

KLASIFIKASI SENTIMEN DAN HUBUNGAN KARAKTERISTIK TEKS TERHADAP TINGKAT ENGAGEMENT TWEET PADA LAYANAN SHOPEE

Ricki Maulana Abdillah, Yulian Findawati, Irwan Alnarus Kautsar, Yunianita Rahmawati

Informatika, Universitas Muhammadiyah Sidoarjo

Jl. Mojopahit No. 666 B, Sidowayah, Celep, Kec. Sidoarjo, Kabupaten Sidoarjo, Jawa Timur 61215, Indonesia
221080200127@umsida.ac.id

ABSTRAK

Media sosial *Twitter (X)* menjadi salah satu sarana utama pengguna dalam menyampaikan opini, keluhan, maupun pengalaman terhadap layanan *e-commerce* seperti *Shopee*. Tingginya volume percakapan yang bersifat tidak terstruktur mendorong perlunya analisis otomatis untuk memahami sentimen pengguna serta keterkaitannya dengan tingkat *engagement*. Permasalahan penelitian ini berfokus pada proses klasifikasi sentimen *tweet* terkait *Shopee* serta analisis hubungan karakteristik teks terhadap tingkat *engagement*. Penelitian ini bertujuan untuk membangun model klasifikasi sentimen dan mengidentifikasi kecenderungan pola *engagement* berdasarkan fitur pada konten *tweet*. Metode yang digunakan meliputi *Naïve Bayes* dan *Support Vector Machine (SVM)* sebagai algoritma klasifikasi sentimen, pembagian data dengan rasio 80:20 antara data latih dan data uji, serta *FP-Growth*, *WordCloud*, dan *Crosstab* untuk mendukung analisis *engagement*. Hasil penelitian menunjukkan bahwa kedua model menghasilkan nilai akurasi sebesar 68%, dengan *SVM* memiliki performa yang lebih seimbang pada setiap kelas sentimen. Nilai akurasi tersebut dipengaruhi oleh ketidakseimbangan distribusi kelas serta tumpang tindih karakteristik teks pada sentimen netral yang memiliki kemiripan dengan sentimen positif dan negatif. Analisis *FP-Growth* menunjukkan pola dominan pada sentimen positif dengan *engagement* rendah, sementara *engagement* tinggi tidak membentuk pola asosiasi yang kuat. Temuan ini mengindikasikan bahwa *engagement* tidak hanya dipengaruhi oleh sentimen, tetapi juga oleh faktor eksternal dan karakteristik akun pengguna.

Kata kunci : *Sentiment Analysis, Engagement, Naive Bayes, SVM, FP-Growth, Twitter*

1. 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]

Analisis sentimen digunakan untuk mengklasifikasikan opini pengguna ke dalam kategori positif, netral, atau negatif. Namun, sentimen aja 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* [4][5][6] [7][8]

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 *Naive 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*. [9]

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Penelitian terkait analisis sentimen pada *Twitter* telah banyak dilakukan dengan pendekatan machine learning berbasis teks. Puspitasari [4] menerapkan metode *Naïve Bayes* untuk menganalisis sentimen pengguna terhadap layanan *e-commerce*, termasuk *Shopee*, dengan memanfaatkan data *tweet* sebagai sumber utama. Hasil penelitian tersebut menunjukkan bahwa *Naïve Bayes* mampu menghasilkan tingkat akurasi yang cukup baik dalam mengklasifikasikan sentimen pengguna. Namun demikian, penelitian tersebut masih terbatas pada klasifikasi polaritas sentimen dan belum mengaitkan hasil analisis dengan karakteristik teks maupun tingkat *engagement* pengguna.

