

Klasifikasi Perilaku Konsumen Pasca Boikot Produk Israel Menggunakan Naïve Bayes dan SVM

DOI: <http://dx.doi.org/10.35889/jutisi.v15i1.3429>

Creative Commons License 4.0 (CC BY – NC)

Wildan Mauli Darajat¹, Herbert Siregar^{2*}, Rasim³, Munir⁴

Ilmu Komputer, Universitas Pendidikan Indonesia, Bandung, Indonesia

*e-mail Corresponding Author. herbert@upi.edu

Abstract

The ongoing Israel–Palestine conflict has contributed to the rise of consumer boycott movements directed at products associated with Israel. These reactions are prominently articulated on social media platforms and indicate changing patterns in consumer attitudes and behavior. This research seeks to analyze and classify public sentiment in Indonesia regarding the boycott issue by employing Natural Language Processing (NLP) techniques in combination with Machine Learning methods, specifically Naïve Bayes and Support Vector Machine (SVM). The dataset comprises user-generated comments obtained from TikTok and Instagram between October 2023 and September 2024 through web scraping procedures. The data were subsequently subjected to manual annotation, text preprocessing, and feature extraction using the TF-IDF weighting scheme. The dataset was partitioned into 80% training data and 20% testing data, and model performance was assessed using accuracy, precision, recall, and F1-score metrics. Experimental results indicate that the SVM model outperformed Naïve Bayes on the training set, achieving an accuracy of 81% and demonstrating stronger generalization in detecting positive sentiment. In contrast, the Naïve Bayes classifier attained an accuracy of 78%, showing consistent performance and superior capability in identifying negative sentiment. These results underscore the significance of selecting classification algorithms that are well suited to the distributional characteristics of sentiment data derived from social media.

Keywords: Text Classification; Support Vector Machine; Naïve Bayes; Boycott; Consumer Behavior

Abstrak

Konflik Israel–Palestina yang terus berlangsung telah mendorong munculnya gerakan boikot konsumen terhadap produk-produk yang memiliki keterkaitan dengan Israel. Respons tersebut banyak diekspresikan melalui platform media sosial dan mencerminkan perubahan pola sikap serta perilaku konsumen. Penelitian ini bertujuan untuk menganalisis dan mengklasifikasikan sentimen masyarakat Indonesia terhadap isu boikot tersebut dengan menerapkan teknik *Natural Language Processing* (NLP) yang dikombinasikan dengan metode *Machine Learning* (ML), yaitu *Naïve Bayes* dan *Support Vector Machine* (SVM). Dataset penelitian terdiri atas komentar pengguna yang dikumpulkan dari platform TikTok dan Instagram selama periode Oktober 2023 hingga September 2024 melalui teknik web scraping. Data selanjutnya melalui proses anotasi manual, praproses teks, serta ekstraksi fitur menggunakan skema pembobotan TF-IDF. Dataset dibagi menjadi 80% data latih dan 20% data uji, dengan kinerja model dievaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil eksperimen menunjukkan bahwa model SVM menghasilkan performa yang lebih unggul pada data latih dengan tingkat akurasi sebesar 81% serta memiliki kemampuan generalisasi yang lebih baik dalam mendeteksi sentimen positif. Sementara itu, algoritma *Naïve Bayes* mencapai akurasi sebesar 78% dan menunjukkan kinerja yang konsisten serta lebih efektif dalam mengidentifikasi sentimen negatif. Temuan ini menegaskan pentingnya pemilihan algoritma klasifikasi yang sesuai dengan karakteristik distribusi data sentimen yang bersumber dari media sosial.

Kata kunci: Klasifikasi Teks; Support Vector Machine; Naïve Bayes; Boikot; Perilaku Konsumen

1. Pendahuluan

Konflik Israel–Palestina merupakan isu geopolitik berkepanjangan [1]. Tidak hanya berdampak pada stabilitas kawasan Timur Tengah, tetapi juga memicu dinamika sosial, politik, dan ekonomi di tingkat global [2]. Eskalasi konflik pada periode Oktober 2023 hingga Mei 2024 khususnya setelah serangan Israel di Rafah, memunculkan gelombang solidaritas internasional yang masif [3]. Solidaritas tersebut tidak hanya diekspresikan melalui aksi demonstrasi, tetapi juga melalui gerakan boikot terhadap produk yang diasosiasikan dengan Israel sebagai bentuk tekanan moral dan politik [4]. Fenomena boikot ini menjadi penting untuk diteliti karena mencerminkan perubahan pola sikap dan perilaku masyarakat global terhadap isu kemanusiaan yang dimediasi oleh ruang digital.

Di Indonesia, dukungan terhadap Palestina terartikulasikan secara kuat melalui media sosial, yang berfungsi sebagai ruang utama penyebaran informasi, mobilisasi opini publik, serta ajakan boikot produk tertentu. Gerakan *Boycott, Divestment, and Sanctions* (BDS) mengalami peningkatan penerimaan publik dan momentum setelah diterbitkannya Fatwa No. 83 Tahun 2023 oleh Majelis Ulama Indonesia (MUI), yang menyerukan kepada umat Islam untuk menghindari konsumsi produk yang memiliki keterkaitan dengan Israel. Kondisi ini mendorong peningkatan signifikan dalam keterlibatan daring, yang tercermin dari meningkatnya jumlah komentar, opini, serta respons emosional dari masyarakat Indonesia di berbagai *platform* media sosial [5]. Namun, besarnya volume data teks yang dihasilkan membuat analisis manual terhadap sentimen dan pola perilaku konsumen menjadi tidak efisien, sehingga diperlukan pendekatan komputasional yang mampu mengolah data dalam skala besar secara sistematis dan objektif [6].

Untuk mengatasi tantangan yang ditimbulkan oleh sifat data opini publik di media sosial yang berskala besar dan tidak terstruktur, penelitian ini menerapkan pendekatan *Natural Language Processing* (NLP) yang dipadukan dengan metode *Machine Learning* (ML) [7] untuk mengklasifikasikan sentimen masyarakat Indonesia terhadap isu boikot produk Israel. NLP berperan dalam mengekstraksi fitur dari komentar netizen yang beragam, sedangkan ML berperan dalam mengidentifikasi pola sentimen pada data berskala besar [8]. Pendekatan ini selaras dengan hasil penelitian yang dilaporkan oleh Rodríguez-Ibáñez et al. [9] dan Alsemaree et al. [10] yang menyoroti efektivitas integrasi NLP dan ML dalam analisis sentimen media sosial, termasuk pada bahasa dengan kompleksitas linguistik tinggi. Hal ini diperkuat oleh survei Chandan et al. [11] yang menyatakan bahwa masih menjadi pendekatan dominan dalam analisis sentimen karena kemampuannya mentransformasikan data teks tidak terstruktur menjadi informasi terklasifikasi. Pendekatan komputasional juga telah diterapkan untuk analisis perilaku pengguna, seperti ditunjukkan oleh Aviano et al. [12] melalui penerapan *association rule mining* pada data aktivitas pengguna berskala besar. Berdasarkan pertimbangan tersebut, penelitian ini menggunakan *Naïve Bayes* karena efisiensi komputasinya, serta *Support Vector Machine* (SVM) karena kemampuannya membentuk batas keputusan yang optimal.

