

Artikel Ilmiah.pdf

by Turnitin Official

Submission date: 30-Jan-2026 03:36PM (UTC+0900)

Submission ID: 2866715403

File name: Artikel_Ilmiah.pdf (703.12K)

Word count: 5448

Character count: 34807

Sentiment Classification and the Relationship of Textual Characteristics to Tweet Engagement Levels on Shopee Services

[Klasifikasi Sentimen Dan Hubungan Karakteristik Teks Terhadap Tingkat Engagement Tweet Pada Layanan Shopee]

Ricki Maulana Abdillah ^{*,1)}, Yulian Findawati ²⁾, Irwan Alnarus Kautsar ³⁾, Yunianita Rahmawati ⁴⁾

^{1,2,3,4)} Program Studi Teknik Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

*Email Penulis Korespondensi: yulianfindawati@umsida.ac.id

Abstract. Social media Twitter (X) is widely used by users to express opinions, complaints, and experiences related to e-commerce services such as Shopee. The high volume of conversations creates the need for automated analysis to understand user sentiment and engagement behavior. This study focuses on classifying sentiment in Shopee-related tweets and examining the relationship between textual characteristics and engagement levels. Sentiment classification models were developed using Naïve Bayes and Support Vector Machine (SVM) algorithms with an 80:20 train-test data split. Engagement analysis was supported by FP-Growth, WordCloud, and Crosstab techniques. The results show that both models achieved an accuracy of 68%, with SVM demonstrating more balanced performance across sentiment classes. FP-Growth analysis revealed dominant patterns in positive sentiment with low engagement, while high engagement showed no strong association patterns. These findings suggest that engagement is influenced not only by sentiment, but also by external factors and account characteristics.

Keywords: Analysis Sentiment, Engagement, Naive Bayes, SVM, FP-Growth, Twitter

Abstrak. Media sosial Twitter (X) banyak dimanfaatkan pengguna untuk menyampaikan opini, keluhan, dan pengalaman terhadap layanan e-commerce seperti Shopee. Tingginya volume percakapan menuntut adanya analisis otomatis untuk memahami sentimen dan tingkat engagement pengguna. Penelitian ini berfokus pada klasifikasi sentimen tweet terkait Shopee serta analisis hubungan karakteristik teks terhadap tingkat engagement. Model klasifikasi sentimen dibangun menggunakan algoritma Naïve Bayes dan Support Vector Machine (SVM) dengan pembagian data latih dan uji sebesar 80:20. Analisis engagement didukung oleh metode FP-Growth, WordCloud, dan Crosstab. Hasil penelitian menunjukkan bahwa kedua model mencapai tingkat akurasi sebesar 68%, dengan SVM menunjukkan kinerja yang lebih seimbang pada setiap kelas sentimen. Analisis FP-Growth mengungkapkan pola dominan pada sentimen positif dengan tingkat engagement rendah, sementara pada engagement tinggi tidak ditemukan pola asosiasi yang kuat. Temuan ini menunjukkan bahwa engagement tidak hanya dipengaruhi oleh sentimen, tetapi juga oleh faktor eksternal dan karakteristik akun.

Kata Kunci: Analisis Sentimen, Engagement, Naive Bayes, SVM, FP-Growth, Twitter

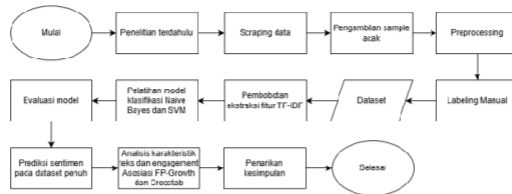
I. PENDAHULUAN

Perkembangan media sosial telah mengubah cara masyarakat menyampaikan opini terhadap suatu produk atau layanan. Twitter (X) menjadi salah satu platform yang banyak digunakan pengguna untuk menyampaikan pengalaman, baik berupa pujian maupun keluhan terhadap layanan e-commerce. Shopee sebagai salah satu platform e-commerce terbesar di Indonesia memiliki volume percakapan yang tinggi di media sosial, sehingga menarik untuk dianalisis lebih lanjut. Kondisi ini menjadikan Twitter sebagai sumber data yang potensial untuk memahami persepsi dan respons pengguna terhadap layanan Shopee.[1][2][3][4][5]

Analisis sentimen digunakan untuk mengklasifikasikan opini pengguna ke dalam kategori positif, netral, atau negatif. Namun, sentimen saja belum cukup untuk menggambarkan dampak suatu konten, karena tingkat keterlibatan pengguna engagement seperti like, retweet, reply, dan views juga punya peran penting. Dalam praktiknya, tweet dengan sentimen positif belum tentu menghasilkan engagement yang tinggi, begitu pula tweet dengan sentimen negatif yang tidak selalu mendapatkan respons rendah. Hal ini menunjukkan adanya permasalahan dalam memahami hubungan antara sentimen dan karakteristik teks terhadap tingkat engagement pengguna di Twitter.[6][7][8][9][10]

Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk mengklasifikasikan sentimen tweet terkait layanan Shopee serta menganalisis hubungan antara karakteristik teks dengan tingkat engagement. Untuk mencapai tujuan tersebut, metode Naïve Bayes dan Support Vector Machine digunakan untuk klasifikasi sentimen, sedangkan FP-Growth, WordCloud, dan Crosstab digunakan untuk analisis karakteristik teks dan pola engagement berdasarkan fitur tweet. Rangkaian metode tersebut diharapkan mampu memberikan gambaran yang lebih komprehensif mengenai keterkaitan sentimen, karakteristik teks, dan tingkat engagement pada tweet terkait layanan Shopee.[11][12][13][14][15][16]

II. METODE



Gambar 1. Alur Tahapan Penelitian

Penelitian ini dilaksanakan melalui beberapa tahapan sistematis guna mencapai hasil yang diharapkan, sebagaimana ditampilkan pada Gambar 1.