Penelitian lain dilakukan oleh Hutapea dan Maharani yang menggunakan *Support Vector Machine (SVM)* berbasis *Word2Vec* untuk mengklasifikasikan sentimen *tweet* terkait layanan *Shopee*. Pendekatan ini terbukti mampu memberikan performa yang lebih baik

dalam menangkap konteks semantik teks [10][11][12]. Meskipun demikian, fokus penelitian tersebut masih terbatas pada analisis sentimen, tanpa mengkaji hubungan antara sentimen dengan interaksi pengguna seperti *like*, *retweet*, *reply*, dan *views*. Selain itu, karakteristik teks seperti panjang *tweet*, penggunaan *emoji*, dan struktur kalimat belum dianalisis secara mendalam. [1]

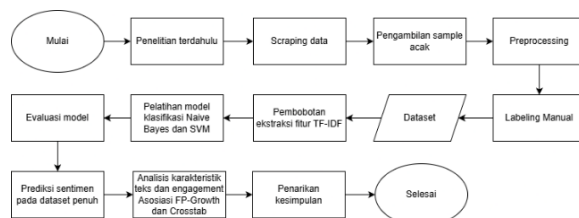
Berdasarkan penelitian terdahulu tersebut, dapat disimpulkan bahwa sebagian besar studi masih berfokus pada klasifikasi sentimen menggunakan algoritma machine learning seperti *Naïve Bayes* dan *Support Vector Machine*. Analisis yang mengaitkan sentimen dengan karakteristik teks serta tingkat *engagement* pengguna masih relatif terbatas. Oleh karena itu, diperlukan penelitian lanjutan yang tidak hanya mengklasifikasikan sentimen, tetapi juga menganalisis hubungan antara karakteristik teks dan tingkat *engagement* pengguna pada tweet terkait layanan Shopee.

2.2. Kebaruan Penelitian

Kebaruan penelitian ini terletak pada integrasi antara analisis klasifikasi sentimen dan analisis karakteristik teks terhadap tingkat *engagement*. Penelitian ini mengombinasikan metode klasifikasi sentimen (*Naïve Bayes* dan *SVM*) dengan analisis pola menggunakan *FP-Growth* dan *Crosstab* untuk mengidentifikasi hubungan antara struktur teks dan tingkat *engagement* pada tweet tentang layanan Shopee.[13][14][15]

3. METODE PENELITIAN

3.1. Alur Penelitian



Gambar 1. Alur Tahapan Penelitian

Alur penelitian tahapan pada gambar 1, dimulai dari penelitian terdahulu, scraping data *Twitter*, pengambilan sampel acak, *preprocessing*, pelabelan manual, ekstraksi fitur *TF-IDF*, pelatihan dan evaluasi model, prediksi sentimen pada *dataset* penuh, analisis *engagement*, hingga penarikan kesimpulan.

3.2. Pengumpulan Data

Pengumpulan data dilakukan dengan menerapkan teknik web scraping pada tweet publik di platform *Twitter* (X). Metode ini digunakan sebagai alternatif pengambilan data mengingat adanya keterbatasan akses terhadap *API* resmi *Twitter*. Proses scraping dilakukan menggunakan tools yang dikembangkan oleh penulis dengan memanfaatkan

library *Selenium* pada bahasa pemrograman *Python* untuk mengotomatisasi interaksi pada *browser*.

Pengambilan tweet dilakukan berdasarkan kumpulan kata kunci yang merepresentasikan berbagai aspek layanan Shopee, seperti 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 menggunakan parameter *since* dan *until* pada query pencarian. Data hasil scraping kemudian disimpan dalam format *CSV* dan *JSON*, dengan total sebanyak 18.863 tweet berhasil dikumpulkan.

Tabel 1. Dataset

Dataset
1770660522595336544,2024-03-21T03:55:40.000Z,Somay kukus voc 10k jasa isi kuesioner,slythryoongs,True,@ShopeeID,"Aku reply ulangg, trnyta yg awal diapus yaa postingannya. Ini aku co kmren, sbnrnya co barang lain di akun adekku. Semuanya pake voucher video, ini ngebantu bngt buat aku yg gprnh dpt voucher yg rebutan . Makasih shopee ehehe... 🐱❤️",Shopee rebutan barang, 0, 0, 0, 119, True, 2, image(2), False, 0,, 233,2,4,True,True,shopee,False,0,,False,,False,0,,https://x.com/slythryoongs/status/1770660522595336544

Pada tabel 1 terdapat 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.

3.3. Pra-pemrosesan Data

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.

