

Analisis Sentimen Ulasan Pelanggan Terhadap Tomoro Coffe Menggunakan Metode Random Forest Berbasis Text Mining

¹Desta Anna Pratama Diyandi, ²Chandra Kirana

^{1*,2}Teknik Informatika, Institut Sains dan Bisnis Atma Luhur, Bangka Belitung, Indonesia

*Korespondensi: 2211500041@mahasiswa.atmaluhur.ac.id

Submit : 16 Jun 2026 | Diterima : 06 Jul 2026 | Terbit : 01 Sept 2026

ABSTRACT

Using a Random Forest technique based on text mining, this research seeks to examine consumer sentiment about Tomoro Coffee reviews. There were a total of 927 reviews culled from various online sources; 546 were favorable and 381 were unfavorable. Data collection, sentiment labeling, and text preparation (including cleaning, stopword removal, tokenizing, and stemming) were the study steps that came before data transformation using TF-IDF. After then, the dataset was divided into two halves, with training and testing data comprising 80:20 of the total. The Random Forest method was then used for classification. The accuracy was 80.65%, precision was 83.58%, recall was 70.10%, and the F1-score was 79.50%, according to the test findings. It was clear from the confusion matrix that the model excelled at detecting positive emotion as opposed to negative. Good, comfortable, and pleasant were the most often used positive emotion terms in the word cloud, while application, voucher, promotion, and payment were the most commonly used negative sentiment words. Based on the findings, it seems that Tomoro Coffee's product and service quality may be evaluated using the Random Forest approach, which is based on text mining. This method is effective for categorizing customer review sentiment.

Keywords: Sentiment Analysis, Tomoro Coffee, Random Forest, Text Mining, TF-IDF.

ABSTRAK

Dengan menggunakan teknik Random Forest berbasis *text mining*, penelitian ini bertujuan untuk menguji sentimen konsumen tentang ulasan Tomoro Coffee. Sebanyak 927 ulasan dikumpulkan dari berbagai sumber daring; 546 ulasan bersifat positif dan 381 ulasan bersifat negatif. Pengumpulan data, pelabelan sentimen, dan persiapan teks (termasuk pembersihan, penghapusan stopwords, tokenisasi, dan stemming) merupakan langkah-langkah penelitian sebelum transformasi data menggunakan TF-IDF. Setelah itu, dataset dibagi menjadi dua bagian, dengan data pelatihan dan pengujian masing-masing terdiri dari 80:20 dari total data. Metode Random Forest kemudian digunakan untuk klasifikasi. Akurasi yang diperoleh adalah 80,65%, presisi 83,58%, recall 70,10%, dan F1-score 79,50%, menurut hasil pengujian. Terlihat jelas dari matriks kebingungan bahwa model unggul dalam mendeteksi emosi positif dibandingkan emosi negatif. Baik, nyaman, dan menyenangkan adalah istilah emosi positif yang paling sering digunakan dalam word cloud, sedangkan aplikasi, voucher, promosi, dan pembayaran adalah kata-kata sentimen negatif yang paling umum digunakan. Berdasarkan temuan tersebut, tampaknya kualitas produk dan layanan Tomoro Coffee dapat dievaluasi menggunakan pendekatan Random Forest, yang berbasis pada text mining. Metode ini efektif untuk mengkategorikan sentimen ulasan pelanggan.

Kata kunci: Analisis sentimen, Tomoro Coffee, Random Forest, Text Mining, TF-IDF

PENDAHULUAN

Cara konsumen menyuarakan pendapat mereka tentang barang dan jasa telah berevolusi sebagai respons terhadap perkembangan TIK. Pelanggan dapat memposting evaluasi secara bebas dan cepat di berbagai platform digital, termasuk Play Store, media sosial, dan aplikasi online. Ulasan pelanggan adalah cara yang bagus untuk mengetahui bagaimana perasaan orang tentang layanan yang mereka terima, serta aspek positif atau negatif dari layanan tersebut. Dalam industri coffee shop, ulasan pelanggan dapat dimanfaatkan sebagai

bahan evaluasi untuk memahami kualitas produk, pelayanan, harga, maupun kenyamanan tempat secara lebih objektif. Penelitian sebelumnya menunjukkan bahwa data ulasan pada Play Store dapat digunakan untuk mengidentifikasi respons pelanggan terhadap layanan coffee shop melalui pendekatan klasifikasi sentimen menggunakan algoritma Random Forest (Isnayni, Saputra, and Hastono 2024).

Pertumbuhan industri coffee shop di Indonesia menunjukkan tren yang terus meningkat seiring berkembangnya budaya konsumsi kopi dan perubahan gaya hidup masyarakat. Kondisi tersebut mendorong munculnya berbagai merek coffee shop yang bersaing dalam menawarkan kualitas produk dan pelayanan terbaik. Salah satu merek yang berkembang pesat adalah Tomoro Coffee yang mengadopsi konsep layanan modern berbasis teknologi digital. Seiring meningkatnya jumlah pelanggan, volume ulasan yang dihasilkan pada berbagai platform digital juga semakin besar. Analisis manual menjadi kurang berhasil karena volume data yang terus meningkat. Metode ini memakan waktu, bersifat subjektif, dan mungkin mengabaikan informasi penting dalam umpan balik pelanggan. (Caesar et al. 2026). Oleh sebab itu, diperlukan pendekatan berbasis komputasi yang mampu melakukan pengolahan data ulasan secara otomatis, sistematis, dan terukur.