Penelitian ini bertujuan untuk memanfaatkan teknik NLP dan metode ML dalam mengklasifikasikan sentimen masyarakat Indonesia terkait isu boikot produk yang memiliki afiliasi dengan Israel berdasarkan komentar yang muncul di media sosial. Fokus utama penelitian ini adalah melakukan evaluasi komparatif terhadap efektivitas algoritma *Naïve Bayes* dan SVM dalam mengidentifikasi pola sentimen positif dan negatif. Hasil penelitian ini diharapkan dapat memberikan kontribusi empiris dalam menjelaskan dinamika perilaku konsumen setelah terjadinya boikot, sekaligus memperluas kajian mengenai penerapan klasifikasi teks berbasis ML pada konteks isu sosial dan politik. Selain itu, temuan yang diperoleh dapat dimanfaatkan sebagai referensi oleh kalangan akademisi, pembuat kebijakan, serta pelaku industri dalam mengoptimalkan pemanfaatan analisis sentimen digital sebagai dasar pengambilan keputusan yang berorientasi pada data.

2. Tinjauan Pustaka

Penelitian terkait analisis perilaku konsumen menunjukkan perkembangan signifikan seiring kemajuan ML dan ketersediaan data berskala besar. Chaubey et al. [13] memprediksi perilaku pembelian pelanggan menggunakan berbagai teknik klasifikasi ML dengan mengombinasikan algoritma tradisional dan *ensemble stacking*. Model KnnSgd terbukti menjadi pendekatan dengan tingkat akurasi paling tinggi, yakni mencapai 92,42%.

Aspek loyalitas pelanggan diteliti oleh Mouli et al. [14] melalui prediksi *customer churn* menggunakan sepuluh algoritma klasifikasi, termasuk *Naïve Bayes* dan SVM, dengan optimasi *hyperparameter GridSearchCV*. Variabel profil pelanggan diproses melalui tahapan *pre-*

processing yang ketat, Berdasarkan hasil pengujian, algoritma Random Forest menunjukkan kinerja paling unggul dibandingkan metode lain yang digunakan dalam penelitian tersebut.

Pitka et al. [15] menganalisis perilaku konsumen daring dari perspektif temporal dan pola asosiasi menggunakan *Decision Tree*, aturan asosiasi GUHA, dan *Formal Concept Analysis* (FCA) pada dataset transaksi jangka panjang. Hasilnya menunjukkan bahwa *Decision Tree* efektif dalam klasifikasi atribut target, GUHA unggul dalam mengidentifikasi dependensi antar-atribut, dan FCA berperan dalam memvisualisasikan hubungan karakteristik produk dan konsumen.

Pendekatan *neuromarketing* dikaji oleh Ouzir et al. [16] dengan memanfaatkan sinyal *electroencephalography* (EEG) untuk mengklasifikasikan preferensi konsumen *e-commerce*. Penelitian ini menemukan bahwa aktivitas EEG pada hemisfer kanan berkorelasi dengan respons suka, dengan *Neural Network* dan *Gradient Boosting* mencapai nilai AUC tertinggi masing-masing sebesar 76,61% dan 75,33%.

Berdasarkan tinjauan pustaka tersebut, penelitian terdahulu umumnya berfokus pada analisis perilaku konsumen berbasis data terstruktur seperti transaksi, profil pelanggan, sinyal biologis, dan sensor fisik. Berbeda dari studi-studi tersebut, penelitian ini mengusulkan pendekatan analisis sentimen berbasis NLP terhadap data teks tidak terstruktur dari media sosial untuk menangkap dimensi sikap, emosi, dan opini konsumen, sehingga menawarkan kebaruan konseptual dan metodologis dalam kajian analisis perilaku konsumen modern.

3. Metodologi

Metode penelitian yang diterapkan dalam studi ini mengacu pada pendekatan kuantitatif. memanfaatkan data berupa komentar netizen yang diperoleh dari *platform* media sosial khususnya TikTok dan Instagram sebagai sumber data utama yang akan dianalisis. Alur metodologi penelitian yang dimulai dari pengumpulan data komentar di media sosial dan pelabelan data secara manual, dilanjutkan dengan praproses untuk memastikan konsistensi data, pembobotan fitur menggunakan TF-IDF, kemudian sentimen diklasifikasikan dengan algoritma *Naïve Bayes* dan SVM. Tahap akhir berupa evaluasi model dilakukan untuk menilai kinerja klasifikasi berdasarkan metrik evaluasi.

3.1 Pengumpulan Data

Data penelitian diperoleh dari TikTok dan Instagram, fokus pada komentar terkait boikot produk Israel. Pengumpulan dilakukan melalui *scraping* untuk mengekstraksi data secara terstruktur yang kemudian disimpan dalam format Excel sebagai dataset utama.

3.2 Pelabelan Data

Tahap pelabelan dilakukan secara manual setelah data terkumpul atau masih berupa data mentah untuk memudahkan kategori klasifikasi sehingga komentar yang bersifat sarkas dapat dideteksi lebih baik. Komentar diklasifikasikan secara manual menjadi dua kategori, yang terdiri atas sentimen positif (mendukung boikot) dan sentimen negatif (menentang boikot).

3.3 Praproses Data

Praproses data bertujuan menghilangkan komponen komentar yang tidak berkaitan dengan kebutuhan analisis [17]. Tahapan praproses teks meliputi normalisasi huruf (*case folding*), pembersihan data teks (*text cleaning*), pemisahan kata (*tokenization*), penghapusan kata tidak bermakna (*stopword removal*), serta proses reduksi kata ke bentuk dasarnya (*stemming*), yang diakhiri dengan eliminasi teks yang tidak relevan.