Tabel 2. Pra-pemrosesan Data

Content	Preprocessed
Aku reply ulangg, trnyta yg awal diapus yaa postingannya. Ini aku co kmren, sbnrnya co barang	reply ulangg trnyta awal diapus postingannya co kmren sbnrnya co barang lain akun adekku semua

Content	Preprocessed
lain di akun adekku. Semuanya pake voucher video, ini ngebantu bngt buat aku yg gprnh dpt voucher yg rebutan . Makasih shopee ehehe... 🐱❤	pake voucher video ngebantu bngt buat gprnh dpt voucher rebut makasih shopee ehehe

Tahapan pra-pemrosesan dilakukan pada tabel 2 untuk meningkatkan kualitas data teks sebelum dianalisis. Proses ini meliputi *case folding*, *cleaning*, *tokenisasi* memakai *NLTK*, *stopword removal*, dan *stemming* memakai *Sastrawi*. Beberapa tahapan pra-pemrosesan dengan memanfaatkan *library NLTK* dan *Sastrawi* pada bahasa pemrograman *Python*.

3.4. Pelabelan Data dan Pembagian Dataset

Tabel 3. Pelabelan

Preprocessed	Label Sentimen
reply ulangg tnyta awal diapus postingannya co kmren sbnrnya co barang lain akun adekku semua pake voucher video ngebantu bngt buat gprnh dpt voucher rebut makasih shopee ehehe	Positif
buka shopee perlu hp proccesor paling bagus jaring atas 100mbps ping 1 aplikasi lain tutup dulu alias berat bgt anjm	Negatif
spaylater susah tembus via vaqr dll sini coba via co shopee bisa glo bantu sampai tembus	Netral

Pada Tabel 3 diatas 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.

3.5. Ekstraksi Fitur 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.

Berdasarkan 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, dimana *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.

3.6. Naïve Bayes

Naïve 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 *Naïve Bayes* dengan memanfaatkan *library Naïve Bayes* dari *scikit-learn*. *Library* tersebut secara internal menangani perhitungan probabilitas berdasarkan representasi fitur *TF-IDF* yang telah dibentuk sebelumnya.

3.7. Support Vector Machine (SVM)

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.

3.8. 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.

3.9. FP-Growth

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:

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

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

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

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.

3.10. Crosstab

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:

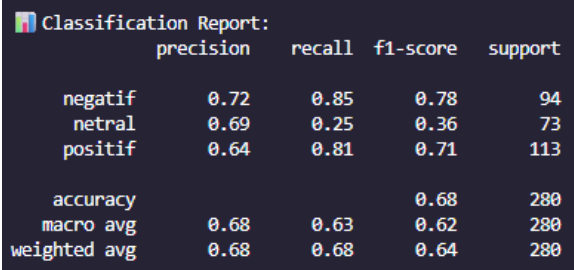
$$f_{ij} = \text{jumlah data yang termasuk kategori } i \text{ pada } j \quad (11)$$

variabel pertama dan kategori *j* pada variabel kedua 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.

4. HASIL DAN PEMBAHASAN

4.1. Evaluasi Model Naïve 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 *Naïve Bayes* dan *Support Vector Machine (SVM)* sama-sama menghasilkan nilai akurasi sebesar 68%.



	precision	recall	f1-score	support
negatif	0.72	0.85	0.78	94
netral	0.69	0.25	0.36	73
positif	0.64	0.81	0.71	113
accuracy			0.68	280
macro avg	0.68	0.63	0.62	280
weighted avg	0.68	0.68	0.64	280

Gambar 2. Classification Report Naïve Bayes

Berdasarkan *classification report*, model *Naïve Bayes* menunjukkan performa terbaik pada kelas negatif dengan nilai *recall* sebesar 0,85, yang menandakan bahwa sebagian besar tweet negatif dapat dikenali dengan baik. Kelas positif juga menunjukkan performa yang cukup baik dengan nilai *F1-score* sebesar 0,71. Sebaliknya, kelas netral memiliki performa paling rendah dengan nilai *F1-score* sebesar 0,36, yang menunjukkan bahwa model masih mengalami kesulitan dalam membedakan sentimen netral dari kelas sentimen lainnya.