Analisis sentimen adalah salah satu metode yang dapat digunakan untuk memahami sentimen konsumen. Metode untuk menemukan, mengekstrak, dan mengkategorikan perasaan atau pandangan dari data teks dikenal sebagai analisis sentimen (Romadoni, Umaidah, and Sari 2020). Untuk memberikan gambaran umum tentang bagaimana konsumen memandang suatu produk atau layanan, analisis sentimen berupaya mengkategorikan komentar ke dalam kategori positif, negatif, dan netral dalam konteks ulasan pelanggan (Jurnal and Informasi 2025), (Amalia, Sawit, and Timur n.d.). Pada industri coffee shop, informasi mengenai sentimen pelanggan menjadi penting karena tingkat kepuasan konsumen dipengaruhi oleh berbagai faktor, seperti kualitas rasa produk, mutu pelayanan, kecepatan penyajian, fasilitas, suasana tempat, serta kesesuaian harga.

Penggunaan algoritma text mining, yang mengambil data dari teks tidak terstruktur, sangat penting untuk analisis sentimen. Dengan bantuan text mining, dimungkinkan untuk mengungkap informasi dan pola yang sebelumnya tidak diketahui dalam basis data teks yang sangat besar (Julians, Herman, and Manongga 2024). Dalam pengolahan ulasan pelanggan, tahapan text mining umumnya diawali dengan proses preprocessing yang meliputi tokenization, case folding, filtering, stemming, dan pembobotan kata (Firdaus, Wijoyo, and Purnomo 2025), (Morama, Ratnawati, and Arwani 2022). Salah satu metode pembobotan yang banyak digunakan adalah TF-IDF, yaitu metode yang digunakan untuk mengukur tingkat kepentingan suatu kata dalam sebuah dokumen relatif terhadap keseluruhan kumpulan dokumen (Jurnal and Informasi 2025). Melalui penerapan TF-IDF, data teks dapat direpresentasikan dalam bentuk numerik sehingga dapat diproses menggunakan algoritma machine learning.

Berbagai penelitian terdahulu telah menunjukkan keberhasilan penerapan analisis sentimen pada beragam objek penelitian. Penelitian pada Gunthem Premium Coffee membuktikan bahwa pendekatan berbasis TF-IDF mampu mengelompokkan ulasan pelanggan ke dalam kategori sentimen positif, negatif, dan netral (Jurnal and Informasi 2025). Selain itu, analisis sentimen juga telah diterapkan untuk mengidentifikasi opini pengguna terhadap layanan digital dan aplikasi berbasis daring (Romadoni et al. 2020). Dalam proses klasifikasi teks, beberapa algoritma yang umum digunakan antara lain Naïve Bayes, SVM, KNN, dan Random Forest (Pradhana et al. 2025), (Anggina, Setiawan, and Bachtiar 2022). Pada objek penelitian Tomoro Coffee, studi sebelumnya menggunakan algoritma SVM untuk menganalisis sentimen pelanggan (Caesar et al. 2026), sedangkan penelitian pada coffee shop lainnya lebih banyak memanfaatkan algoritma Naïve Bayes maupun pendekatan lexicon-based (Pradhana et al. 2025), (Anggina et al. 2022).

Kinerja Random Forest yang kuat pada data berdimensi tinggi menjadikannya teknik klasifikasi yang populer. Cara kerja teknik ini adalah dengan membangun banyak pohon keputusan dan kemudian menggunakan metode pemungutan suara mayoritas untuk menentukan hasil kategorisasi (Amalia et al. n.d.), (Morama et al. 2022). Kelebihan Random Forest meliputi ketahanannya terhadap overfitting, kemampuannya untuk memproses data yang bising, dan konsistensi kinerjanya di berbagai dataset (Amalia et al. n.d.), (Firdaus et al. 2025). Untuk riset sentimen aplikasi digital dan layanan online, Random Forest telah terbukti memberikan hasil klasifikasi yang baik dalam sejumlah studi sebelumnya (Firdaus et al. 2025), (No, Insany, and Kharisma 2024), (Rania et al. n.d.), (Larasati, Ratnawati, and Hanggara 2022).

Selain itu, algoritma ini juga dinilai mampu menangani variasi penggunaan bahasa dalam data teks dan menghasilkan tingkat akurasi yang kompetitif dibandingkan metode klasifikasi lainnya (Amalia et al. n.d.), (No et al. 2024).

Meskipun demikian, masih terdapat sejumlah keterbatasan dalam penelitian sebelumnya. Beberapa penelitian menggunakan jumlah dataset yang relatif terbatas, menghasilkan tingkat akurasi yang belum optimal, serta hanya berfokus pada klasifikasi sentimen positif dan negatif tanpa mempertimbangkan kategori netral (Isnayni et al. 2024), (Dharma, Candra, and Turnip 2023), (Wijayanto et al. 2023), (Anggina et al. 2022). Selain itu, penelitian mengenai analisis sentimen ulasan pelanggan Tomoro Coffee dengan pendekatan Random Forest berbasis text mining masih belum banyak ditemukan. Sebagian besar penelitian pada objek Tomoro Coffee masih berfokus pada penggunaan algoritma SVM (Caesar et al. 2026). Kondisi tersebut menunjukkan adanya peluang penelitian untuk mengevaluasi kemampuan Random Forest dalam mengklasifikasikan sentimen pelanggan pada objek yang sama. (Amalia et al. n.d.), (No et al. 2024), (Rania et al. n.d.), (Larasati et al. 2022).