2.1 Penelitian Terdahulu

Penelitian mengenai analisis sentimen pada platform Twitter telah banyak dilakukan dengan memanfaatkan pendekatan machine learning berbasis teks, khususnya dalam konteks layanan e-commerce. Beberapa studi menggunakan algoritma Naïve Bayes untuk mengklasifikasi sentimen pengguna terhadap layanan seperti Shopee dan menunjukkan bahwa metode tersebut mampu menghasilkan tingkat akurasi yang cukup baik dalam membedakan sentimen positif, negatif, dan netral. Namun, penelitian-penelitian tersebut umumnya hanya berfokus pada klasifikasi polaritas sentimen tanpa mengaitkannya dengan karakteristik teks maupun tingkat keterlibatan pengguna.

Pendekatan lain menggunakan Support Vector Machine dengan representasi fitur berbasis Word2Vec juga telah diterapkan untuk meningkatkan pemahaman konteks semantik dalam tweet. Hasilnya menunjukkan performa klasifikasi yang lebih stabil, tetapi analisis masih terbatas pada sentimen semata. Aspek interaksi pengguna seperti jumlah like, retweet, reply, dan views, serta karakteristik teks seperti panjang tweet, penggunaan emoji, dan struktur kalimat, belum dianalisis secara komprehensif.

Berdasarkan kondisi tersebut, kebaruan penelitian ini terletak pada pengintegrasian analisis klasifikasi sentimen dengan analisis karakteristik teks terhadap tingkat engagement. Penelitian ini tidak hanya membangun model sentimen menggunakan Naïve Bayes dan Support Vector Machine, tetapi juga mengombinasikannya dengan analisis pola menggunakan FP-Growth dan Crosstab untuk mengidentifikasi hubungan antara fitur teks dan tingkat engagement pada tweet yang berkaitan dengan layanan Shopee.

2.2 Scraping data

Pengumpulan data dilakukan dengan menerapkan teknik web scraping terhadap tweet publik pada platform Twitter (X). Pendekatan ini dipilih sebagai alternatif pengambilan data karena adanya keterbatasan akses terhadap Application Programming Interface (API) resmi Twitter. Proses scraping dilakukan menggunakan tools yang dikembangkan oleh penulis dengan memanfaatkan library Selenium pada bahasa pemrograman Python untuk mengotomatisasi interaksi melalui browser.

Pengambilan data tweet dilakukan berdasarkan kumpulan kata kunci yang merepresentasikan berbagai aspek layanan Shopee, meliputi pengalaman pengguna, keluhan layanan, aktivitas promosi, serta opini yang bersifat emosional. Rentang waktu pengambilan data ditetapkan mulai 1 Januari 2022 hingga 13 Oktober 2025 dengan memanfaatkan parameter *since* dan *until* pada query pencarian. Data hasil scraping selanjutnya disimpan dalam format CSV dan JSON, dengan total sebanyak 18.863 tweet berhasil dikumpulkan.

Mencakup setiap tweet yang diperoleh memiliki 33 fitur, yang mencakup identitas tweet dan waktu publikasi (Tweet ID, Date), informasi akun (Username, Display Name, Verified, Verified Type), konten dan konteks tweet (Content, Query, Panjang Teks, Huruf Kapital, Jumlah Emoji, Kata Persuasif, Layanan Disebut, Layanan Disebut Nama), metrik engagement (Like, Retweet, Reply, Views), serta karakteristik struktur tweet seperti keberadaan media (Has Media, Total Media, Media Type), tautan (Mengandung Link, Jumlah Tautan, Link dalam Konten), hashtag (Hashtag, Jumlah Hashtags, Daftar Hashtags), mention (Mention, Jumlah Mention, Daftar Mention), dan status percakapan (Replying, Replying To). Seluruh fitur tersebut digunakan sebagai dasar analisis sentimen dan keterkaitannya dengan tingkat engagement tweet.

2.3 Pengambilan Sample Acak

Sebelum dilakukan tahapan pra-pemrosesan, dari total 18.863 data hasil scraping dilakukan pengambilan 1.400 data sampel secara acak untuk keperluan pelatihan dan evaluasi model klasifikasi sentimen. Pemilihan sampel

dilakukan untuk memungkinkan proses pelabelan manual dan evaluasi model secara terkontrol. Preprocessing dilakukan pada fitur kolom Content yaitu isi pada tweet.

