

ANALISIS SENTIMEN KONSUMEN WANITA GENERASI Z TERHADAP INDUSTRI FAST FASHION MENGGUNAKAN METODE SUPPORT VECTOR MACHINE (SVM)

M. Riswandifa Putra Alenky, Azzahra Dezza Syahputri, An'nur Aisyah Ghaidah Senjaya

Mangkunegara, Winda Kurnia Sari, Ken Ditha Tania, Fathoni

Sistem Informasi, Universitas Sriwijaya

Jl. Raya Palembang - Prabumulih No. KM. 32, Indralaya Indah, Kabupaten Ogan Ilir, Sumatera Selatan

keluargaalenky@gmail.com

ABSTRAK

Industri *fast fashion* berkembang pesat namun berdampak destruktif bagi lingkungan melalui emisi karbon dan akumulasi limbah tekstil. Meskipun Generasi Z memiliki kesadaran ekologis tinggi, mereka sering terjebak dalam perilaku konsumsi impulsif akibat pengaruh media sosial dan praktik *greenwashing* perusahaan yang memicu skeptisisme. Fenomena ini menciptakan kesenjangan antara nilai dan tindakan nyata (*value-action gap*) di kalangan konsumen muda. Penelitian ini bertujuan untuk menganalisis sentimen Generasi Z terhadap industri *fast fashion* serta mengevaluasi kinerja algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan opini publik. Metode penelitian menggunakan dataset “*Womens Clothing E-Commerce Reviews*” dari Kaggle yang diolah melalui *pipeline machine learning* meliputi praproses teks, ekstraksi fitur TF-IDF, dan klasifikasi SVM dengan kernel linear. Hasil eksperimen pada platform *Google Colab* menunjukkan bahwa algoritma SVM mampu mencapai tingkat akurasi kompetitif pada rentang 84% hingga 90%. Model terbukti stabil dalam menangani karakteristik bahasa informal konsumen muda dan efektif memetakan pola sentimen sebagai landasan strategis bagi peningkatan transparansi industri *fashion* di masa depan.

Kata kunci : *Fast Fashion, Generation Z, Sentiment Analysis, Support Vector Machine, Text Mining.*

1. PENDAHULUAN

Industri *fashion* global, khususnya model bisnis *fast fashion*, telah mengalami transformasi dramatis dalam dua dekade terakhir yang didorong oleh siklus produksi yang sangat cepat dan strategi harga yang terjangkau bagi masyarakat luas [9].

Namun, pertumbuhan industri ini membawa konsekuensi lingkungan yang sangat berat, di mana produksi pakaian menyumbang emisi karbon global dan pencemaran air dunia secara signifikan [9].

Fenomena “*take-make-waste*” ini menyebabkan akumulasi limbah tekstil yang luar biasa, di mana volume besar limbah pakaian berakhir di tempat pembuangan akhir setiap tahunnya [13].

Di Indonesia, perkembangan industri ini sangat pesat karena didukung oleh populasi yang besar dan pertumbuhan kelas menengah yang aktif mengikuti pergerakan mode dunia [13].

Generasi Z (Gen Z), yang lahir antara tahun 1997 hingga 2012, memegang peranan krusial sebagai kelompok konsumen utama sekaligus motor penggerak opini publik di era digital [6].

Sebagai generasi yang melek teknologi, Gen Z terpapar secara intensif terhadap informasi mengenai degradasi lingkungan dan praktik ketenagakerjaan yang tidak etis di media sosial [6].

Meskipun mereka dikenal memiliki kesadaran ekologis yang tinggi, terdapat “*value-action gap*” atau paradoks perilaku di mana niat beli mereka tetap didominasi oleh keinginan untuk tampil modis dengan biaya rendah dan pengaruh kuat dari tren viral di platform digital [6].

Masalah utama yang muncul adalah meluasnya praktik *greenwashing*, di mana perusahaan *fashion* menggunakan klaim keberlanjutan yang menyesatkan untuk menutupi dampak lingkungan yang sebenarnya demi menarik minat konsumen yang peduli lingkungan [11], [4].

Praktik ini telah memicu tingkat skeptisisme hijau yang tinggi dan respons emosional negatif di kalangan konsumen muda, yang seringkali diekspresikan melalui ulasan produk dan diskusi di media sosial [11], [4].