3.4 Ekstraksi Fitur Berbasis TF-IDF

Pada tahap ini, data teks yang telah melalui proses praproses dikonversi ke dalam bentuk numerik melalui penerapan teknik TF-IDF. Metode ini memberikan bobot numerik pada setiap kata dengan mempertimbangkan frekuensi kemunculannya serta tingkat kepentingannya relatif terhadap dokumen. Hasil dari proses ini berupa matriks fitur numerik yang selanjutnya digunakan sebagai input dalam proses klasifikasi sentimen menggunakan algoritma *Naïve Bayes* dan SVM.

$$TF_{IDF(w,s,D)} = tf(w,s) * \log\left(\frac{N}{dfw}\right) \quad (1)$$

Rumus tersebut digunakan untuk memBobot kata berdasarkan kemunculannya dalam dokumen relatif terhadap keseluruhan dataset. Perhitungan dilakukan dengan mengalikan frekuensi kata dalam dokumen $tf(w, s)$ dengan logaritma dari rasio total dokumen (N) terhadap jumlah dokumen yang memuat kata tersebut (dfw). Suatu kata yang memiliki frekuensi kemunculan tinggi namun jarang muncul pada dokumen lain akan diberikan bobot yang lebih besar [19].

3.5 Klasifikasi Naive Bayes dan SVM

Dalam penelitian ini, klasifikasi perilaku konsumen dibuat dengan menggunakan *Naive Bayes* dan SVM. Setelah penerapan representasi fitur berbasis TF-IDF. Dataset dibagi menjadi 80% data latih dan 20% data uji guna menilai kemampuan model dalam melakukan generalisasi terhadap data yang belum pernah dilihat sebelumnya. Hasil prediksi yang diperoleh dari data uji digunakan sebagai dasar evaluasi kinerja serta analisis perbandingan kedua algoritma. Algoritma *Naive Bayes* merupakan pendekatan klasifikasi probabilistik yang diturunkan dari *Teorema Bayes* dan didasarkan pada asumsi independensi bersyarat antar fitur, sehingga memungkinkan proses pembelajaran yang sederhana serta efisien secara komputasi.

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

Naive Bayes menentukan probabilistik posterior suatu kelas dengan memanfaatkan informasi probabilistik data prediktor, $P(C | X)$, dengan memanfaatkan probabilistik prior kelas $P(C)$, probabilistik prediktor terhadap kelas $P(X | C)$, serta probabilistik prior prediktor $P(X)$ [21].

Pada SVM dengan kernel linear, proses klasifikasi dilakukan dengan membentuk sebuah *hyperplane* optimal yang berfungsi sebagai batas pemisah dengan dua kelas utama yang merepresentasikan sentimen positif dan sentimen negatif pada ruang fitur.

$$\begin{aligned} x_i \cdot w + b &= 0 \text{ (hyperplane)} \\ x_i \cdot w + b &\geq +1, y_i = +1 \text{ (kelas positif)} \\ x_i \cdot w + b &\leq -1, y_i = -1 \text{ (kelas negatif)} \end{aligned} \quad (3)$$

Pada persamaan tersebut, x_i merepresentasikan vektor fitur, w adalah vektor bobot, dan b merupakan bias. Data diklasifikasikan sebagai kelas positif ($y_i = +1$) jika $x_i \cdot w + b \geq +1$, dan sebagai kelas negatif ($y_i = -1$) jika $x_i \cdot w + b \leq -1$, sehingga SVM membentuk batas pemisah optimal untuk klasifikasi [22].

3.6 Evaluasi Model

Evaluasi efektivitas model dilakukan dengan membandingkan hasil prediksi model terhadap data uji yang telah diberi label secara manual menggunakan *confusion matrix*. Alat evaluasi ini mengukur hasil klasifikasi berdasarkan empat kondisi utama, yaitu *True Positive* (TP) yang menunjukkan data positif terklasifikasi dengan benar, *True Negative* (TN) yang merepresentasikan data negatif terprediksi secara tepat, *False Positive* (FP) ketika data negatif keliru diprediksi sebagai positif, serta *False Negative* (FN) ketika data positif salah diklasifikasikan sebagai negatif [23].

Nilai *accuracy* mencerminkan tingkat ketepatan prediksi model secara keseluruhan dan dihitung sebagai perbandingan antara jumlah prediksi yang benar dengan total data uji [24].

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

Precision menunjukkan tingkat akurasi prediksi pada suatu kelas sentimen berdasarkan rasio prediksi benar terhadap total prediksi kelas tersebut [25].

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

Recall menggambarkan kemampuan model dalam mengidentifikasi seluruh data yang sebenarnya termasuk ke dalam kelas positif [26].

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

F1-Score menggabungkan *precision* dan *recall* melalui rata rata harmonik, sehingga memberi bobot yang seimbang pada kedua metrik tersebut untuk menilai keseluruhan performa model [27].

$$F1 - Score = 2 \frac{precision \cdot recall}{precision + recall} \quad (7)$$

4. Hasil dan Pembahasan

Bab ini menguraikan temuan penelitian dan pembahasannya berdasarkan data yang telah dianalisis terkait proses klasifikasi sentimen konsumen pasca boikot produk Israel menggunakan algoritma *Naïve Bayes* dan SVM. Pembahasan mencakup temuan-temuan utama dalam studi, analisis kinerja model berdasarkan metrik evaluasi, serta interpretasi pola sentimen yang diperoleh dari data komentar media sosial.

4.1 Pengumpulan Data

Data penelitian diperoleh melalui pelaksanaan tahapan pengambilan data untuk memperoleh komentar yang relevan dengan topik “Boikot Produk Israel” pada periode meningkatnya konflik Palestina–Israel, yakni dari bulan Oktober 2023 hingga September 2024. Di bawah ini menampilkan gambar hasil dari *scraping data* menggunakan *Apify*.

createTime	createTimeISO	detailedMention	diggCount	input	likedBy	Apinned	replyCorr	submittedVideoUrl	text	uid	uniqueId	videoWebUrl
1,753E+09	2025-07-18T06:56:01.000Z	2	https://www	false	false	0	https://www.tiktok.com/@maaf dari pintu nya ga	750797250761	www.dokter.oc	https://www.t		
1,717E+09	2024-06-03T21:38:24.000Z	4085	https://www	false	false	40	https://www.tiktok.com/@Aku gak pernah Starbuc	723930339574	bintang0994	https://www.t		
1,717E+09	2024-06-04T00:38:34.000Z	5538	https://www	false	false	52	https://www.tiktok.com/@kalian tau gak sih kala	695954370570	rafiasharij	https://www.t		
1,717E+09	2024-06-04T06:01:21.000Z	877	https://www	false	false	60	https://www.tiktok.com/@demo Starbuck, tp me	678236877132	salamrayaketaj	https://www.t		
1,717E+09	2024-06-03T20:15:45.000Z	306	https://www	false	false	12	https://www.tiktok.com/@Otw beli starbuck	685548968124	kentangmekdi	https://www.t		
1,717E+09	2024-06-03T21:04:05.000Z	492	https://www	false	false	12	https://www.tiktok.com/@sinti**	709280100666	ayuswairata	https://www.t		
1,718E+09	2024-06-04T23:13:35.000Z	20	https://www	false	false	0	https://www.tiktok.com/@ @ liena :klo mau boiki	689518276461	peiraa	https://www.t		
1,718E+09	2024-06-04T14:52:35.000Z	744	https://www	false	false	18	https://www.tiktok.com/@ alhamdulillah berkat	696323217238	faisalalmark	https://www.t		
1,717E+09	2024-06-04T01:07:50.000Z	551	https://www	false	false	17	https://www.tiktok.com/@yokkk luurr ramaikan	695994515131	the_buj4nkbro	https://www.t		