2.4 Preprocessing

Tahapan preprocessing dilakukan untuk memastikan data teks berada dalam kondisi yang seragam dan siap dianalisis oleh algoritma klasifikasi. Case folding diterapkan untuk mengubah seluruh huruf menjadi huruf kecil sehingga perbedaan penulisan tidak dianggap sebagai kata yang berbeda. Proses cleaning bertujuan untuk menghilangkan elemen yang tidak relevan seperti URL, mention, hashtag, angka, emoji tertentu, dan tanda baca yang tidak memiliki kontribusi langsung terhadap makna sentimen. Tokenisasi dilakukan menggunakan library NLTK untuk memecah teks menjadi unit kata agar setiap kata dapat diproses secara terpisah.

Stopword removal digunakan untuk menghapus kata-kata umum yang sering muncul tetapi tidak memiliki nilai informatif terhadap sentimen, sehingga fokus analisis tertuju pada kata-kata bermakna. Selanjutnya, stemming dilakukan menggunakan library Sastrawi untuk mengembalikan kata ke bentuk dasarnya, sehingga variasi kata dengan makna serupa dapat diperlakukan sebagai satu entitas. Seluruh tahapan preprocessing ini diimplementasikan menggunakan bahasa pemrograman Python dengan bantuan library NLTK dan Sastrawi guna meningkatkan kualitas representasi teks dan akurasi hasil analisis.

2.5 Labeling Manual Dataset

Sebanyak 1.400 data yang telah melalui proses pra-pemrosesan kemudian diberi label sentimen secara manual ke dalam tiga kelas, yaitu positif, netral, dan negatif. Proses pelabelan dilakukan dengan mempertimbangkan konteks isi tweet secara menyeluruh untuk memastikan konsistensi dan akurasi label, ditemukan 565 positif, 472 negatif, 363 netral.

Data berlabel keseluruhan 1400 selanjutnya digunakan untuk melatih model klasifikasi sentimen dengan skema pembagian data 80:20, di mana 80% data (1.120 tweet) digunakan sebagai data latih dan 20% data (280 tweet) digunakan sebagai data uji. Pembagian ini bertujuan untuk mengevaluasi performa model secara objektif terhadap data yang tidak dilibatkan dalam proses pelatihan.

2.6 Pembobotan TF-IDF

Ekstraksi fitur teks dilakukan menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF) untuk merepresentasikan data teks ke dalam bentuk numerik. Pembobotan TF-IDF digunakan untuk memberikan penekanan pada kata-kata yang memiliki tingkat kepentingan lebih tinggi dalam suatu dokumen dibandingkan keseluruhan korpus. Secara konseptual, pembobotan TF-IDF dapat dirumuskan sebagai berikut:

$$TF-IDF(t, d) = TF(t, d) \times \log \left(\frac{N}{DF(t)} \right) \quad (1)$$

Dengan $TF(t, d)$ menyatakan frekuensi kemunculan kata t pada dokumen d , $DF(t)$ menunjukkan jumlah dokumen yang mengandung kata t , dan N merupakan jumlah total dokumen dalam korpus. Proses ekstraksi fitur TF-IDF dilakukan menggunakan tools dengan parameter TF-IDF dengan memanfaatkan library TF-IDF dari scikit-learn, baik data latih maupun data uji diproses menggunakan representasi TF-IDF yang sama, di mana vocabulary dibangun berdasarkan data latih dan selanjutnya diterapkan pada data uji. Representasi numerik ini kemudian digunakan sebagai masukan untuk proses pelatihan dan evaluasi model klasifikasi sentimen.

2.7 Pelatihan Model Naive Bayes & SVM

Naive Bayes merupakan metode klasifikasi berbasis probabilistik yang mengacu pada Teorema Bayes untuk menentukan kecenderungan kelas suatu dokumen teks. Secara konseptual dapat dijelaskan melalui rumus berikut:

$$P(H | X) = \frac{P(X | H)P(H)}{P(X)} \quad (2)$$

Rumus menjelaskan bahwa probabilitas suatu tweet termasuk ke dalam kelas sentimen tertentu ditentukan oleh peluang kemunculan fitur teks pada kelas tersebut serta peluang awal dari masing-masing kelas sentimen. Nilai $P(X|H)$ merepresentasikan seberapa besar kemungkinan kata atau fitur muncul pada suatu kelas, sedangkan $P(H)$ menunjukkan distribusi awal kelas sentimen dalam data, dan $P(X)$ berperan sebagai faktor normalisasi. Dalam penelitian ini, proses klasifikasi sentimen tidak dilakukan melalui perhitungan manual berdasarkan rumus tersebut, melainkan diimplementasikan menggunakan parameter Naive Bayes dengan memanfaatkan library Naive Bayes dari scikit-learn. Library tersebut secara internal menangani perhitungan probabilitas berdasarkan representasi fitur TF-IDF yang telah dibentuk sebelumnya.

Support Vector Machine (SVM) merupakan metode klasifikasi terawasi yang bertujuan untuk mencari hyperplane optimal dalam memisahkan data ke dalam kelas-kelas tertentu. Untuk kernel linear, fungsi keputusan SVM secara teoritis dapat dijelaskan sebagai berikut:

$$f(x) = w^T x + b \quad (3)$$

Persamaan tersebut menggambarkan fungsi keputusan pada Support Vector Machine, di mana vektor fitur x diproyeksikan ke dalam ruang pemisah melalui bobot w dan bias b untuk menentukan posisi data terhadap hyperplane. Nilai keluaran fungsi $f(x)$ digunakan sebagai dasar penentuan kelas sentimen, dengan tanda positif atau negatif menunjukkan kecenderungan kelas yang berbeda. Pada penelitian ini, algoritma SVM diimplementasikan menggunakan parameter SVM dengan memanfaatkan library SVM dari scikit-learn dengan kernel linear. Proses pelatihan dan pengujian model dilakukan sepenuhnya melalui fungsi bawaan library tanpa melakukan perhitungan matematis secara manual.