Untuk meneliti sentimen ulasan pelanggan terhadap Tomoro Coffee, penelitian ini menyarankan penggunaan pendekatan Random Forest yang berbasis pada text mining. Tujuan penelitian ini adalah untuk mengklasifikasikan ulasan pelanggan ke dalam tiga jenis sentimen: positif, negatif, dan netral. Penelitian ini juga akan menilai seberapa baik algoritma Random Forest dalam klasifikasi teks. Terakhir, penelitian ini akan memberikan informasi yang dapat digunakan Tomoro Coffee untuk meningkatkan kualitas produk dan layanan mereka. Kontribusi penelitian ini terletak pada penerapan Random Forest pada data ulasan pelanggan Tomoro Coffee yang masih terbatas dikaji dalam penelitian sebelumnya, sehingga diharapkan dapat memberikan referensi empiris mengenai efektivitas algoritma tersebut dalam analisis sentimen berbasis text mining.

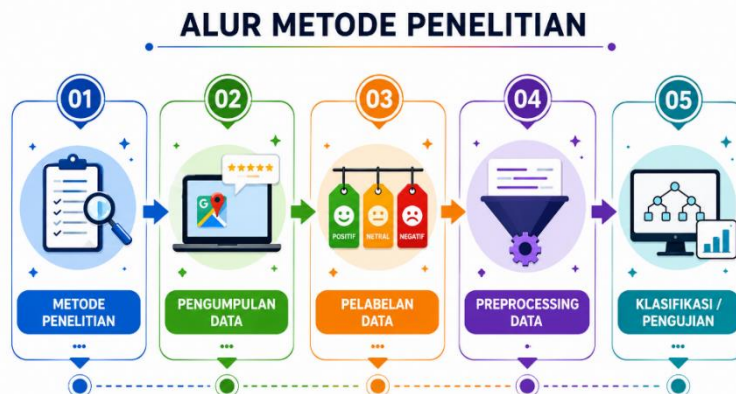
METODE PENELITIAN

Penelitian ini mengkaji sentimen ulasan pelanggan Tomoro Coffee menggunakan metode kuantitatif komputasional. Hal ini dilakukan dengan menggunakan teknik text mining dan algoritma pembelajaran mesin, khususnya Random Forest. Karena penelitian ini berkaitan dengan pengukuran objektif data teks yang ditransformasikan secara numerik, yang memungkinkan prosedur kategorisasi yang sistematis dan terukur, maka metode kuantitatif digunakan.

Secara umum, tahapan penelitian dimulai dari pengumpulan data ulasan, pemberian label sentimen, preprocessing teks, transformasi data menggunakan TF-IDF, hingga proses klasifikasi dengan Random Forest. Untuk menganalisis lebih lanjut kemampuan model dalam membedakan sentimen dalam ulasan pelanggan, selanjutnya peneliti menggunakan ukuran akurasi, presisi, recall, dan F1-score untuk memeriksa hasil kategorisasi.

Alur Penelitian

Berikut alur penelitian yang selaras dengan struktur yang di inginkan:



Gambar 1. Alur Penelitian

Pengumpulan Data

Karena adanya korelasi langsung antara kualitas data dan hasil analisis, fase pengumpulan data merupakan langkah pertama yang penting dalam setiap proyek penelitian. Evaluasi pelanggan terhadap Tomoro Coffee yang ditemukan di saluran digital, khususnya Google Play Store, digunakan sebagai data dalam penelitian ini. Google Play Store dipilih karena merupakan platform populer tempat pengguna memberikan ulasan dan peringkat setelah berinteraksi dengan suatu layanan atau perusahaan.

Ulasan pelanggan pada dasarnya berupa teks bebas yang tidak terstruktur. Di dalamnya terkandung berbagai informasi yang berkaitan dengan pengalaman pelanggan, seperti cita rasa kopi, pelayanan barista, kebersihan tempat, kenyamanan suasana, kecepatan penyajian, hingga kesesuaian harga. Karakteristik data seperti ini sangat sesuai untuk dianalisis menggunakan text mining karena mengandung opini, persepsi, dan evaluasi subjektif dari pelanggan.

Secara teknis, data dapat dikumpulkan melalui web scraping atau dengan cara manual menggunakan perangkat lunak tertentu, selama tetap memperhatikan etika dalam pengelolaan data publik. Data yang dikumpulkan umumnya mencakup isi review, rating, waktu ulasan, serta identitas nonpribadi pengguna apabila tersedia. Dalam penelitian ini, perhatian utama difokuskan pada teks ulasan karena bagian tersebut merupakan sumber utama dalam analisis sentimen.

Setelah data terkumpul, data disusun ke dalam format dataset, seperti CSV atau Excel, agar dapat diproses pada tahap berikutnya. Dataset tersebut kemudian digunakan sebagai dasar untuk proses pelabelan, preprocessing, dan klasifikasi. Oleh karena itu, pengumpulan data berfungsi sebagai prosedur pengumpulan informasi dan dasar utama untuk menentukan validitas keseluruhan penelitian.

