. M. D. Maysanjaya, and I. G. A. A. D. Indradewi, “Analisis Sentimen Program Mbkm Pada Media Sosial Twitter Menggunakan KNN Dan SMOTE,” *SINTECH (Science Inf. Technol. J.)*, vol. 6, no. 2, pp. 89–98, Aug. 2023, doi: 10.31598/sintechjournal.v6i2.1372.