2.8 Evaluasi Model

Evaluasi performa model dilakukan untuk menilai kemampuan generalisasi model dalam mengklasifikasikan data uji. Proses evaluasi didukung oleh confusion matrix yang menggambarkan perbandingan antara hasil prediksi model dan label aktual. Confusion matrix tersebut terdiri dari beberapa komponen utama, yaitu terdiri dari komponen True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Berdasarkan nilai-nilai tersebut, performa model dianalisis menggunakan metrik accuracy, precision, recall, dan F1-score.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$Precision = \frac{TP}{TP + FP} \quad (5)$$

$$Recall = \frac{TP}{TP + FN} \quad (6)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (7)$$

Secara teoritis, accuracy mengukur tingkat ketepatan prediksi secara keseluruhan, precision menunjukkan ketelitian model dalam memprediksi suatu kelas, recall menggambarkan kemampuan model dalam menemukan seluruh data yang relevan, sedangkan F1-score digunakan untuk menyeimbangkan nilai precision dan recall. Dalam penelitian ini, seluruh metrik evaluasi dihitung menggunakan parameter classification report dan confusion matrix dari scikit-learn sehingga proses evaluasi dilakukan secara otomatis dan konsisten.

2.9 Analisis Karakteristik Teks dan Engagement FP-Growth dan Crosstab

FP-Growth merupakan algoritma pencarian pola asosiasi yang digunakan untuk menemukan keterkaitan antar fitur tanpa membangkitkan kandidat itemset. Secara konseptual, analisis asosiasi pada FP-Growth didasarkan pada ukuran support, confidence, dan lift[9], yang masing-masing dapat dijelaskan sebagai berikut:

$$Support(X) = \frac{Jumlah\ tweet\ yang\ memuat\ X}{Total\ tweet} \quad (8)$$

$$Confidence(X \rightarrow Y) = \frac{Support(X \cup Y)}{Support(X)} \quad (9)$$

$$Lift(X \rightarrow Y) = \frac{Confidence(X \rightarrow Y)}{Support(Y)} \quad (10)$$

Rumus-rumus tersebut digunakan sebagai dasar interpretasi pola yang dihasilkan. Pada penelitian ini, analisis FP-Growth diimplementasikan menggunakan parameter FP-Growth dengan memanfaatkan library mlxtend. Library ini digunakan untuk mengidentifikasi pola hubungan antara karakteristik teks dan kategori engagement (rendah, sedang, dan tinggi) secara efisien.

Crosstab atau tabulasi silang digunakan untuk melihat hubungan antara dua variabel kategori secara deskriptif. Nilai pada setiap sel Crosstab menunjukkan jumlah data yang berada pada kombinasi kategori tertentu. Secara konseptual, nilai sel Crosstab dapat dinyatakan sebagai:

$$f_{ij} = \text{jumlah data yang termasuk kategori } i \text{ pada variabel pertama dan kategori } j \text{ pada variabel kedua} \quad (11)$$

Dengan i menyatakan kategori pada variabel pertama (misalnya tingkat engagement) dan j menyatakan kategori pada variabel kedua (misalnya karakteristik teks). Dalam penelitian ini, Crosstab digunakan untuk memperkuat interpretasi hasil FP-Growth dengan menunjukkan distribusi keterkaitan antara fitur teks dan tingkat engagement tweet secara kuantitatif.

III. HASIL DAN PEMBAHASAN

3.1 Evaluasi Model Naive Bayes dan SVM

Evaluasi model dilakukan menggunakan data uji sebanyak 280 tweet yang diperoleh dari hasil pembagian data berlabel dengan rasio 80:20. Performa model dianalisis menggunakan classification report dan confusion matrix untuk melihat kemampuan klasifikasi pada setiap kelas sentimen. Hasil pengujian menunjukkan bahwa model Naive Bayes dan Support Vector Machine (SVM) sama-sama menghasilkan nilai akurasi sebesar 68%. Nilai akurasi tersebut dipengaruhi oleh ketidakseimbangan distribusi data pada setiap kelas sentimen serta adanya tumpang tindih karakteristik teks, khususnya pada kelas netral yang memiliki kemiripan fitur dengan sentimen positif dan negatif, sehingga meningkatkan tingkat kesalahan klasifikasi.