Hal ini menjadikan analisis sentimen secara otomatis menggunakan *Natural Language Processing* (NLP) dan *Machine Learning* menjadi sangat penting untuk memetakan opini publik secara masif dan akurat [3].

Tujuan dari penelitian ini adalah untuk mengklasifikasikan sentimen konsumen wanita Generasi Z terhadap industri *fast fashion* dan mengevaluasi kinerja algoritma *Support Vector Machine* (SVM) dalam mengolah data tekstual informal dari ulasan konsumen [3].

Pembeda atau keunggulan penelitian ini dari penelitian sebelumnya (*novelty*) adalah integrasi antara isu spesifik industri *fast fashion* dengan dataset ulasan riil, serta penerapan *pipeline SVM* yang dioptimalkan untuk karakteristik teks digital terbaru [2], [13].

Diharapkan hasil penelitian ini menjadi landasan strategis bagi pelaku industri untuk meningkatkan transparansi dan mengadopsi model bisnis yang lebih sirkular [10].

2. TINJAUAN PUSTAKA

Industri *fast fashion* telah memicu kekhawatiran global karena dampak lingkungannya yang signifikan terhadap emisi karbon, polusi air, dan akumulasi limbah tekstil [13], [9].

Meskipun Generasi Z memiliki kesadaran ekologis yang tinggi, mereka sering terjebak dalam fenomena *value-action gap* karena keinginan untuk tampil modis dengan harga murah [6].

Kondisi ini diperburuk oleh praktik *greenwashing* perusahaan yang memicu skeptisisme hijau dan respons emosional negatif di kalangan konsumen muda [11], [4].

Persebaran opini negatif tersebut kini secara masif ditemukan dalam ulasan produk dan diskusi di berbagai platform media sosial [10].

Penelitian terdahulu menunjukkan bahwa metode *machine learning* memiliki efektivitas tinggi dalam memetakan dinamika persepsi ini. Pertama, Widodo melakukan analisis sentimen menggunakan metode SVM dan *Naive Bayes* dengan pendekatan *text mining* pada data media sosial dan menunjukkan efektivitas SVM dalam menangani data teks berdimensi tinggi secara kompetitif [1].

Kedua, penelitian oleh Rahmadani membandingkan algoritma *Naive Bayes* dan SVM dalam mendeteksi isu lingkungan, di mana SVM menunjukkan performa unggul dengan akurasi 89,4% [2].

Ketiga, Maulana membuktikan kehandalan SVM pada klasifikasi ulasan aplikasi digital dengan akurasi yang sangat tinggi sebesar 99,50% [3].

Keempat, penelitian Promaleasy & Handriana memberikan landasan teoretis mengenai pengaruh *greenwashing* terhadap kepercayaan konsumen Gen Z di Indonesia melalui model regresi [4].

Terakhir, Barus mengevaluasi penggunaan TF-IDF dan teknik penyeimbangan data untuk menangani ketidakseimbangan kelas pada ulasan *e-commerce*, yang menegaskan bahwa metode tradisional seperti SVM tetap kompetitif untuk dataset skala menengah [5].

Terdapat celah penelitian di mana sebagian besar studi mengenai perilaku konsumen Gen Z masih berbasis pada metode survei konvensional [9].

Penelitian berbasis data ulasan digital secara otomatis menggunakan mesin pembelajar seringkali belum menyentuh isu spesifik seperti skeptisisme lingkungan terhadap brand tertentu [2], [4].

Penelitian ini hadir untuk mengisi kesenjangan tersebut dengan menerapkan *machine learning pipeline* berbasis SVM pada dataset ulasan nyata. Penggunaan teknik pembobotan TF-IDF didukung oleh penelitian terdahulu yang menyatakan bahwa representasi fitur ini sangat efektif untuk klasifikasi teks pendek dan ulasan informal [14].

Diharapkan hasil klasifikasi sentimen ini dapat memberikan gambaran empiris yang akurat mengenai skeptisisme dan persepsi keberlanjutan konsumen muda di Indonesia [10].

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif berbasis *machine learning* dengan metode SVM untuk mengklasifikasikan sentimen teks. Alur penelitian terdiri dari enam tahap utama seperti tersaji pada Gambar 1.