Pelabelan Data

Setelah data berhasil dikumpulkan, tahap selanjutnya adalah pelabelan data, yaitu proses pemberian kategori sentimen pada setiap ulasan berdasarkan isi opini yang disampaikan. Proses ini dilakukan agar data dapat digunakan sebagai data latih dalam metode supervised learning dengan kategori sentimen positif, negatif, dan apabila diperlukan netral.

Makna setiap ulasan diperiksa secara manual untuk menyelesaikan proses pelabelan. Sebuah ulasan dianggap positif jika mengungkapkan kebahagiaan, kepuasan, atau pengalaman yang menyenangkan; ulasan negatif adalah ulasan yang mengungkapkan ketidakpuasan, kritik, atau keluhan. Ulasan dianggap netral jika memberikan informasi yang bermanfaat tanpa menunjukkan prasangka yang jelas.

Keakuratan klasifikasi sangat bergantung pada label, sehingga langkah ini sangat penting. Apabila proses pelabelan dilakukan secara tidak konsisten, maka model akan mempelajari pola yang keliru sehingga akurasi yang dihasilkan menurun. Oleh karena itu, pelabelan harus dilakukan secara cermat berdasarkan konteks dan makna ulasan. Dataset yang telah diberi label inilah yang kemudian dijadikan masukan utama dalam proses pembelajaran Random Forest.

Preprocessing Data

Data yang dikumpulkan dari ulasan pelanggan di platform digital seringkali belum diproses dan mungkin mencakup ejaan yang tidak teratur, simbol, kosakata non-standar, akronim, tanda baca, dan karakteristik bahasa informal. Karena itu, persiapan diperlukan untuk mempersiapkan data agar siap dikategorikan. Salah satu bagian terpenting dari text mining adalah pra-pemrosesan, yaitu proses pembersihan, normalisasi, dan pengorganisasian data teks sebagai persiapan untuk ekstraksi informasi yang efektif.

Tahapan preprocessing dalam penelitian ini meliputi case folding, tokenisasi, filtering, stemming, dan transformasi TF-IDF.

Case Folding

Case folding merupakan proses mengubah seluruh karakter huruf dalam teks menjadi huruf kecil. Tahap ini dilakukan untuk menyeragamkan bentuk kata sehingga perbedaan penggunaan huruf besar dan huruf kecil tidak memengaruhi hasil analisis. Sebagai contoh, kata "Enak", "ENAK", dan "enak" akan diseragamkan menjadi "enak". Langkah ini penting untuk menghindari munculnya duplikasi term yang sebenarnya memiliki makna sama.

Tokenisasi

Tokenisasi adalah proses memecah kalimat menjadi unit-unit kata atau token. Sebagai contoh, kalimat “kopi tomoro enak dan nyaman” akan diuraikan menjadi beberapa token, yaitu [kopi, tomoro, enak, dan, nyaman]. Proses ini bertujuan agar setiap kata dapat dianalisis secara individual pada tahapan berikutnya.

Filtering

Filtering atau stopwords removal dilakukan untuk menghilangkan kata-kata umum yang tidak memberikan kontribusi besar terhadap makna sentimen, seperti kata hubung “dan”, “yang”, preposisi “di, ke, dari, serta kata fungsional lainnya.” Dengan menghapus stopwords, proses analisis menjadi lebih terfokus pada kata-kata yang benar-benar merepresentasikan sentimen, seperti “enak, cepat, ramah, buruk, atau lama”.

Stemming

Stemming merupakan proses mengubah kata berimbuhan menjadi bentuk kata dasar. Dalam bahasa Indonesia, tahap ini sangat penting karena ulasan pelanggan sering menggunakan kata yang telah mendapat imbuhan, seperti “pelayanannya”, “makanannya”, “menyenangkan”, atau “kecepatannya”. Dengan stemming, kata-kata tersebut diubah ke bentuk dasar agar tidak dianggap sebagai fitur yang berbeda. Proses ini membantu meningkatkan konsistensi representasi teks dan memperbaiki kualitas klasifikasi.

Transformasi TF-IDF

Agar algoritma pembelajaran mesin dapat menggunakan teks sebagai input setelah tahap pra-pemrosesan selesai, teks tersebut harus diubah menjadi bentuk numerik. Penelitian ini menggunakan teknik TF-IDF untuk melakukan transformasi tersebut. TF-IDF adalah metode pembobotan kata yang menunjukkan seberapa signifikan suatu kata dalam kaitannya dengan keseluruhan korpus.

Rumus TF-IDF terdiri atas dua komponen utama, yaitu TF dan IDF.

$$TF(t, d) = \frac{f_{t,d}}{\sum_k f_{k,d}} \quad (1)$$

$$IDF(t) = \log \left(\frac{N}{df_t} \right) \quad (2)$$

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (3)$$

Keterangan:

- $TF(t, d)$ = Frekuensi term t pada dokumen d
- $f_{t,d}$ = Jumlah kemunculan term t dalam dokumen d
- N = Jumlah seluruh dokumen
- df_t = Jumlah dokumen yang mengandung term t

Secara teori, kata-kata yang umum dalam satu teks tetapi jarang ditemukan dalam teks lain seharusnya memiliki bobot yang lebih besar. Dengan demikian, TF-IDF mampu menonjolkan kata-kata yang memiliki relevansi tinggi terhadap sentimen ulasan.