Classification Report:				
	precision	recall	f1-score	support
negatif	0.72	0.45	0.58	94
netral	0.49	0.25	0.35	73
positif	0.64	0.81	0.72	113
accuracy			0.68	280
macro avg	0.60	0.50	0.60	280
weighted avg	0.68	0.68	0.68	280

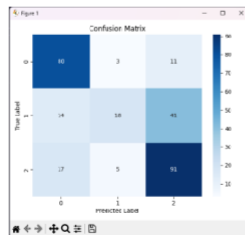
(a)

Classification Report:				
	precision	recall	f1-score	support
negatif	0.72	0.81	0.76	94
netral	0.56	0.44	0.49	73
positif	0.78	0.73	0.77	113
accuracy			0.68	280
macro avg	0.69	0.66	0.66	280
weighted avg	0.67	0.68	0.67	280

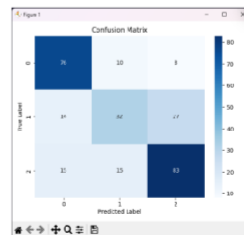
(b)

Gambar 2. Classification Report Naive Bayes dan SVM

Berdasarkan classification report pada Gambar 2(a), model Naive Bayes menunjukkan performa terbaik pada kelas negatif dengan nilai recall sebesar 85%, yang menandakan bahwa sebagian besar tweet negatif berhasil diklasifikasikan dengan baik. Kelas positif juga memiliki performa yang cukup baik dengan nilai F1-score sebesar 71%. Sebaliknya, kelas netral menunjukkan performa terendah dengan nilai F1-score sebesar 36%, yang mengindikasikan bahwa model masih mengalami kesulitan dalam membedakan sentimen netral dari sentimen positif dan negatif. Sementara itu, berdasarkan classification report pada Gambar 2(b), model Support Vector Machine (SVM) memperlihatkan distribusi performa yang lebih merata pada setiap kelas sentimen. Kelas positif memperoleh nilai F1-score sebesar 72%, sedangkan kelas negatif mencapai nilai F1-score sebesar 76%. Performa kelas netral pada model SVM juga mengalami peningkatan dibandingkan dengan Naive Bayes, dengan nilai F1-score sebesar 49%, meskipun tetap menjadi kelas dengan performa terendah.



(a)

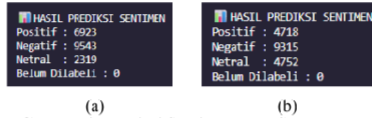


(b)

Gambar 3. Confusion Matrix Naive Bayes dan SVM

Hasil analisis tersebut diperkuat oleh confusion matrix pada kedua model. Pada model Naïve Bayes Gambar 3 (a), dari total 73 tweet bersentimen netral, hanya 18 tweet yang berhasil diklasifikasikan dengan benar sebagai netral, sementara 14 tweet salah diklasifikasikan sebagai negatif dan 41 tweet diprediksi sebagai positif. Kondisi ini menunjukkan adanya kemiripan karakteristik teks antara sentimen netral dengan sentimen positif maupun negatif, sehingga sulit dipisahkan oleh Naïve Bayes yang bekerja berdasarkan asumsi independensi antar kata. Sementara itu, berdasarkan confusion matrix pada model Support Vector Machine (SVM) Gambar 3 (b), dari 73 data netral sebanyak 32 data berhasil diklasifikasikan dengan benar, sedangkan 14 data salah diprediksi sebagai negatif dan 27 data sebagai positif. Jumlah kesalahan klasifikasi pada model SVM lebih kecil dibandingkan Naïve Bayes, yang mengindikasikan bahwa SVM memiliki kemampuan pemisahan kelas yang lebih baik dalam ruang fitur TF-IDF, terutama dalam menangani data dengan pola yang saling tumpang tindih.

3.2 Prediksi Sentimen Data Skala Besar



Gambar 4. Prediksi Sentimen Data Skala Besar

Berdasarkan hasil prediksi label sentimen yang ditunjukkan pada gambar 4, model Support Vector Machine (SVM) (b) dan Naïve Bayes (a) menghasilkan distribusi sentimen yang didominasi oleh sentimen positif dan negatif, sementara sentimen netral muncul dalam jumlah lebih kecil. Pada prediksi menggunakan SVM, diperoleh 4.718 tweet bersentimen positif, 9.315 tweet bersentimen negatif, dan 4.752 tweet bersentimen netral dari total 18.863 data. Sementara itu, prediksi menggunakan Naïve Bayes menghasilkan 6.923 tweet bersentimen positif, 9.543 tweet bersentimen negatif, dan 2.319 tweet bersentimen netral. Perbedaan distribusi ini menunjukkan bahwa SVM cenderung mempertahankan proporsi sentimen netral lebih besar, sedangkan Naïve Bayes lebih agresif dalam mengklasifikasikan data ke kelas positif dan negatif. Temuan ini konsisten dengan hasil evaluasi sebelumnya dan memperkuat pemilihan SVM sebagai model dengan perilaku klasifikasi yang lebih stabil dan konservatif pada data berskala besar.

Perlu dicatat bahwa jumlah data yang digunakan pada tahap prediksi mengalami sedikit pengurangan dibandingkan jumlah data mentah awal. Pengurangan ini terjadi karena pada tahap pra-proses teks, terdapat sejumlah tweet yang menghasilkan nilai kosong (null) pada kolom teks hasil pra-proses, sehingga tidak dapat direpresentasikan dalam bentuk vektor TF-IDF. Tweet-tweet tersebut umumnya hanya mengandung simbol, tautan, emoji, atau karakter non-teks yang seluruhnya terhapus selama proses pembersihan data. Oleh karena itu, data dengan nilai kosong tersebut dihapus (drop) untuk menghindari kesalahan pada proses transformasi fitur dan prediksi model, sehingga hanya data yang valid secara tekstual yang digunakan dalam analisis sentimen.