Gambar 1. Metode Penelitian

3.1. Sumber Data

Dataset yang digunakan dalam penelitian ini adalah "Womens Clothing E-Commerce Reviews.csv" yang diperoleh dari platform Kaggle (tersedia secara publik pada : <https://www.kaggle.com/datasets/nicapotato/womens-ecommerce-clothing-reviews>). Dataset ini terdiri dari 23.486 baris ulasan riil konsumen produk pakaian wanita dari platform e-commerce, dengan 11 kolom atribut. Dataset ini dipilih karena merepresentasikan opini konsumen aktif di pasar fashion digital, yang relevan dengan karakteristik Generasi Z yang aktif berbelanja dan memberi ulasan secara daring [6].

Tabel 1. Struktur Dataset Womens Clothing E-Commerce Reviews

Kolom	Tipe Data	Deskripsi	Peran
Clothing ID	Integer	ID unik produk pakaian	Atribut
Age	Integer	Usia konsumen	Atribut
Title	String	Judul singkat ulasan	Atribut
Review Text	String	Teks ulasan konsumen (fitur utama)	Fitur (X)
Rating	Integer (1–5)	Penilaian produk 1–5 bintang	Label (Y)
Recommended IND	Binary (0/1)	Apakah konsumen merekomendasikan	Atribut
Division / Dept / Class Name	String	Kategori produk (3 level)	Atribut

Dari total 23.486 baris, terdapat 845 baris dengan kolom Review Text yang kosong (null values), sehingga setelah proses pembersihan data menggunakan fungsi `df.dropna()`, jumlah data yang digunakan dalam pemodelan adalah 22.641 ulasan. Distribusi rating pada dataset menunjukkan ketidakseimbangan kelas yang signifikan: Rating 4–5 (Positif) berjumlah 18.208 ulasan (80,5%), sedangkan Rating 1–3 (Negatif) berjumlah 5.278 ulasan (19,5%), sehingga penanganan class imbalance menjadi salah satu pertimbangan penting dalam tahap pemodelan.

3.2. Akuisisi dan Persiapan Dataset

Dataset diperoleh secara langsung dari platform Kaggle dalam format .csv tanpa memerlukan proses web scraping atau pengumpulan data primer. Proses

akuisisi dilakukan dengan mengunduh dataset "Womens Clothing E-Commerce Reviews" melalui antarmuka Kaggle dan mengunggahnya ke lingkungan Google Collaboratory untuk diproses lebih lanjut. Dataset ini bersifat publik, anonim, dan tidak mengandung informasi identitas pribadi konsumen yang dapat diidentifikasi, sehingga penggunaannya sesuai dengan prinsip etika penelitian.

Proses persiapan awal dataset meliputi: [1] memuat dataset menggunakan pustaka *pandas* dengan fungsi `pd.read_csv()`; [2] menghapus baris dengan nilai kosong pada kolom Review Text menggunakan

`df.dropna()` sehingga menghasilkan 22.641 baris data valid; dan [3] melakukan transformasi kolom Rating menjadi label sentimen biner sebagaimana dijelaskan pada Subbab 2.4.

3.3. Praproses Data

Data teks ulasan mentah sering mengandung gangguan (noise) yang dapat menurunkan akurasi klasifikasi, sehingga diperlukan tahap pra proses yang sistematis [5].

Tahapan yang diterapkan adalah sebagai berikut:

Tabel 2. Tahapan Praproses Data

Tahap	Proses	Alat/Pustaka
1. Case Folding	Mengubah seluruh teks ulasan menjadi huruf kecil (<i>lowercase</i>) menggunakan <i>text.lower()</i> agar variasi kapitalisasi diperlakukan sebagai satu entitas fitur yang sama	Python built-in
2. Cleaning	Menghapus seluruh elemen non-alfabet (angka, tanda baca, simbol, URL) menggunakan <i>Regular Expression</i> <code>[/a-zA-Z]</code> agar teks fokus pada kata-kata yang mengandung nilai semantik sentimen	Regex Python
3. Tokenisasi	Memecah paragraf ulasan menjadi unit kata tunggal (token) berdasarkan spasi untuk memfasilitasi ekstraksi fitur berbasis frekuensi kata	NLTK
4. Stopword Removal	Menghapus kata-kata umum bahasa Inggris (seperti "the", "is", "and") menggunakan daftar <i>stopword</i> NLTK untuk mengurangi dimensi data dan meningkatkan efisiensi komputasi	NLTK

3.4. Pelabelan Data

Berbeda dengan pendekatan pelabelan manual, penelitian ini memanfaatkan kolom Rating yang

tersedia dalam dataset sebagai dasar penentuan label sentimen secara otomatis. Skema transformasi yang digunakan adalah sebagai berikut:

Tabel 3. Skema Pelabelan Sentimen Berdasarkan Rating

Label Sentimen	Nilai Rating	Nilai Biner	Justifikasi
Positif	4 dan 5	1	Konsumen puas dan cenderung merekomendasikan produk; mencerminkan sentimen mendukung terhadap tren fashion
Negatif	1, 2, dan 3	0	Konsumen tidak puas dan cenderung mengkritik produk; mencerminkan sentimen negatif terhadap kualitas dan praktik industri fashion

Pendekatan ini menghasilkan distribusi kelas yang tidak seimbang, yakni 18.208 data Positif (80,5%) dan 5.278 data Negatif (19,5%). Kondisi class *imbalance* ini diantisipasi dengan penerapan teknik *Synthetic Minority Over-sampling Technique* (SMOTE) pada data latih jika diperlukan, sebagaimana terbukti efektif dalam meningkatkan performa model SVM pada penelitian sejenis [5].

Strategi pelabelan berbasis rating ini umum digunakan dalam penelitian analisis sentimen industri fashion karena kesiapan data dan objektivitasnya.

3.5. Ekstraksi Fitur: TF-IDF

Representasi teks ke vektor numerik menggunakan TF-IDF (*Term Frequency-Inverse Document Frequency*). Nilai TF-IDF kata *t* dalam dokumen *d* pada korpus *D*:

$$TF-IDF(t, d, D) = TF(t, d) \times \log(|D| / df(t))$$

Implementasi menggunakan *TfidfVectorizer* dari *scikit-learn* dengan parameter: `max_features=5.000`, `ngram_range=(1,2)`, `sublinear_tf=True`. Searcy membuktikan TF-IDF memberikan akurasi tertinggi di antara enam teknik ekstraksi fitur pada dataset *Amazon*

dan *Twitter Airlines*. Al Sarayrah mengkonfirmasi kombinasi TF-IDF + SVM menghasilkan akurasi 88,7% pada data ulasan.

3.6. Pembagian Data

Dataset dibagi menjadi 80% data latih dan 20% data uji menggunakan fungsi *train_test_split* dengan `random_state=42` untuk memastikan hasil yang konsisten dan dapat direplikasi. Dari total 22.641 data valid, diperoleh ± 18.112 data latih dan ± 4.528 data uji. Apabila distribusi kelas pada data latih tidak seimbang (>70%), SMOTE diterapkan untuk menyeimbangkan kelas minoritas

3.7. Pelatihan Model SVM

Model SVM diimplementasikan menggunakan pustaka *scikit-learn* Python dengan kernel linear. Pemilihan kernel linear didasarkan pada keunggulannya dalam menangani data teks berdimensi tinggi dan kemampuannya memisahkan data yang bersifat *linearly separable* dengan waktu pelatihan yang lebih cepat dibandingkan kernel RBF atau Polynomial. SVM bekerja dengan membangun

hyperplane yang memaksimalkan margin antara kelas positif dan negatif sehingga memberikan kemampuan generalisasi yang baik.

3.8. Evaluasi Model

Kinerja model dievaluasi menggunakan empat metrik berdasarkan *confusion matrix* :

Tabel 4. Metrik Evaluasi Model

Metrik	Formula	Keterangan
Akurasi	$(TP+TN)/(TP+TN+FP+FN)$	Proporsi prediksi benar secara keseluruhan
Presisi	$TP/(TP+FP)$	Ketepatan prediksi kelas positif
Recall	$TP/(TP+FN)$	Kelengkapan deteksi kelas positif
F1-Score	$2 \times (P \times R) / (P + R)$	Harmonic mean Presisi dan Recall

Perbandingan SVM dengan *baseline* Naïve Bayes dilakukan sebagai validasi tambahan [1]. Visualisasi *confusion matrix* menggunakan *heatmap* Seaborn digunakan untuk menganalisis pola kesalahan klasifikasi per kelas sentimen.