Gambar 1. Hasil *Scraping Data*

Data penelitian dikumpulkan melalui proses *scraping* menggunakan layanan *Apify* dengan menargetkan komentar pada platform TikTok dan Instagram. Dari proses tersebut diperoleh sebanyak 25.092 komentar yang digunakan sebagai dataset utama penelitian.

4.2 Pelabelan Data

Pelabelan data dilakukan sebelum praproses untuk mempertahankan makna asli komentar. Komentar dari TikTok dan Instagram dianalisis manual dan pelabelan data dilakukan secara sistematis ke dalam dua perilaku yang berbeda, yang terdiri atas sikap mendukung boikot dan sikap menentangnya. Pelabelan dilakukan dengan menilai setiap komentar agar struktur kalimat saat praproses tidak memengaruhi penilaian. Contoh pelabelan disajikan pada Tabel 1.

Tabel 1. Pelabelan Data

Komentar	Pelabelan
Ya Allah ternyata ini urusan serius nti di Akherat 🤔 ya	Positif
boikotnya pilih" 😊 kalau butuh dipake 😊	Negatif

4.3 Praproses Data

Tahapan awal untuk pengolahan data merupakan awal proses pengolahan data yang bertujuan untuk menyiapkan sekumpulan data yang akurat dan terintegrasi [28]. Tahapan praproses dalam penelitian ini terdiri atas beberapa langkah seperti yang akan dijelaskan dibawah ini.

4.3.1 Menghapus Kolom

Dataset dibersihkan dengan menghapus atribut yang tidak relevan, seperti nama pengguna, jumlah tanda suka, dan kolom pendukung lainnya, sehingga analisis difokuskan pada kolom text yang kemudian diubah namanya menjadi kolom komentar.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 27102 entries, 0 to 27101
Data columns (total 1 columns):
#   Column      Non-Null Count  Dtype
---  ---
0   komentar    25092 non-null  object
dtypes: object(1)
memory usage: 211.9+ KB
```

Gambar 2. Informasi *DataFrame* Setelah Menghapus Kolom

Gambar 2 menunjukkan ringkasan struktur dataset dalam bentuk *DataFrame* pandas yang terdiri dari 27.102 entri dengan satu kolom utama, yaitu komentar. Dari seluruh entri yang tersedia, sebanyak 25.092 data memiliki nilai non-null, yang menandakan bahwa terdapat sebagian data yang kosong dan perlu di hapus melalui tahapan *text elimination* di akhir tahapan pra-proses.

4.3.2 Case Folding

Proses ini bertujuan mengonversi seluruh huruf pada komentar ke dalam bentuk huruf kecil [29]. Tujuan dari proses ini adalah untuk memastikan konsistensi dan keseragaman penulisan teks.

Tabel 2. Hasil Penerapan *Case Folding*

Komentar	Case Folding
Ya Allah ternyata ini urusan srius nti di Akherat 🤔 ya	ya allah ternyata ini urusan srius nti di akherat 🤔 ya

4.3.3 Text Cleaning

Proses pembersihan data secara otomatis untuk memastikan kualitas dataset agar siap dianalisis. Proses ini mencakup penghapusan elemen-elemen yang tidak relevan [30]. Guna memberikan ilustrasi yang lebih jelas terkait hasil pembersihan data, disajikan Tabel 3 yang menampilkan perbandingan isi komentar sebelum dan sesudah proses pembersihan.

Tabel 3. *Text Cleaning*

Kondisi	Komentar	Cleaning
Menghapus baris baru \n	temen ku yg dari pakistan katanya perminta maafan mcd sana juga ga sebut nama palestina. \nkayaknya emang perminta maafannya udah di atur dari pusatnya	temen ku yg dari pakistan katanya perminta maafan mcd sana juga ga sebut nama palestina kayaknya emang perminta maafannya udah di atur dari pusatnya
Menghapus <i>mention</i>	Bayangin lagi jalan ke McD, terus di luar ada beginian @rosie	bayangin lagi jalan ke mcd terus di luar ada beginian
Menghapus URL	Boikot krn ga mampu beli, KALAU GRATIS YA BEREPUT \n 🤔 https://vt.tiktok.com/ZSrvyRAKx/ [foto]	boikot krn ga mampu beli kalau gratis ya berebut foto
Menghapus <i>hashtag</i>	#freepalestine masih banyak produk lokal yang lebih enak dan ramah di saku	masih banyak produk lokal yang lebih enak dan ramah di saku
Menghapus emoji	Ya Allah ternyata ini urusan srius nti di Akherat 🤔 ya	ya allah ternyata ini urusan srius nti di akherat ya
Menghapus karakter tersembunyi	Insha Allah ..Istiqomah boikot.\n 🤔 PS ❤️ 🤔 PS 🤔	insha allah istiqomah boikot
Menghapus tanda baca	balik ke UMKM ajah sih, kayak martabak, pecel, soto, rawon, \nroti kukus, es cendol, es dawet, es buah, es campur..msh byk kok makanan yg bnr2 asli Ind	balik ke umkm ajah sih kayak martabak pecel soto rawon roti kukus es cendol es dawet es buah es campur msh byk kok makanan yg bnr asli ind
Merapikan spasi	Postingan ga mendidik.... pikirin orang2 yg kerja di lokasi yg diboikot. Itu orang Indonesia juga!!!	postingan ga mendidik pikirin orang2 yg kerja di lokasi yg diboikot itu orang indonesia juga

4.3.4 Normalization

Proses normalisasi bertujuan untuk mengonversi bentuk leksikal yang tidak baku, seperti singkatan dan bahasa informal, ke dalam istilah baku yang sesuai [31].

Tabel 4. Normalization

Cleaning	Normalisasi
ya allah ternyata ini urusan srius nti di akherat ya	ya allah ternyata ini urusan serius nanti di akhirat ya

4.3.5 Tokenizing

Pada tahap *tokenizing*, teks komentar diuraikan menjadi kata-kata individual atau token untuk memudahkan analisis [32].

Tabel 5. Tokenizing

Normalisasi	Tokenizing
ya allah ternyata ini urusan serius nanti di akhirat ya	['ya', 'allah', 'ternyata', 'ini', 'urusan', 'serius', 'nanti', 'di', 'akhirat', 'ya']

4.3.6 Stopword Removal

Stopword removal bertujuan menghilangkan kata yang tidak relevan dari data teks.. Bertujuan mengurangi *noise* pada kata-kata yang tidak relevan [33].

Tabel 6. Stopword

Tokenizing	Stopword
['ya', 'allah', 'ternyata', 'ini', 'urusan', 'serius', 'nanti', 'di', 'akhirat', 'ya']	['allah', 'urusan', 'serius', 'akhirat']

4.3.7 Stemming

Stemming dilakukan dengan tujuan menyederhanakan kata menjadi bentuk dasar agar perbedaan variasi kata dapat dikurangi dan membantu model menangkap makna inti teks secara konsisten. Hasil *stemming* kemudian disusun kembali dalam bentuk kalimat agar dapat diproses pada tahap pembobotan TF-IDF [34].

Tabel 7. Stemming

Stopword	Stemming
['allah', 'urusan', 'serius', 'akhirat']	allah urus serius akhirat

4.3.8 Text Elimination

Pada tahap ini dilakukan proses penghapusan teks, yaitu komentar kosong dan komentar duplikat. Komentar kosong dihapus karena tidak mengandung informasi yang dapat digunakan untuk proses klasifikasi sentimen. Sementara itu, komentar duplikat dibersihkan untuk mencegah terjadinya bias pada model akibat adanya pengulangan data yang sama.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 19209 entries, 0 to 19208
Data columns (total 2 columns):
#   Column      Non-Null Count  Dtype
---  ---
0   stemming    19209 non-null  object
1   pelabelan   19209 non-null  object
dtypes: object(2)
memory usage: 300.3+ KB
```