Klasifikasi dan Pengujian

Tahap klasifikasi merupakan inti dari penelitian ini karena pada bagian inilah model Random Forest digunakan untuk mengidentifikasi sentimen ulasan pelanggan. Dataset pelatihan dan pengujian dibuat dari dataset yang telah diproses sebelumnya setelah transformasi TF-IDF. Segmentasi ini memungkinkan model untuk dilatih pada sebagian data sebelum dievaluasi kemampuannya untuk mengkategorikan data yang tidak terduga.

Random Forest bekerja dengan membangun sejumlah decision tree secara acak dari sampel data latih. Setiap pohon keputusan menghasilkan prediksi masing-masing, kemudian hasil akhir ditentukan berdasarkan voting mayoritas dari seluruh pohon. Kemampuan untuk mengurangi varians model dan bahaya overfitting, yang umum terjadi pada pohon keputusan tunggal, merupakan keunggulan dari strategi ini.

Secara matematis, fungsi prediksi Random Forest dapat dituliskan sebagai:

$$\hat{y} = \text{mode}(h_1(x), h_2(x), \dots, h_B(x)) \quad (4)$$

Keterangan:

- $\hat{y}$ = hasil prediksi akhir
- $h_1(x), h_2(x), \dots, h_B(x)$ = prediksi dari masing-masing pohon keputusan
- B = jumlah pohon dalam hutan
- mode = kelas yang paling banyak muncul

Dalam proses pengujian, hasil klasifikasi dievaluasi menggunakan sejumlah metrik performa, berupa "accuracy, precision, recall, dan F1-score." Metrik ini digunakan untuk mengevaluasi sejauh mana model mampu mengenali kelas sentimen dengan benar.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (5)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (6)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (7)$$

$$F1 = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (8)$$

Keterangan:

- True Positive (TP) merupakan kondisi ketika data yang sebenarnya bernilai positif berhasil diprediksi secara tepat sebagai positif.
- True Negative (TN) merupakan kondisi ketika data yang sebenarnya bernilai negatif berhasil diprediksi secara tepat sebagai negatif.
- False Positive (FP) merupakan kondisi ketika data yang sebenarnya bernilai negatif tetapi diprediksi sebagai positif.
- False Negative (FN) merupakan kondisi ketika data yang sebenarnya bernilai positif tetapi diprediksi sebagai negatif.

Selain metrik evaluasi tersebut, confusion matrix juga digunakan untuk memberikan gambaran rinci mengenai performa model pada setiap kelas. Melalui evaluasi ini, peneliti dapat mengetahui apakah model lebih baik dalam mengidentifikasi sentimen "positif, negatif, atau netral," sekaligus menilai keseimbangan performa antar kelas.

Hasil pengujian kemudian dianalisis untuk mengetahui efektivitas Random Forest dalam mengklasifikasikan ulasan pelanggan Tomoro Coffee. Analisis ini penting untuk menilai apakah metode yang digunakan telah sesuai dengan karakteristik data teks serta apakah hasilnya dapat dijadikan dasar untuk interpretasi sentimen pelanggan secara lebih luas.

HASIL DAN PEMBAHASAN

Deskripsi Dataset Penelitian

Ulasan pengguna Tomoro Coffee yang diberi label sentimen digunakan dalam penelitian ini. Ulasan tersebut diambil dari Play Store. Dari total 927 evaluasi, 546 memiliki emosi positif dan 381 memiliki emosi negatif. Pendekatan Random Forest, yang berbasis pada text mining, digunakan pada dataset ini sebagai sumber utama untuk analisis sentimen.

Kami memeriksa kualitas data untuk memastikan tidak ada nilai yang hilang untuk atribut apa pun sebelum kami mulai mengklasifikasikan. Ditetapkan bahwa dataset tersebut lengkap karena pemeriksaan `isna()` mengembalikan respons False untuk setiap atribut. Kondisi ini menunjukkan bahwa data telah memenuhi persyaratan untuk diproses lebih lanjut tanpa memerlukan tahapan penanganan *missing value*.

Tabel 1. Distribusi Dataset Sentimen

Label Sentimen	Jumlah Data	Persentase
Positif	546	58,90%
Negatif	381	41,10%
Total	927	100%

***	userName	score	at	content
0	Nakano Itsuki	5	2026-05-19 06:18:50	TOMORO COFFEE YANG TERBAIKKKK
1	Delsy Mutiara Diva	1	2026-05-18 09:34:28	kami beli voc tomoro, gabisa digunakan mana ga...
2	Dzaki Fikri	1	2026-05-17 11:17:55	tolong voucher duo healing seriesnya segera di...
3	Valleta Natasya	2	2026-05-17 11:17:36	Tolong di perbaiki dong, masa voucher discount...
4	V (Palen)	5	2026-05-17 03:37:53	Good :>

Gambar 2. Grafik Distribusi Sentimen

Berdasarkan Tabel 4.1, distribusi sentimen dalam dataset didominasi oleh ulasan positif sebanyak 58,90%, sedangkan ulasan negatif sebesar 41,10%. Dominasi sentimen positif mengindikasikan bahwa sebagian besar pelanggan memberikan penilaian yang baik terhadap produk maupun layanan yang ditawarkan oleh Tomoro Coffee. Namun demikian, proporsi sentimen negatif yang masih cukup besar menunjukkan adanya sejumlah pelanggan yang menyampaikan ketidakpuasan terhadap aspek tertentu dari layanan yang diterima.