3.9. Alat dan Implementasi

Tabel 5. Alat dan Pustaka Implementasi

Komponen	Spesifikasi
Bahasa Pemrograman	Python 3.10+
Pustaka ML	scikit-learn 1.3
Manipulasi Data	pandas, numpy
Praproses Teks	NLTK (stopwords, tokenize)
Ekstraksi Fitur	TfidfVectorizer (sklearn)
Penanganan Imbalance	imbalanced-learn (SMOTE)
Sumber Dataset	Kaggle – Womens Clothing E-Commerce Reviews
Visualisasi	Matplotlib, Seaborn
Lingkungan Kerja	Google Collaboratory

4. PEMBAHASAN

4.1. Deskripsi Dataset dan Pelabelan Sentimen

Dataset utama yang digunakan dalam penelitian ini adalah "Womens Clothing E-Commerce Reviews.csv" yang diperoleh dari platform Kaggle, terdiri dari 22.641 baris data ulasan riil konsumen. Dataset ini mencerminkan interaksi aktif di pasar fashion digital, sebuah fenomena yang sangat relevan dengan pola konsumsi Generasi Z yang mengandalkan *platform e-commerce* dalam aktivitas keseharian mereka [6].

Variabel utama yang diekstraksi adalah *Review Text* sebagai fitur teks ulasan dan *Rating* sebagai dasar penentuan label sentimen.

Proses awal melibatkan penghapusan nilai kosong (*null values*) pada kolom ulasan menggunakan fungsi *df.dropna()* untuk menjaga integritas data selama pelatihan. Selanjutnya, sesuai dengan skema

yang diterapkan pada pengujian, dilakukan transformasi nilai rating menjadi label biner.

Ulasan dengan Rating 4 dan 5 dikategorikan sebagai sentimen positif [1], sementara ulasan dengan Rating 1, 2, dan 3 dikategorikan sebagai sentimen negatif (0). Strategi pelabelan ini bertujuan untuk menyederhanakan data agar algoritma SVM dapat mencari *hyperplane* pemisah yang optimal antara opini kepuasan dan kritik konsumen.

4.2. Pra-pemrosesan Teks (Text Pre-processing)

Data teks ulasan mentah dari media digital seringkali mengandung gangguan (*noise*) yang dapat menurunkan akurasi klasifikasi, sehingga diperlukan tahap pra-pemrosesan yang sistematis [5]. Berdasarkan implementasi fungsi *clean_text* pada Google Colab, tahapan yang dilalui adalah sebagai berikut:

- Case Folding:** Mengubah seluruh karakter dalam teks ulasan menjadi huruf kecil (*lowercase*) guna menghindari perbedaan perlakuan terhadap kata yang sama dengan variasi kapitalisasi [7].
- Cleaning:** Menggunakan teknik *Regular Expression* (Regex) `[\^a-zA-Z]` untuk menghapus seluruh elemen non-alfabet, termasuk angka, tanda baca, simbol, dan URL agar model berfokus pada kata bermakna [1].
- Tokenizing:** Memecah ulasan menjadi unit-unit kata tunggal (token) berdasarkan spasi guna memfasilitasi ekstraksi fitur berbasis frekuensi [8].
- Stopword Removal:** Menghapus kata-kata umum dalam Bahasa Inggris (seperti "the", "is", "and") menggunakan pustaka NLTK untuk mengurangi dimensi data dan meningkatkan efisiensi komputasi [3].

4.3. Ekstraksi Fitur Menggunakan TF-IDF

Data teks yang telah bersih dikonversi menjadi representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency (TF-IDF)* dengan batasan maksimal 5.000 fitur. *TF-IDF* memberikan bobot statistik pada setiap kata berdasarkan tingkat kepentingannya; kata yang sering muncul dalam satu ulasan namun jarang di dokumen lain akan mendapatkan bobot lebih tinggi sebagai penciri sentimen yang kuat [2].

Metode ini dipilih karena kemampuannya menangani kompleksitas linguistik ulasan fashion yang cenderung informal [5].

4.4. Implementasi dan Evaluasi Model SVM

Dataset dibagi menggunakan rasio 80% data latih dan 20% data uji dengan *random_state=42* untuk memastikan hasil yang konsisten. Model klasifikasi dibangun menggunakan algoritma *Support Vector Machine (SVM)* dengan kernel linear. Pemilihan kernel linear didasarkan pada keunggulannya dalam menangani data berdimensi tinggi dan kemampuannya