Gambar 3. Informasi DataFrame Setelah Praproses

Berdasarkan ringkasan DataFrame pada Gambar 3, dataset akhir terdiri dari 19.209 entri valid dengan dua variabel utama, yaitu kolom *stemming* sebagai representasi teks hasil praproses dan kolom pelabelan sebagai target klasifikasi perilaku konsumen.

4.4 Representasi Fitur TF-IDF

Teknik TF-IDF digunakan untuk mengonversi komentar menjadi representasi angka dengan menangkap tingkat kepentingan relatif setiap kata dalam dokumen. Kata-kata yang memiliki frekuensi tinggi pada dokumen tertentu namun jarang ditemukan pada dokumen lain akan diberikan bobot yang lebih besar, sehingga memungkinkan model untuk lebih efektif membedakan istilah-istilah yang bermakna dalam interpretasi sentimen [35].

	allah	urus	serius	akhirat	keren	abang	kurang	teman	usaha	mikro
0	0.922858	1.2216	1.442197	1.477973	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
1	0.000000	0.0000	0.000000	0.000000	2.753213	2.65992	0.000000	0.000000	0.000000	0.000000
2	0.000000	0.0000	0.000000	0.000000	0.000000	0.000000	1.736391	3.345965	0.000000	0.000000
3	0.000000	0.0000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.154161	0.224393

Gambar 4. Hasil Pembobotan TF-IDF

Gambar 4 menampilkan matriks fitur numerik hasil pembobotan TF-IDF, di mana setiap baris merepresentasikan satu komentar dari total 19.209 data yang telah dibersihkan, sedangkan setiap kolom merepresentasikan kata kunci unik hasil proses *stemming*. Nilai numerik pada setiap sel menunjukkan bobot kepentingan kata dalam suatu komentar, di mana nilai misalkan 1.442197 (serius) merefleksikan tingkat signifikansi kata berdasarkan frekuensi kemunculan dan kelangkaannya di seluruh korpus, sementara nilai nol menunjukkan bahwa kata tersebut tidak muncul pada komentar terkait.

4.5 Proses Training dan Evaluasi Model Klasifikasi

Tahap klasifikasi diawali dengan analisis distribusi kategori sentimen untuk memastikan keseimbangan data. Hasil visualisasi pada Gambar 6 menunjukkan bahwa dataset terdiri atas 9.373 komentar (48,79%) dengan sentimen positif dan 9.836 komentar (51,21%) dengan sentimen negatif. Distribusi yang relatif seimbang ini menunjukkan tidak adanya ketimpangan kelas yang signifikan, sehingga tidak perlu penyeimbangan data.

4.5.1 Training Model Klasifikasi Menggunakan Data Latih

Data dibagi ke dalam data latih dan data uji dengan proporsi masing-masing sebesar 80% dan 20%. Pembagian ini menunjukkan bahwa distribusi sentimen positif dan negatif pada kedua subset relatif seimbang, dengan proporsi sentimen negatif yang sedikit lebih tinggi dibandingkan sentimen positif. Persentase masing-masing sebesar 51,21% dan 48,79%. Kondisi ini menandakan tidak adanya ketimpangan kelas yang signifikan, sehingga proses pelatihan dan pengujian model dapat dilakukan tanpa penerapan teknik penyeimbangan data tambahan serta memungkinkan evaluasi kinerja model yang lebih objektif.

4.5.2 Sampel Hasil Klasifikasi pada Data Latih

Data latih (80%) yang telah digunakan pada proses *training* selanjutnya menghasilkan prediksi label sentimen oleh algoritma *Naïve Bayes* dan SVM. Tabel 8 menyajikan sampel hasil klasifikasi pada data latih yang menampilkan perbandingan antara label manual dan label prediksi dari kedua algoritma tersebut.

Tabel 8. Sampel Hasil Klasifikasi Menggunakan Data Latih

Stemming	Label Manual	Prediksi NB	Prediksi SVM
keluh sanggup beli coba ada gratis merek starbucks antre	0	0	0
serah orang produk dukung israel enak	0	1	0
alhamdulillah detik tahan ayam goreng krispi jual tetangga enak	1	1	1
ikut majelis ulama indonesia buat baik manusia muka bumi pegang	1	0	0
haram kecap bango royco pakai bumbu masako ajinomoto	1	0	0

Tabel 8 menampilkan sampel hasil klasifikasi menggunakan data latih (80%) dengan membandingkan label manual dan hasil prediksi kedua algoritma berdasarkan pada kolom

stemming yang merupakan komentar bersih setelah melewati seluruh tahapan praproses. Secara umum, kedua model mampu mengklasifikasikan sentimen sesuai dengan label manual.