Dari sudut pandang klasifikasi, distribusi data tersebut tergolong cukup seimbang sehingga mampu mendukung proses pembelajaran model secara optimal. Ketidakseimbangan data yang terlalu tinggi berpotensi menyebabkan model lebih cenderung memprediksi kelas mayoritas. Oleh karena itu, komposisi dataset yang digunakan dalam penelitian ini dinilai cukup representatif untuk mendukung proses klasifikasi sentimen.

Hasil Preprocessing Data

Dalam penelitian berbasis text mining, langkah pra-pemrosesan sangat penting karena bertujuan untuk membersihkan, mengatur, dan mempersiapkan data teks mentah untuk analisis. Pra-pemrosesan dalam penelitian ini mencakup banyak langkah, termasuk membersihkan teks, menghilangkan kata-kata penghenti (stopwords), tokenisasi, dan stemming.

Text Cleaning

Selama langkah pembersihan teks, semua karakter khusus dihilangkan dan teks diubah menjadi huruf kecil. Ini termasuk tanda baca, simbol, angka, URL, dan emoji.

Tabel 2 Contoh Hasil Text Cleaning

Sebelum Cleaning	Setelah Cleaning
TOMORO COFFEE YANG TERBAIKKKK	tomoro coffee yang terbaikkkk
Good :>	good

	content	Label	text_clean
0	TOMORO COFFEE YANG TERBAIKKKK	Positif	tomoro coffee yang terbaikkkk
1	kami beli voc tomoro, gabisa digunakan mana ga...	Negatif	kami beli voc tomoro gabisa digunakan mana gab...
2	tolong voucher duo healing seriesnya segera di...	Negatif	tolong voucher duo healing seriesnya segera di...
3	Tolong di perbaiki dong, masa voucher discount...	Negatif	tolong di perbaiki dong masa voucher discount ...
4	Good :>	Positif	good
6	plis perbaiki metode pembayaran nya, saya mau ...	Negatif	plis perbaiki metode pembayaran nya saya mau b...
7	AYO TAMBAHKAN DIAPLIKASI PEMBAYARAN MELALUI QR...	Positif	ayo tambahkan diaplikasi pembayaran melalui qr...

Gambar 3. Hasil Text Cleaning

Hasil proses ini menghasilkan bentuk teks yang lebih seragam sehingga memudahkan sistem dalam mengenali pola kata pada tahap berikutnya. Normalisasi teks diperlukan karena sistem komputer akan menganggap kata yang memiliki perbedaan huruf kapital sebagai kata yang berbeda meskipun memiliki makna yang sama.

Stopword Removal

Tahap berikutnya adalah *stopword removal*, yaitu proses menghapus kata-kata umum yang tidak memiliki kontribusi signifikan dalam penentuan sentimen.

Tabel 3. Contoh Hasil Stopword Removal

Sebelum Stopword Removal	Setelah Stopword Removal
kami beli voc tomoro gabisa digunakan mana ga	beli voc tomoro gabisa digunakan
Sebelum Stopword Removal	Setelah Stopword Removal
tolong voucher duo healing seriesnya segera diperbaiki	voucher duo healing series diperbaiki

	content	Label	text_clean	text_StopWord
0	TOMORO COFFEE YANG TERBAIKKKK	Positif	tomoro coffee yang terbaikkkk	tomoro coffee terbaikkkk
1	kami beli voc tomoro, gabisa digunakan mana ga...	Negatif	kami beli voc tomoro gabisa digunakan mana gab...	beli voc tomoro gabisa gabisa return sesuai ke...
2	tolong voucher duo healing seriesnya segera di...	Negatif	tolong voucher duo healing seriesnya segera di...	tolong voucher duo healing seriesnya diperbaik...
3	Tolong di perbaiki dong, masa voucher discount...	Negatif	tolong di perbaiki dong masa voucher discount ...	tolong perbaiki voucher discount collab pakai
4	Good :>	Positif	good	
6	plis perbaiki metode pembayaran nya, saya mau ...	Negatif	plis perbaiki metode pembayaran nya saya mau b...	plis perbaiki metode pembayaran bayar pakai da...
7	AYO TAMBAHKAN DIAPLIKASI PEMBAYARAN MELALUI QR...	Positif	ayo tambahkan diaplikasi pembayaran melalui qr...	ayo tambahkan diaplikasi pembayaran qris minnnn

Gambar 4. Hasil Stopword Removal

Proses ini membantu mengurangi jumlah fitur yang tidak relevan sehingga model dapat lebih fokus pada kata-kata yang memiliki kandungan informasi sentimen. Selain itu, pengurangan jumlah fitur juga berkontribusi terhadap peningkatan efisiensi komputasi pada tahap klasifikasi.

Tokenizing

Tokenizing merupakan proses pemisahan teks menjadi unit-unit kata yang lebih kecil atau token.

Tabel 4. Contoh Hasil Tokenizing

Kalimat		Hasil Token			
tomoro coffee yang terbaikkkk		[tomoro, coffee, terbaikkkk]			
content	Label	text_clean	text_StopWord	text_tokens	
0 TOMORO COFFEE YANG TERBAIKKKK	Positif	tomoro coffee yang terbaikkkk	tomoro coffee terbaikkkk	[tomoro, coffee, terbaikkkk]	
1 kami beli voc tomoro, gabisa digunakan mana ga...	Negatif	kami beli voc tomoro gabisa digunakan mana gab...	beli voc tomoro gabisa gabisa return sesuai ke...	[beli, voc, tomoro, gabisa, gabisa, return, se...	
2 tolong voucher duo healing seriesnya segera di...	Negatif	tolong voucher duo healing seriesnya segera di...	tolong voucher duo healing seriesnya diperbaik...	[tolong, voucher, duo, healing, seriesnya, dip...	
3 Tolong di perbaiki dong, masa voucher discount...	Negatif	tolong di perbaiki dong masa voucher discount ...	tolong perbaiki voucher discount collab pakai	[tolong, perbaiki, voucher, discount, collab, ...	