3.3 Klasifikasi Engagement FP-Growth

Tingkat engagement diperoleh dari akumulasi jumlah like, retweet, reply, dan views pada setiap tweet. Nilai engagement tersebut kemudian dikelompokkan ke dalam tiga kategori, yaitu engagement rendah ($\leq 58,67$), engagement sedang ($58,67-271,34$), dan engagement tinggi ($> 271,34$). Pengelompokan ini menghasilkan tiga kelas engagement, di mana setiap tweet memiliki nilai numerik serta label kategori engagement yang sesuai, dengan total data yang digunakan sebanyak 1400 tweet. Sebagaimana hasil kategori engagement pada Tabel 1.

Tabel 1. Kategori Engagement

Angka Engagement	Kategori Engagement
189094	Tinggi
866	Tinggi
139	Sedang
58	Rendah
0	Rendah

Analisis FP-Growth menunjukkan bahwa pola asosiasi yang signifikan hanya muncul pada kategori engagement rendah. Pola dominan memperlihatkan bahwa tweet bersentimen positif, akun tidak terverifikasi, tanpa mention dan

tautan, serta memiliki satu media cenderung masuk ke dalam engagement rendah, dengan nilai confidence > 0,50 dan lift > 1. Sementara itu, pada kategori engagement sedang dan tinggi tidak ditemukan pola yang memenuhi ambang signifikansi, yang mengindikasikan bahwa engagement tinggi tidak dapat dijelaskan hanya oleh karakteristik teks, melainkan dipengaruhi oleh faktor lain di luar konten. Sebagaimana ditampilkan pada Tabel 2

Tabel 2. FP-Growth Engagement

Sentimen	Kategori Engagement	Antecedents	Support	Confidence	Lift
Positif	Rendah	frozenset({'Total Media_1', 'label_sentimen_positif', 'Jumlah Mention_0', 'Mengandung Link_False', 'Verified_False'})	0.05214	0.50694	6.22563
Positif	Rendah	frozenset({'Mention_False', 'Total Media_1', 'label_sentimen_positif', 'Jumlah Tautan_0', 'Verified_False'})	0.05214	0.50694	6.22563

Analisis FP-Growth menunjukkan bahwa pola asosiasi yang signifikan hanya muncul pada kategori engagement rendah. Pola dominan memperlihatkan bahwa tweet bersentimen positif, akun tidak terverifikasi, tanpa mention dan tautan, serta memiliki satu media cenderung masuk ke dalam engagement rendah, dengan nilai confidence > 0,50 dan lift > 1. Sementara itu, pada kategori engagement sedang dan tinggi tidak ditemukan pola yang memenuhi ambang signifikansi, yang mengindikasikan bahwa engagement tinggi tidak dapat dijelaskan hanya oleh karakteristik teks, melainkan dipengaruhi oleh faktor lain di luar konten.

3.4 Analisis WordCloud



Gambar 5. Analisis WordCloud

Berdasarkan visualisasi WordCloud pada Gambar 5(a), sentimen negatif dengan engagement rendah didominasi oleh kata-kata yang merepresentasikan keluhan terhadap layanan, seperti susah, berat, dan lama, yang menggambarkan pengalaman pengguna yang kurang memuaskan. Pada Gambar 5(b), sentimen netral dengan engagement sedang memperlihatkan kemunculan kata-kata yang bersifat informatif, seperti voucher, order, dan bayar, yang menunjukkan bahwa tweet netral umumnya digunakan untuk menyampaikan informasi atau aktivitas transaksi tanpa muatan emosi yang kuat. Sementara itu, pada Gambar 5(c), sentimen positif dengan engagement tinggi didominasi oleh kata-kata yang berkaitan dengan promosi dan kepuasan pengguna, seperti diskon, promo, gratis ongkir, dan murah. Pola ini mengindikasikan bahwa konten bermuatan positif, khususnya yang mengandung unsur promosi, cenderung menarik perhatian pengguna dan menghasilkan tingkat engagement yang lebih tinggi.

Kemunculan kata Shopee yang relatif dominan pada seluruh visualisasi WordCloud disebabkan oleh ruang lingkup penelitian yang secara khusus memfokuskan analisis pada tweet yang berkaitan dengan layanan Shopee. Seluruh data yang digunakan diperoleh melalui proses pengambilan tweet berbasis kata kunci dan konteks yang relevan dengan Shopee, sehingga penyebutan nama layanan tersebut secara alami menjadi bagian dari percakapan pengguna, baik dalam bentuk keluhan, informasi, maupun apresiasi. Kondisi ini menunjukkan bahwa kata Shopee berperan sebagai konteks utama dalam pembentukan sentimen dan bukan sebagai indikator polaritas sentimen tertentu, melainkan sebagai penanda topik yang merepresentasikan objek penelitian.

3.5 Analisis Crosstab

Tabel 3. Crosstab Engagement Rendah

Fitur	Label Sentimen	Frekuensi	Persentase (%)
Has Media	Positif	218	50
Huruf Kapital	Positif	218	1.28
Jumlah Emoji	Positif	218	6.67

Berdasarkan hasil Crosstab pada kategori engagement rendah, terdapat 218 tweet bersentimen positif. Kelompok ini menunjukkan kecenderungan penggunaan media dan hashtag yang relatif tinggi, masing-masing sebesar 50%. Namun, fitur lain seperti penggunaan huruf kapital, panjang teks, dan jumlah emoji hanya muncul dalam proporsi yang sangat kecil. Hal ini menunjukkan bahwa meskipun jumlah tweet positif cukup besar, karakteristik teks yang digunakan belum cukup kuat untuk mendorong interaksi pengguna yang lebih tinggi.