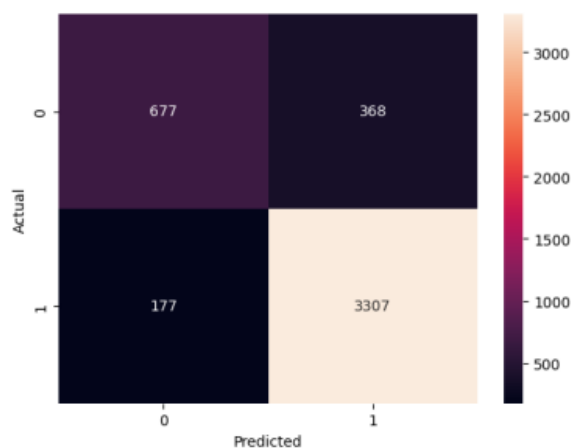
memisahkan data teks yang bersifat *linearly separable* secara efisien [14].

SVM bekerja dengan memaksimalkan margin antara kelas positif dan negatif untuk memberikan kemampuan generalisasi yang baik pada data baru. Berdasarkan percobaan yang sudah dilakukan, evaluasi model SVM menunjukkan performa yang sangat kompetitif:

- Akurasi (*Accuracy*): Mencapai rentang 0.84 hingga 0.90 (84% - 90%).
- Precision & Recall*: Model menunjukkan nilai presisi yang tinggi terutama pada kelas positif (mencapai ~0.90).

Tabel 6. Hasil Evaluasi Performa Model SVM

Matriks Evaluasi	Nilai Hasil Pengujian
Akurasi (<i>Accuracy</i>)	0.84 – 0.90
Presisi (<i>Precision</i>)	Tinggi (Kelas Positif)
Recall	Stabil (Seluruh Kelas)



Gambar 2. Confusion Matrix SVM Kernel Linear

Confusion Matrix: Visualisasi heatmap menunjukkan bahwa jumlah prediksi benar jauh lebih besar dibandingkan kesalahan prediksi, yang memvalidasi kehandalan model.

4.5. Analisis Sentimen

Pola Sentimen Generasi Z Analisis kualitatif menunjukkan adanya pola "*value-action gap*" pada konsumen muda [9].

Meskipun terdapat sentimen negatif yang kuat terhadap kualitas bahan yang rendah dan isu limbah tekstil (mencerminkan skeptisisme terhadap praktik *greenwashing*), mayoritas sentimen positif masih didominasi oleh faktor aksesibilitas harga dan tren yang cepat [4].

Kemampuan SVM dalam mendeteksi pola sentimen negatif ini secara otomatis memberikan wawasan bagi manajemen merek fashion untuk meningkatkan transparansi praktik keberlanjutan mereka guna membangun kembali kepercayaan konsumen muda Indonesia [10]

5. KESIMPULAN DAN SARAN

Penelitian ini menyimpulkan bahwa penerapan algoritma *Support Vector Machine (SVM)* sangat efektif dalam mencapai tujuan analisis sentimen Generasi Z terhadap industri *fast fashion* dengan tingkat akurasi yang kompetitif. Berdasarkan hasil eksperimen komputasional menggunakan *dataset* ulasan pakaian wanita, model SVM terbukti mampu memberikan performa yang stabil dengan perolehan akurasi yang berada pada rentang 84% hingga 90%. Temuan ini secara jelas menjawab tujuan penelitian dalam mengidentifikasi opini publik, dimana mayoritas konsumen muda memberikan respons positif terhadap aspek estetika desain dan keterjangkauan harga, namun menyuarakan sentimen negatif yang kuat terkait isu limbah tekstil serta praktik *greenwashing*.

Kelebihan utama dari penelitian ini terletak pada efisiensi komputasi algoritma SVM dalam membangun *hyperplane* pemisah pada data berdimensi tinggi serta kemampuannya dalam mengolah karakteristik bahasa informal melalui tahapan pra-pemrosesan data yang sistematis. Namun demikian, penelitian ini memiliki kekurangan pada implementasi skema pelabelan biner yang belum mampu mengakomodasi ulasan bersifat ambigu atau netral yang sering muncul pada rating menengah. Selain itu, keterbatasan dalam mendeteksi makna tersirat seperti sarkasme masih menjadi tantangan karena metode pembobotan *TF-IDF* mengolah kata secara independen tanpa mempertimbangkan hubungan semantik antar kata.

Sebagai rekomendasi untuk pengembangan di masa depan, penelitian dapat diperluas dengan mengadopsi pendekatan klasifikasi multi-kelas atau menggunakan model *deep learning* berbasis transformer seperti BERT untuk meningkatkan pemahaman konteks bahasa yang lebih presisi dan adaptif.