Gambar 5. Hasil Tokenizing

Melalui proses ini, setiap kata dapat diperlakukan sebagai fitur independen yang selanjutnya akan dihitung bobotnya menggunakan metode TF-IDF. Tahapan ini menjadi dasar dalam proses ekstraksi fitur pada analisis teks.

Stemming

Tahap *stemming* dilakukan untuk mengubah kata yang memiliki imbuhan menjadi bentuk dasarnya.

Tabel 5. Contoh Hasil Stemming

Kata Awal		Kata Dasar				
pembayaran		bayar				
diperbaiki		baik				
pelayanan		layan				
	content	Label	text_clean	text_StopWord	text_tokens	text_steamindo
0	TOMORO COFFEE YANG TERBAIKKKK	Positif	tomoro coffee yang terbaikkkk	tomoro coffee terbaikkkk	[tomoro, coffee, terbaikkkk]	tomoro coffee terbaikkkk
1	kami beli voc tomoro, gabisa digunakan mana ga...	Negatif	kami beli voc tomoro gabisa digunakan mana gab...	beli voc tomoro gabisa gabisa return sesuai ke...	[beli, voc, tomoro, gabisa, gabisa, return, se...	beli voc tomoro gabisa gabisa return sesuai te...
2	tolong voucher duo healing seriesnya segera di...	Negatif	tolong voucher duo healing seriesnya segera di...	tolong voucher duo healing seriesnya diperbaiki...	[tolong, voucher, duo, healing, seriesnya, dip...	tolong voucher duo healing seriesnya baik baik...
3	Tolong di perbaiki dong, masa voucher discount	Negatif	tolong di perbaiki dong masa voucher discount	tolong perbaiki voucher discount collab pakai	[tolong, perbaiki, voucher, discount, collab	tolong baik voucher discount collab pakai

Gambar 6. Hasil Stemming

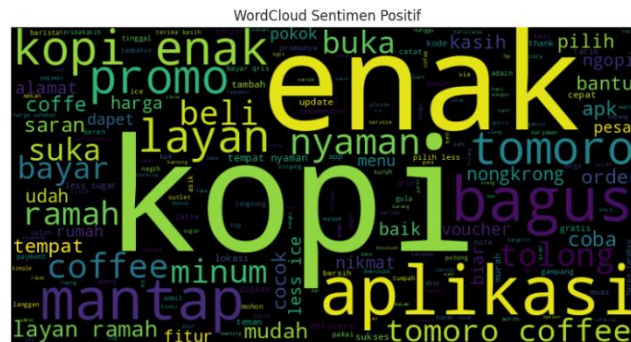
Hasil stemming menunjukkan bahwa berbagai variasi kata yang memiliki makna serupa dapat direduksi menjadi satu bentuk dasar yang sama. Proses ini membantu meningkatkan konsistensi representasi data teks sekaligus mengurangi jumlah fitur yang harus diproses oleh model klasifikasi.

Analisis Word Cloud Sentimen

Dengan menggunakan visualisasi word cloud, analisis eksplorasi dilakukan untuk menemukan

istilah yang paling sering muncul dalam setiap kategori sentimen setelah semua proses pra-pemrosesan selesai.

Word Cloud Sentimen Positif



Gambar 7. Word Cloud Sentimen Positif

Visualisasi *word cloud* pada sentimen positif memperlihatkan dominasi kata-kata seperti *enak*, *kopi*, *bagus*, *mantap*, *nyaman*, *ramah*, dan *aplikasi*.

Kemunculan kata *enak* dengan frekuensi yang tinggi menunjukkan bahwa kualitas rasa produk menjadi faktor utama yang memengaruhi kepuasan pelanggan. Selain itu, kata *nyaman* dan *ramah* mengindikasikan bahwa pelanggan juga memberikan perhatian terhadap suasana tempat dan kualitas pelayanan yang diterima. Temuan ini menunjukkan bahwa pengalaman pelanggan tidak hanya dipengaruhi oleh kualitas produk, tetapi juga oleh aspek pelayanan dan kenyamanan lingkungan.

Word Cloud Sentimen Negatif



Gambar 8. Word Cloud Sentimen Negatif

Pada kategori sentimen negatif, kata-kata yang paling dominan meliputi *aplikasi*, *voucher*, *promo*, *order*, *jam*, *bayar*, dan *beli*.

Dominasi kata-kata tersebut mengindikasikan bahwa sebagian besar keluhan pelanggan berkaitan dengan aspek operasional dan layanan digital. Permasalahan yang sering muncul mencakup penggunaan aplikasi, sistem pembayaran, program promosi, serta proses pemesanan produk. Hasil ini menunjukkan bahwa sumber utama ketidakpuasan pelanggan lebih banyak berasal dari sistem layanan pendukung dibandingkan kualitas produk yang ditawarkan.