Tabel 4. Crosstab Engagement Sedang

Fitur	Label Sentimen	Frekuensi	Persentase (%)
Has Media	Negatif	208	50
Hashtag	Negatif	208	50
Jumlah Hashtags	Negatif	208	7.69

Pada kategori engagement sedang, sebanyak 208 tweet bersentimen negatif mendominasi kelompok ini. Pola yang terbentuk menunjukkan bahwa penggunaan media dan hashtag muncul pada sekitar 50% data, sementara fitur pendukung lainnya berada pada persentase yang rendah. Temuan ini menunjukkan bahwa engagement sedang tidak hanya dipengaruhi oleh polaritas sentimen, tetapi juga oleh keterbatasan variasi karakteristik teks yang digunakan.

Tabel 5. Crosstab Engagement Tinggi

Fitur	Label Sentimen	Frekuensi	Persentase (%)
Has Media	Positif	189	50
Huruf Kapital	Positif	189	1.28
Jumlah Emoji	Positif	189	6.67

Sementara itu, pada kategori engagement tinggi terdapat 189 tweet bersentimen positif dengan pola karakteristik yang relatif serupa dengan engagement rendah, terutama pada penggunaan media dan hashtag. Tidak ditemukan satu fitur teks yang benar-benar dominan dalam mendorong engagement tinggi. Selain itu, sentimen netral tidak muncul sebagai kelompok dominan pada seluruh kategori engagement. Hal ini disebabkan karena distribusi tweet netral tersebar relatif merata dan tidak menunjukkan kecenderungan kuat terhadap karakteristik teks tertentu. Temuan ini menunjukkan bahwa sentimen netral cenderung menghasilkan respons pengguna yang stabil namun tidak ekstrem, serta memperkuat kesimpulan bahwa engagement tinggi lebih dipengaruhi oleh kombinasi faktor lain, termasuk karakteristik akun dan faktor eksternal di luar isi teks.

VII. SIMPULAN

Penelitian ini berhasil melakukan klasifikasi sentimen tweet yang berkaitan dengan layanan Shopee menggunakan metode Naïve Bayes dan Support Vector Machine (SVM) dengan tingkat akurasi sebesar 68%. Nilai akurasi tersebut dipengaruhi oleh karakteristik data teks media sosial yang memiliki perbedaan distribusi antar kelas serta adanya tumpang tindih makna pada sentimen netral yang kerap memiliki kemiripan dengan sentimen positif maupun negatif, sehingga meningkatkan potensi kesalahan klasifikasi. Hasil analisis engagement menunjukkan bahwa sentimen positif tidak selalu diikuti oleh tingkat engagement yang tinggi, serta pola asosiasi yang signifikan hanya ditemukan pada kategori engagement rendah. Temuan ini mengindikasikan bahwa engagement tidak semata-mata ditentukan oleh polaritas sentimen, melainkan juga dipengaruhi oleh faktor lain di luar isi teks, seperti karakteristik akun pengguna dan kondisi eksternal. Penelitian selanjutnya disarankan untuk mempertimbangkan penambahan variabel temporal,

struktur interaksi antar pengguna, serta pendekatan analisis topik guna memperoleh pemahaman yang lebih komprehensif mengenai faktor-faktor yang memengaruhi engagement.

2 UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada seluruh pihak yang telah memberikan dukungan, arahan, dan bantuan selama proses penyusunan penelitian ini, baik secara akademik maupun nonakademik. Kontribusi berupa bimbingan keilmuan, masukan konstruktif, serta fasilitas pendukung sangat membantu penulis dalam menyelesaikan penelitian hingga tahap akhir. Selain itu, dukungan moral dan motivasi yang diberikan selama proses penelitian juga menjadi faktor penting dalam menjaga keberlangsungan dan kualitas penyusunan karya ini. Penulis berharap hasil penelitian ini dapat memberikan manfaat dan menjadi referensi bagi pengembangan kajian selanjutnya.