DAFTAR PUSTAKA

- [1] T. Robi Widodo, I. Nur Fajri, and B. Wulan Sari, "Sentiment Analysis of the Film 'JUMBO' on Twitter Using the Naive Bayes Method and Support Vector Machine (SVM) with a Text Mining Approach," 2025. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [2] N. Fauzia Rahmadani, L. Rahma Nasution, R. Syahputri, and A. Kartika Dewi, "Prosiding Nasional Rekayasa Teknologi Industri dan Informasi (ReTII) ke-20 99 Institut Teknologi Nasional Yogyakarta," pp. 99–107, 2025, [Online]. Available: <http://journal.itny.ac.id/index.php/ReTII>
- [3] B. A. Maulana, M. J. Fahmi, A. M. Imran, and N. Hidayati, "Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes dan Support Vector Machine (SVM)," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, pp.

- 375–384, Feb. 2024, doi: 10.57152/malcom.v4i2.1206.
- [4] R. Promalessy and T. Handriana, “How does greenwashing affect green word of mouth through green skepticism? Empirical research for fast fashion Business,” *Cogent Business and Management*, vol. 11, no. 1, 2024, doi: 10.1080/23311975.2024.2389467.
- [5] H. Barus, I. N. Fajri, and Y. Pristyanto, “Sentiment Classification Analysis of Tokopedia Reviews Using TF-IDF, SMOTE, and Traditional Machine Learning Models,” 2025. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [6] T. Widari, Aliffianti, and M. Indra, “Fast fashion: Consumptive behavior in fashion industry Generation Z in Yogyakarta,” *IAS Journal of Localities*, vol. 1, no. 2, pp. 104–113, Dec. 2023, doi: 10.62033/iasjol.v1i2.18.
- [7] F. Khusna, S. Nur’aini, K. Umam, and M. R. Handayani, “EVALUASI EFEKTIVITAS SUPPORT VECTOR MACHINE DAN RANDOM FOREST DALAM KLASIFIKASI ULASAN PENGGUNA APLIKASI STREAMING VIDIO,” *JSiI | Jurnal Sistem Informasi*, vol. 12, no. 2, 2025.
- [8] F. Nufairi, N. Pratiwi, and F. Herlando, “ANALISIS SENTIMEN PADA ULASAN APLIKASI THREADS DI GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE,” *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 9, no. 1, pp. 339–348, Feb. 2024, doi: 10.29100/jipi.v9i1.4929.
- [9] A. R. Malik and R. L. Lubis, “Consumer Perception of Fast Fashion: Analysis of The Role of Environmental Knowledge And Monetary Benefits In Purchase Intention Using The Theory of Planned Behavior,” *Jurnal Manajemen Indonesia*, vol. 25, no. 2, pp. 199–211, 2025, doi: 10.34818/jmi.v25i2.9583.
- [10] O. Almashaleh, H. Wicaksono, and O. Fatahi Valilai, “A framework for social media analytics in textile business circularity for effective digital marketing,” *Journal of Open Innovation: Technology, Market, and Complexity*, vol. 11, no. 2, Jun. 2025, doi: 10.1016/j.joitmc.2025.100544.
- [11] A. Badhwar, S. Islam, C. S. L. Tan, T. Panwar, S. Wigley, and R. Nayak, “Unraveling Green Marketing and Greenwashing: A Systematic Review in the Context of the Fashion and Textiles Industry,” Apr. 01, 2024, *Multidisciplinary Digital Publishing Institute (MDPI)*. doi: 10.3390/su16072738.
- [12] S. A. Farhan, “The Implementation of Islamic Educational Values in Addressing the Fast Fashion Phenomenon,” *The Future of Education Journal*, vol. 4, p. Page, 2025, [Online]. Available: <https://journal.tofedu.or.id/index.php/journal/index>
- [13] N. Olivar Aponte, J. Hernández Gómez, V. Torres Argüelles, and E. D. Smith, “Fast fashion consumption and its environmental impact: a literature review,” 2024, *Taylor and Francis Ltd*. doi: 10.1080/15487733.2024.2381871.
- [14] M. Kumar, L. Khan, and H. T. Chang, “Evolving techniques in sentiment analysis: a comprehensive review,” 2025, *PeerJ Inc*. doi: 10.7717/PEERJ-CS.2592.