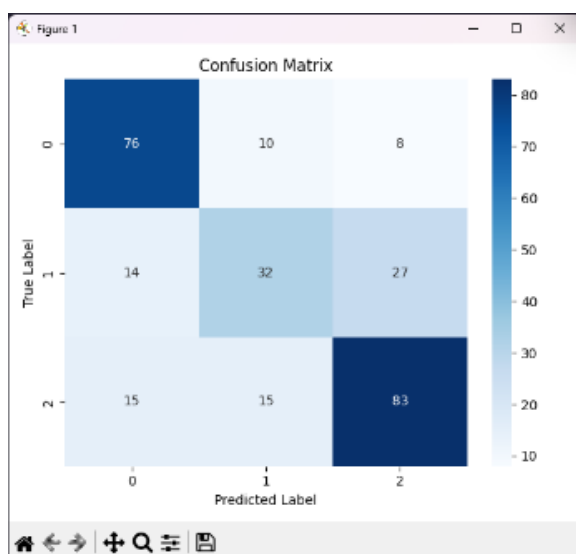


Gambar 3. Confusion Matrix Naïve Bayes

Hasil tersebut diperkuat oleh *confusion matrix* Naïve Bayes. Dari total 73 *tweet* bersentimen netral, hanya 18 *tweet* yang berhasil diklasifikasikan dengan benar sebagai netral. Sebanyak 14 *tweet* netral salah diklasifikasikan sebagai negatif, sementara 41 *tweet* lainnya diprediksi sebagai positif. Kondisi ini menunjukkan adanya kemiripan karakteristik teks antara *tweet* netral dengan *tweet* positif maupun negatif, sehingga sulit dipisahkan oleh Naïve Bayes yang bekerja dengan asumsi independensi antar kata.

	precision	recall	f1-score	support
negatif	0.72	0.81	0.76	94
netral	0.56	0.44	0.49	73
positif	0.70	0.73	0.72	113
accuracy			0.68	280
macro avg	0.66	0.66	0.66	280
weighted avg	0.67	0.68	0.67	280

Gambar 4. Classification Report SVM

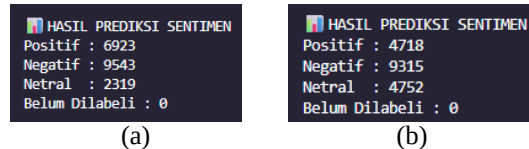


Gambar 5. Confusion Matrix SVM

Pada model SVM, classification report menunjukkan distribusi performa yang lebih merata pada setiap kelas sentimen. Kelas positif memperoleh nilai F1-score sebesar 0,72, sedangkan kelas negatif mencapai nilai F1-score sebesar 0,76. Performa kelas netral juga mengalami peningkatan dibandingkan dengan Naïve Bayes, dengan nilai F1-score sebesar 0,49, meskipun masih menjadi kelas dengan performa terendah.

Berdasarkan *confusion matrix SVM*, dari 73 data netral, sebanyak 32 *tweet* berhasil diklasifikasikan dengan benar. Sisanya masing-masing 14 *tweet* salah diprediksi sebagai negatif dan 27 *tweet* salah diprediksi sebagai positif. Jumlah kesalahan ini lebih kecil dibandingkan Naïve Bayes, yang menunjukkan bahwa SVM memiliki kemampuan pemisahan kelas yang lebih baik dalam ruang fitur *TF-IDF*, terutama dalam menangani data dengan pola yang saling tumpang tindih.

4.2. Prediksi Sentimen Data Skala Besar



Gambar 6. Naïve Bayes dan SVM

Berdasarkan hasil prediksi label sentimen yang ditunjukkan pada gambar 6 (a) dan (b), model (b) Support Vector Machine (SVM) dan (a) Naïve Bayes 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.

4.3. Klasifikasi Engagement FP-Growth

Q1 (Rendah ≤): 58.67
 Q2 (Sedang ≤): 271.34
 Kategori: Rendah, Sedang, Tinggi

Gambar 7. Kategori Engagement Total

Gambar 7 menunjukkan tingkat *engagement* yang diperoleh dari akumulasi jumlah like, retweet, reply, dan views pada setiap seluruh *tweet*. Nilai

tersebut kemudian dikelompokkan ke dalam beberapa kategori dengan batas tertentu, yaitu *engagement* rendah untuk nilai $\leq 58,67$, *engagement* sedang pada rentang 58,67 hingga 271,34, dan *engagement* tinggi untuk nilai di atas 271,34. Pengelompokan ini menghasilkan tiga kategori *engagement*, yaitu rendah, sedang, dan tinggi, sebagaimana ditampilkan pada tabel 4, di mana setiap tweet memiliki nilai numerik *engagement* serta label kategori *engagement* yang sesuai, dengan total data yang digunakan sebanyak 1400 *tweet*.