4.5.3 Evaluasi Performa Model pada Data Latih

Setelah disajikan sampel hasil klasifikasi pada data latih, langkah berikutnya penilaian kinerja model menggunakan data latih (80%). Evaluasi ini bertujuan untuk mengetahui seberapa efektif algoritma Naïve Bayes dan SVM dalam mengklasifikasikan sentimen dengan memanfaatkan pola yang diperoleh selama proses pelatihan. Penilaian kinerja model dilakukan menggunakan metrik evaluasi standar, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*.

Naive Bayes Training Accuracy : 0.7879872453959784				
SVM Training Accuracy : 0.8170755515064749				
===== NAÏVE BAYES =====				
	precision	recall	f1-score	support
Negatif (0)	0.7781	0.8197	0.7984	7869
Positif (1)	0.7995	0.7547	0.7765	7498
accuracy			0.7880	15367
macro avg	0.7888	0.7872	0.7874	15367
weighted avg	0.7886	0.7880	0.7877	15367
===== SVM =====				
	precision	recall	f1-score	support
Negatif (0)	0.8119	0.8366	0.8241	7869
Positif (1)	0.8228	0.7966	0.8095	7498
accuracy			0.8171	15367
macro avg	0.8174	0.8166	0.8168	15367
weighted avg	0.8172	0.8171	0.8170	15367

Gambar 5. Evaluasi Model pada Data Latih

Hasil evaluasi menunjukkan SVM memiliki kinerja yang lebih baik dibandingkan Naïve Bayes, sebagaimana ditunjukkan oleh nilai akurasi yang lebih tinggi sebesar 81,71% dan lebih tinggi dibandingkan akurasi Naïve Bayes sebesar 78,80%, sehingga menandakan bahwa SVM lebih optimal dalam menangkap pola sentimen pada data latih.

Berdasarkan matriks keseluruhan, algoritma SVM secara konsisten lebih baik pada kedua kelas sentimen. Pada kelas negatif, SVM mencapai *F1-score* (0,8241), lebih tinggi dibandingkan Naïve Bayes (0,7984), yang menandakan kemampuan pemisahan kelas yang lebih kuat. Sementara itu, pada kelas positif, SVM memperoleh *F1-score* (0,8095), juga melampaui Naïve Bayes (0,7765) yang menegaskan keunggulan SVM dalam mengenali variasi sentimen positif pada data latih.

4.6 Pengujian dan Evaluasi Model Klasifikasi

Pada tahapan ini membahas proses pengujian dan evaluasi model klasifikasi menggunakan data uji (20%). Bertujuan untuk menilai kemampuan generalisasi model Naïve Bayes dan SVM dalam mengklasifikasikan sentimen pada data baru serta mengevaluasi kinerja model secara objektif.

4.6.1 Pengujian Model Menggunakan Data Uji

Pengujian model menggunakan data uji (20%) menghasilkan prediksi label sentimen terhadap data yang tidak pernah digunakan pada tahap *training*. Hasil prediksi tersebut menjadi output utama dari tahap pengujian dan selanjutnya digunakan sebagai dasar penyajian sampel hasil klasifikasi serta evaluasi performa model pada tahapan selanjutnya.

4.6.2 Sampel Hasil Klasifikasi pada Data Uji

Hasil dari proses pengujian model berupa prediksi label sentimen pada data uji selanjutnya disajikan dalam bentuk sampel hasil klasifikasi. Tabel 9 menampilkan beberapa contoh hasil klasifikasi pada data uji yang memperlihatkan perbandingan antara label manual dan

label prediksi yang dihasilkan oleh algoritma *Naïve Bayes* dan SVM sebelum dilakukan evaluasi performa.

Tabel 9. Sampel Hasil Klasifikasi Menggunakan Data Uji

Stemming	Label Manual	Prediksi NB	Prediksi SVM
fungsi boikot fomo	0	0	0
inti beli barang habis ganti	1	1	1
tolong produk israel	0	1	1
lapor kuartal usaha turun jual	1	0	1
unilever beli	0	1	1

Tabel 9 menyajikan sampel hasil klasifikasi menggunakan data uji (20%) yang menampilkan perbandingan antara label manual dan label prediksi yang dihasilkan oleh algoritma *Naïve Bayes* dan SVM. Sampel ini menunjukkan bagaimana kedua model mengklasifikasikan sekaligus memperlihatkan adanya kesesuaian dan perbedaan hasil prediksi pada beberapa kasus.

4.6.3 Evaluasi Performa Model pada Data Uji

Tahap akhir adalah evaluasi performa model berdasarkan hasil pengujian pada data uji (20%). Penilaian kinerja model dilakukan dengan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score* sebagai alat evaluasi untuk mengkaji kualitas hasil prediksi serta melakukan analisis perbandingan antara *Naïve Bayes* dan SVM dalam klasifikasi sentimen secara objektif.

Naive Bayes Testing Accuracy : 0.7378969286829776				
SVM Testing Accuracy : 0.7381572097865695				
===== NAÏVE BAYES =====				
	precision	recall	f1-score	support
Negatif (0)	0.7299	0.7748	0.7517	1967
Positif (1)	0.7474	0.6992	0.7225	1875
accuracy			0.7379	3842
macro avg	0.7387	0.7370	0.7371	3842
weighted avg	0.7384	0.7379	0.7374	3842
===== SVM =====				
	precision	recall	f1-score	support
Negatif (0)	0.7364	0.7611	0.7485	1967
Positif (1)	0.7402	0.7141	0.7269	1875
accuracy			0.7382	3842
macro avg	0.7383	0.7376	0.7377	3842
weighted avg	0.7382	0.7382	0.7380	3842

Gambar 6. Evaluasi Model pada Data Uji

Hasil evaluasi performa pada data uji (20%) menunjukkan bahwa algoritma *Naïve Bayes* dan SVM memiliki kinerja yang relatif seimbang dalam mengklasifikasikan sentimen. Nilai *accuracy* yang diperoleh *Naïve Bayes* sebesar 73,79%, sedangkan SVM mencapai 73,82%, yang menunjukkan bahwa kedua model mampu melakukan generalisasi dengan tingkat akurasi yang hampir sama pada data yang tidak digunakan selama proses training.

Evaluasi menggunakan keseluruhan metrik memperlihatkan bahwa *Naïve Bayes* mempunyai nilai *precision* sedikit lebih baik pada kelas positif (0,7474) dibandingkan SVM (0,7402), yang membuktikan *Naïve Bayes* relatif selektif untuk memprediksi nilai positif. Namun, SVM memiliki nilai *recall* kelas positif yang lebih tinggi (0,7141) melampaui *Naïve Bayes* (0,6992), yang menegaskan keunggulan SVM dalam mengidentifikasi data sentimen positif secara menyeluruh.