REFERENSI

- [1] A. N. Puspitasari, Y. Findawati, and Y. Rahmawati, "ANALISIS SENTIMEN TWEET PENGGUNA E-COMMERCE DENGAN MENGGUNAKAN METODE KLASIFIKASI NAIVE BAYES," *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 9, no. 3, pp. 1123–1132, Aug. 2024, doi: 10.29100/jipi.v9i3.4939.
- [2] P. S. Hutapea and W. Maharani, "Sentiment Analysis on Twitter Social Media towards Shopee E-Commerce through Support Vector Machine (SVM) Method," *JINAV: Journal of Information and Visualization*, vol. 4, no. 1, pp. 7–17, Jan. 2023, doi: 10.35877/454ri.jinav1504.
- [3] A. Syah, F. Nurdyansyah, and A. Y. Rahman, "ANALISIS SENTIMEN APLIKASI SHOPEE, TOKOPEDIA, LAZADA DAN BLIBLI MENGGUNAKAN LEKSIKON DAN RANDOM FOREST," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 3S1, Oct. 2024, doi: 10.23960/jitet.v12i3s1.5155.
- [4] R. Rameez, H. A. Rahmani, and E. Yilmaz, "ViralBERT: A User Focused BERT-Based Approach to Virality Prediction," in *UMAP2022 - Adjunct Proceedings of the 30th ACM Conference on User Modeling, Adaptation and Personalization*, Association for Computing Machinery, Inc, Jul. 2022, pp. 85–89, doi: 10.1145/3511047.3536415.
- [5] E. Saquete, J. Zubcoff, Y. Gutierrez, P. Martinez-Barco, and J. Fernández, "Why are some social-media contents more popular than others? Opinion and association rules mining applied to virality patterns discovery," *Expert Syst. Appl.*, vol. 197, Jul. 2022, doi: 10.1016/j.eswa.2022.116676.
- [6] N. Al Maeni, N. Suarna, and W. Prihartono, "ANALISIS SENTIMEN PENGGUNA SHOPEE BERDASARKAN DATA TWEET DARI TWITTER MENGGUNAKAN METODE NAIVE BAYES CLASSIFIER," 2024.
- [7] F. S. Mufidah, S. Winarno, F. Alzami, E. D. Udayanti, and R. R. Sani, "Analisis Sentimen Masyarakat Terhadap Layanan ShopeeFood Melalui Media Sosial Twitter Dengan Algoritma Naive Bayes Classifier," *JOINS (Journal of Information System)*, vol. 7, no. 1, pp. 14–25, May 2022, doi: 10.33633/joins.v7i1.5883.
- [8] S. Jaffali, "Analyzing the Impact of Emotional Content in Tweets on User Engagement and Online Behavior," 2025.
- [9] B. Z. Ramadhan, I. Riza, and I. Maulana, "Analisis Sentimen Ulasan Pada Aplikasi E-Commerce Dengan Menggunakan Algoritma Naive Bayes," 2022. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [10] I. Syahrohimi, S. D. Saputra, R. W. Saputra, V. H. Pranatawijaya, and R. Priskila, "PERBANDINGAN ANALISIS SENTIMEN SETELAH PILPRES 2024 DI TWITTER MENGGUNAKAN ALGORITMA MACHINE LEARNING," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 2, Apr. 2024, doi: 10.23960/jitet.v12i2.4249.
- [11] Z. Purwanti, P. Studi Sistem Informasi, S. Tinggi Ilmu Komputer Cipta Karya Informatika, K. Jakarta Timur, and D. Khusus Ibukota Jakarta, "Pemodelan Text Mining untuk Analisis Sentimen Terhadap Program Makan Siang Gratis di Media Sosial X Menggunakan Algoritma Support Vector Machine (SVM)," 2024. [Online]. Available: <https://journal.stmiki.ac.id>
- [12] M. I. Prayugah and N. Ariyanti, "ANALISIS SENTIMEN PUBLIK PADA PEMERINTAH DALAM SERANGAN RANSOMWARE DENGAN PENDEKATAN SMOTE," 2025. Accessed: Jul. 08, 2025. [Online]. Available: <https://archive.umsida.ac.id/index.php/archive/preprint/view/6971>
- [13] Zaenal and Ratna I, "Sentiment Analysis of OYO App Reviews Using the Support Vector Machine Algorithm Analisis Sentimen terhadap Ulasan Aplikasi OYO menggunakan Algoritma Support Vector Machine Zaenal, Ika Ratna Indra Astutik," 2022. Accessed: Jul. 08, 2025. [Online]. Available: <https://archive.umsida.ac.id/index.php/archive/preprint/view/424/version/416>

- [14] Y. Ma, "Construction and Data Analysis of a New Media Content Popularity Prediction Model Based on Naive Bayes Algorithm," in *Procedia Computer Science*, Elsevier B.V., 2025, pp. 294–302. doi: 10.1016/j.procs.2025.04.207.
- [15] F. Syah Zikri and M. Ikhsan, "The Comparison Between The Apriori Algorithm And The FP-Growth Algorithm In Determining Frequent Pattern," *INOVTEK Polbeng - Seri Informatika*, vol. 10, no. 2, pp. 615–625, Apr. 2025, doi: 10.35314/s1yanj03.
- [16] A. N. Samsi, Y. Findawati, I. A. Kautsar, and U. M. Sidoarjo, "SNESTIK Seminar Nasional Teknik Elektro, Sistem Informasi, dan Teknik Informatika Analisa Sentimen Mahasiswa Terhadap Minat Berorganisasi Dengan Metode Naive Bayes," p. 171, doi: 10.31284/p.sneistik.2022.2712.

Conflict of Interest Statement:

The author declares that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Artikel Ilmiah.pdf

ORIGINALITY REPORT

12%	10%	5%	11%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

PRIMARY SOURCES

1	Submitted to Exeed College Student Paper	10%
2	archive.umsida.ac.id Internet Source	2%
3	Ichsani Mursidah, Remi Sanjaya, Bambang Yulianto, Dhian Sweetania, Puji Sularsih. "Klasifikasi Sentimen Google Play Store Aplikasi ChatGPT Berbahasa Indonesia Berbasis IndoBERT", Jurnal Minfo Polgan, 2025 Publication	2%

Exclude quotes	On	Exclude matches	< 2%
Exclude bibliography	On		