Tabel 4. Kategori Engagement

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

Tabel 5. 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.

4.4. Analisis WordCloud



Gambar 8. WordCloud Negatif Rendah

Analisis *WordCloud* pada sentimen negatif dengan *engagement* rendah sebagaimana ditunjukkan pada Gambar 8, kata-kata yang muncul didominasi oleh keluhan terkait layanan, seperti permasalahan akun, pengiriman, dan aplikasi, yang mencerminkan pengalaman pengguna yang kurang memuaskan seperti kata susah, goblok, nunggu, anjing, kalimat shopee muncul paling besar karena penelitian ini mengacu pada kata kunci utama yaitu layanan shopee.



Gambar 9. WordClous Netral Sedang

Pada sentimen netral dengan *engagement* sedang, *WordCloud* pada Gambar 9 memperlihatkan dominasi kata-kata yang bersifat informatif, seperti voucher, order, dan bayar. Pola ini menunjukkan bahwa tweet netral umumnya digunakan untuk menyampaikan informasi tanpa disertai ekspresi emosi yang kuat, kalimat shopee muncul paling besar karena penelitian ini mengacu pada kata kunci utama yaitu layanan shopee.



Gambar 10. WordCloud Positif Tinggi

Pada sentimen positif dengan *engagement* tinggi, *WordCloud* pada Gambar 10 didominasi oleh kata-kata yang berkaitan dengan promosi dan kepuasan pengguna, seperti diskon, promo, gratis ongkir, murah, dan rekomendasi produk. Visualisasi ini

menunjukkan bahwa konten positif yang mengandung unsur promosi cenderung mendapatkan perhatian yang lebih besar dari pengguna, kalimat shopee muncul paling besar karena penelitian ini mengacu pada kata kunci utama yaitu layanan shopee.

4.5. Analisis Crosstab

Analisis Crosstab dilakukan untuk melihat hubungan antara label sentimen, kategori *engagement*, dan karakteristik teks secara deskriptif. Crosstab dibentuk dengan mengelompokkan data berdasarkan sentimen dan tingkat *engagement*, kemudian menghitung frekuensi serta persentase kemunculan setiap fitur pada masing-masing kelompok.

Tabel 6. 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 7. 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 8. 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.

5. KESIMPULAN DAN SARAN

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%, dengan SVM memiliki performa yang lebih seimbang pada setiap kelas sentimen. Nilai akurasi tersebut dipengaruhi ketidakseimbangan distribusi kelas serta tumpang tindih karakteristik teks pada sentimen netral yang memiliki kemiripan dengan sentimen positif dan negatif. 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 menunjukkan bahwa *engagement* tidak hanya dipengaruhi oleh sentimen, tetapi juga oleh faktor lain di luar isi teks, seperti karakteristik akun dan kondisi eksternal. Penelitian selanjutnya dapat mempertimbangkan penggunaan variabel temporal, jaringan interaksi pengguna, serta analisis topik untuk memperoleh pemahaman yang lebih mendalam mengenai faktor-faktor yang memengaruhi *engagement*.

DAFTAR PUSTAKA

- [1] 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.
- [2] 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 Naïve Bayes Classifier," *JOINS (Journal of Information System)*, vol. 7, no. 1, pp. 14–25, May 2022, doi: 10.33633/joins.v7i1.5883.
- [3] A. Syah, F. Nurdiansyah, 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] 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.
- [5] 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>
- [6] I. Syahrohim, 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.
- [7] S. Jaffali, “Analyzing the Impact of Emotional Content in Tweets on User Engagement and Online Behavior,” 2025.
- [8] E. Saquete, J. Zubcoff, Y. Gutiérrez, P. Martínez-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.
- [9] 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.
- [10] 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>
- [11] 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>
- [12] 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.
- [13] 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.
- [14] N. Al Maeni, N. Suarna, and W. Prihartono, “ANALISIS SENTIMEN PENGGUNA SHOPEE BERDASARKAN DATA TWEET DARI TWITTER MENGGUNAKAN METODE NAIVE BAYES CLASSIFIER,” 2024.
- [15] 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>