4.7 Pembahasan

Berdasarkan hasil evaluasi, SVM memiliki performa lebih unggul pada data latih (80%) dibandingkan *Naïve Bayes* berdasarkan seluruh metrik evaluasi yang menandakan kemampuannya dalam mempelajari pola sentimen berbasis TF-IDF secara lebih optimal. Pada data uji (20%), kedua algoritma mengalami penurunan performa yang wajar dengan tingkat akurasi yang relatif seimbang, namun SVM memiliki *recall* yang lebih tinggi pada kelas positif, sementara *Naïve Bayes* lebih efektif dalam mengenali kelas negatif, sehingga mencerminkan perbedaan karakteristik generalisasi kedua model.

Temuan ini sejalan dengan berbagai penelitian terdahulu yang melaporkan keunggulan SVM dibandingkan *Naïve Bayes* dalam klasifikasi teks. Studi oleh Tabany dan Gueffal [36] menunjukkan bahwa SVM unggul dalam analisis sentimen pada ulasan Amazon, dengan akurasi hingga 93% setelah *hyperparameter tuning*, melebihi *Naïve Bayes* dan algoritma lain. Luo [37] juga melaporkan bahwa SVM *outperform classifier* lain pada klasifikasi teks bahasa Inggris, dengan tingkat klasifikasi melebihi 90% saat jumlah fitur memadai. Selanjutnya, penelitian Elgeldawi et al. [38] menegaskan efektivitas SVM dalam analisis sentimen teks Arab, di mana SVM mencapai akurasi tertinggi 95,62% setelah tuning menggunakan *Bayesian Optimization*. Hasil-hasil ini menegaskan bahwa performa algoritma sangat dipengaruhi oleh optimasi parameter dan karakteristik dataset, sekaligus menunjukkan keunggulan SVM dibandingkan *Naïve Bayes* pada berbagai bahasa dan jenis teks.

Kontribusi utama dari penelitian ini ditunjukkan melalui analisis komparatif kinerja *Naïve Bayes* dan SVM dalam melakukan klasifikasi sentimen pada isu boikot produk, dengan memanfaatkan metrik evaluasi berupa *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil pengujian mengindikasikan bahwa algoritma SVM memiliki kinerja yang lebih stabil dan konsisten, khususnya dalam mengidentifikasi kelas positif sebagai kelas minoritas, sedangkan *Naïve Bayes* cenderung lebih efektif dalam mengenali kelas negatif sebagai kelas mayoritas. Temuan ini menegaskan bahwa pemilihan algoritma klasifikasi sentimen perlu mempertimbangkan karakteristik distribusi data, serta memperkaya pemahaman empiris mengenai penerapan metode *machine learning* pada analisis sentimen isu sosial di media sosial.

5. Simpulan

Berdasarkan hasil evaluasi, dapat disimpulkan bahwa model klasifikasi menggunakan SVM dan *Naïve Bayes* sama-sama mampu memberikan prediksi terhadap data sentimen dengan performa yang cukup baik, namun masing-masing menunjukkan kekuatan yang berbeda. Secara keseluruhan, kedua algoritma menunjukkan performa yang stabil. Akan tetapi, *Naïve Bayes* cenderung lebih fokus pada kelas mayoritas (negatif), sedangkan SVM lebih efektif dalam menangkap pola pada kelas minoritas (positif). Selain itu, hasil analisis sentimen memperlihatkan bahwa kecenderungan opini publik sedikit condong ke kontra-boikot, karena porsi sentimen negatif yang menolak boikot berada sedikit di atas sentimen positif yang mendukung.

Daftar Referensi

- [1] P. Goyal dan P. Soni, "Beyond borders: investigating the impact of the 2023 Israeli–Palestinian conflict on global equity markets," *Journal of Economic Studies*, vol. 51, no. 8, pp. 1714–1731, 2024.
- [2] W. Jia dan A. A. M. Hassan, "Navigating the crossroads: Analyzing China's Middle East policy in the Palestinian-Israeli conflict context," *Contemporary Review of the Middle East*, vol. 12, no. 1, pp. 10–28, 2025.
- [3] F. R. Widyrianto, "Upaya Diplomasi Indonesia dalam merespon Agresi Israel di Gaza Pada Tahun 2023-2024," Universitas Islam Indonesia, 2025.
- [4] H. F. Aked, *Friends of Israel: The backlash against Palestine solidarity*. Verso Books, 2023.
- [5] Y. A. Adi, A. F. Adi, dan M. T. Bahri, "# Boycott as a Social Movement: Evidence from the X Platform.," *Observatorio (OBS*)*, vol. 19, no. 1, pp. 10–28 2025.
- [6] K. Du, F. Xing, R. Mao, dan E. Cambria, "Financial sentiment analysis: Techniques and applications," *ACM Comput. Surv.*, vol. 56, no. 9, pp. 1–42, 2024.
- [7] Y. Liu, F. Han, L. Meng, J. Lai, dan X. Gao, "Textual sentiment classification in tourism research: between manual computing model and machine learning," *Current Issues in Tourism*, pp. 1–22, 2025.
- [8] L. S. Riza dkk., "Comparison of Machine Learning Algorithms for Species Family Classification using DNA Barcode.," *Knowl. Eng. Data Sci.*, vol. 6, no. 2, pp. 231-242, 2023.

- [9] M. Rodriguez-Ibáñez, A. Casáñez-Ventura, F. Castejón-Mateos, dan P.-M. Cuenca-Jiménez, "A review on sentiment analysis from social media platforms," *Expert Syst. Appl.*, vol. 223, pp. 119-131, 2023.
- [10] O. Alsemaree, A. S. Alam, S. S. Gill, dan S. Uhlig, "Sentiment analysis of Arabic social media texts: A machine learning approach to deciphering customer perceptions," *Heliyon*, vol. 10, no. 9, pp. 156-168, 2024.
- [11] M. K. Chandan dan S. Mandal, "A comprehensive survey on sentiment analysis: Framework, techniques, and applications," *Comput. Sci. Rev.*, vol. 58, pp. 100-117, 2025.
- [12] D. Aviano, B. L. Putro, E. P. Nugroho, dan H. Siregar, "Behavioral tracking analysis on learning management system with apriori association rules algorithm," dalam *2017 3rd International Conference on Science in Information Technology (ICSITech)*, pp. 372–377, 2017.
- [13] G. Chaubey, P. R. Gavhane, D. Bisen, dan S. K. Arjaria, "Customer purchasing behavior prediction using machine learning classification techniques," *J. Ambient Intell. Humaniz. Comput.*, vol. 14, no. 12, pp. 16133–16157, 2023.
- [14] K. C. Mouli dkk., "An analysis on classification models for customer churn prediction," *Cogent Eng.*, vol. 11, no. 1, pp. 278-298, 2024.
- [15] T. Pitka dkk., "Time analysis of online consumer behavior by decision trees, GUHA association rules, and formal concept analysis," *Journal of Marketing Analytics*, vol. 13, no. 1, pp. 29–52, 2025.
- [16] M. Ouzir, H. C. Lamrani, R. L. Bradley, dan I. El Moudden, "Neuromarketing and decision-making: Classification of consumer preferences based on changes analysis in the EEG signal of brain regions," *Biomed. Signal Process. Control*, vol. 87, pp. 105-469, 2024.
- [17] T. A. Alghamdi dan N. Javaid, "A survey of preprocessing methods used for analysis of big data originated from smart grids," *Ieee Access*, vol. 10, pp. 29149–29171, 2022.
- [18] H. Liu, X. Chen, dan X. Liu, "A study of the application of weight distributing method combining sentiment dictionary and TF-IDF for text sentiment analysis," *Ieee Access*, vol. 10, pp. 32280–32289, 2022.
- [19] F. Shehzad, A. Rehman, K. Javed, K. A. Alnowibet, H. A. Babri, dan H. T. Rauf, "Binned term count: An alternative to term frequency for text categorization," *Mathematics*, vol. 10, no. 21, pp. 41-64, 2022.
- [20] K. L. Tan, C. P. Lee, K. M. Lim, dan K. S. M. Anbananthen, "Sentiment analysis with ensemble hybrid deep learning model," *IEEE Access*, vol. 10, pp. 103694–103704, 2022.
- [21] M. Ilić, Z. Srdjević, dan B. Srdjević, "Water quality prediction based on Naive Bayes algorithm," *Water Science and Technology*, vol. 85, no. 4, pp. 1027–1039, 2022.
- [22] V. Piccialli dan M. Sciandrone, "Nonlinear optimization and support vector machines," *Ann. Oper. Res.*, vol. 314, no. 1, pp. 15–47, 2022.
- [23] A. Tharwat, "Classification assessment methods," *Applied computing and informatics*, vol. 17, no. 1, pp. 168–192, 2021.
- [24] U. Norinder dan P. Norinder, "Predicting Amazon customer reviews with deep confidence using deep learning and conformal prediction," *Journal of Management Analytics*, vol. 9, no. 1, pp. 1–16, 2022, doi: 10.1080/23270012.2022.2031324.
- [25] Y. Rimal dan N. Sharma, "Ensemble machine learning prediction accuracy: local vs. global precision and recall for multiclass grade performance of engineering students," *Front. Educ. (Lausanne)*, vol. 10, pp. 1–27, 2025, doi: 10.3389/feduc.2025.1571133.
- [26] A. Ogunpola, F. Saeed, S. Basurra, A. M. Albarrak, dan S. N. Qasem, "Machine Learning-Based Predictive Models for Detection of Cardiovascular Diseases," *Diagnostics*, vol. 14, no. 2, pp. 1–31, 2024, doi: 10.3390/diagnostics14020144.
- [27] D. Chicco dan G. Jurman, "The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation," *BMC Genomics*, vol. 21, no. 1, pp. 1376–1397, 2020, doi: 10.1186/s12864-019-6413-7.
- [28] A. J. Bingham, "From Data Management to Actionable Findings: A Five-Phase Process of Qualitative Data Analysis," *Int. J. Qual. Methods*, vol. 22, pp. 1–26, 2023, doi: 10.1177/16094069231183620.
- [29] N. C. C. Brown, P. Weill-Tessier, M. Sekula, A. L. Costache, dan M. Kölling, "Novice Use of the Java Programming Language," *ACM Transactions on Computing Education*, vol. 23, no. 1, pp. 1–29, 2022, doi: 10.1145/3551393.

-
- [30] D. Chicco, L. Oneto, dan E. Tavazzi, "Eleven quick tips for data cleaning and feature engineering," *PLoS Comput. Biol.*, vol. 18, no. 12, pp. 2678-2687, 2022, doi: 10.1371/journal.pcbi.1010718.
- [31] S. Demir dan B. Topcu, "Graph-based Turkish text normalization and its impact on noisy text processing," *Engineering Science and Technology, an International Journal*, vol. 35, pp. 1-28, 2022, doi: 10.1016/j.jestch.2022.101192.
- [32] J. Khan, K. Ahmad, S. K. Jagatheesaperumal, dan K. A. Sohn, "Textual variations in social media text processing applications: challenges, solutions, and trends," *Artif. Intell. Rev.*, vol. 58, no. 3, pp. 3509-3518, 2025, doi: 10.1007/s10462-024-11071-z.
- [33] S. Sarica dan J. Luo, "Stopwords in technical language processing," *PLoS One*, vol. 16, no. 8, pp. 1-17, 2021, doi: 10.1371/journal.pone.0254937.
- [34] Z. Jiang, B. Gao, Y. He, Y. Han, P. Doyle, dan Q. Zhu, "Text Classification Using Novel Term Weighting Scheme-Based Improved TF-IDF for Internet Media Reports," *Math. Probl. Eng.*, vol. 2021, pp. 3908-3918, 2021, doi: 10.1155/2021/6619088.
- [35] H. Liu, X. Chen, dan X. Liu, "A Study of the Application of Weight Distributing Method Combining Sentiment Dictionary and TF-IDF for Text Sentiment Analysis," *IEEE Access*, vol. 10, pp. 32280-32289, 2022, doi: 10.1109/ACCESS.2022.3160172.
- [36] M. Tabany dan M. Gueffal, "Sentiment Analysis and Fake Amazon Reviews Classification Using SVM Supervised Machine Learning Model," *Journal of Advances in Information Technology*, vol. 15, no. 1, pp. 49-58, 2024, doi: 10.12720/jait.15.1.
- [37] X. Luo, "Efficient English text classification using selected Machine Learning Techniques," *Alexandria Engineering Journal*, vol. 60, no. 3, pp. 3401-3409, 2021, doi: 10.1016/j.aej.2021.02.009.
- [38] E. Elgeldawi, A. Sayed, A. R. Galal, dan A. M. Zaki, "Hyperparameter tuning for machine learning algorithms used for arabic sentiment analysis," *Informatics*, vol. 8, no. 4, pp. 1-22, 2021, doi: 10.3390/informatics8040079.