Hasil Transformasi TF-IDF

Setelah tahap pra-pemrosesan selesai, pendekatan TF-IDF digunakan untuk ekstraksi fitur. Tujuan dari teknik ini adalah untuk membuat data dapat digunakan oleh algoritma pembelajaran mesin dengan mengubahnya dari teks menjadi angka.

['aaah' 'abal' 'abik' ... 'you' 'zidan' 'zonk']

Trik di balik TF-IDF adalah memberikan bobot lebih besar pada istilah-istilah yang sering muncul

tetapi tidak sering terjadi secara keseluruhan

Dengan pendekatan tersebut, kata-kata yang memiliki kemampuan lebih tinggi dalam membedakan suatu sentimen akan memperoleh bobot yang lebih besar.

Penerapan TF-IDF menghasilkan ribuan fitur kata yang kemudian digunakan sebagai data masukan bagi algoritma Random Forest. Representasi numerik ini memungkinkan model melakukan pembelajaran pola sentimen secara lebih efektif dan sistematis.

Hasil Klasifikasi Random Forest

Dataset kemudian dibagi menggunakan skema 80:20, yaitu 741 data training dan 186 data testing.

Tabel 6. Pembagian Dataset

Data	Jumlah
Training	741
Testing	186

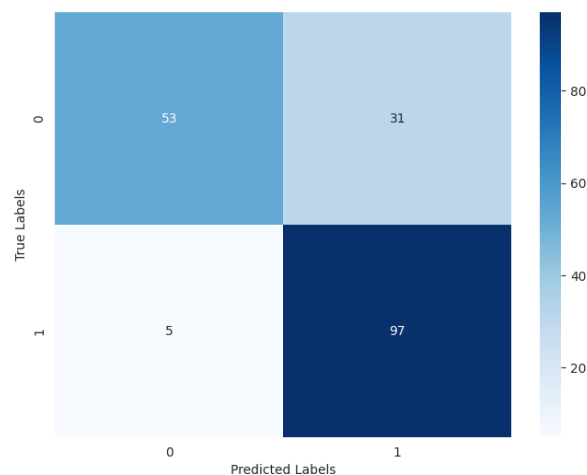
... (741,)
(741,)
(186,)
(186,)

Gambar 9. Pembagian Data 80:20

Berdasarkan hasil pengujian model Random Forest diperoleh hasil sebagai berikut:

Tabel 7 Hasil Evaluasi Random Forest

Metrik	Nilai
Accuracy	80,65%
Precision	83,58%
Recall	79,10%
F1-Score	79,50%



Gambar 10. Grafik Hasil Evaluasi Random Forest

Model Random Forest menunjukkan tingkat akurasi sebesar 80,65% dalam hasil pengujian. Dengan angka ini, kita dapat melihat bahwa model tersebut berhasil mengidentifikasi kategori sentimen dengan benar hampir 80% dari ulasan pelanggan.

Perolehan nilai akurasi tersebut mengindikasikan bahwa Random Forest memiliki kemampuan yang cukup baik dalam mengenali pola sentimen pelanggan berdasarkan fitur hasil transformasi TF-IDF. Selain itu, nilai precision yang lebih tinggi dibandingkan recall menunjukkan bahwa model relatif baik dalam menghasilkan prediksi yang tepat, meskipun masih terdapat beberapa data yang belum berhasil diklasifikasikan sesuai dengan kelas sebenarnya.

Analisis Confusion Matrix

Untuk mengetahui performa model secara lebih rinci digunakan confusion matrix.

Tabel 8 Confusion Matrix

Aktual	Prediksi Negatif	Prediksi Positif
Negatif	53	31
Positif	5	97

Berdasarkan hasil confusion matrix, model berhasil mengklasifikasikan 53 ulasan negatif dan 97 ulasan positif secara benar. Di sisi lain, terdapat 31 ulasan negatif yang diprediksi sebagai positif dan 5 ulasan positif yang diprediksi sebagai negatif.

Temuan ini menunjukkan bahwa model tersebut unggul dalam mengidentifikasi sentimen positif dibandingkan sentimen negatif. Ulasan negatif cenderung memiliki berbagai kualitas yang kompleks, seperti yang terlihat dari banyaknya kesalahan klasifikasi dalam kelompok tersebut. Selain itu, beberapa ulasan negatif kemungkinan mengandung unsur opini campuran sehingga menyulitkan model dalam menentukan kategori sentimen secara tepat.

Pembahasan

Temuan menunjukkan bahwa tingkat akurasi klasifikasi sebesar 80,65% dicapai dengan menggabungkan Random Forest dengan TF-IDF. Hasil menunjukkan bahwa pendekatan tersebut berhasil mengidentifikasi tren sentimen dalam ulasan Tomoro Coffee.

Kata-kata seperti "lezat, kopi, enak, stabil, dan nyaman" mendominasi sentimen positif yang ditemukan dalam studi word cloud, yang berfokus pada kualitas produk dan pengalaman pelanggan. Semua indikasi mengarah pada kesimpulan bahwa rasa produk, suasana, dan layanan adalah tiga faktor terpenting dalam menentukan apakah konsumen puas atau tidak. Sebaliknya, sentimen negatif lebih banyak berkaitan dengan aspek operasional dan layanan digital, seperti aplikasi, voucher, promo, pembayaran, dan proses pemesanan. Temuan tersebut mengindikasikan bahwa peningkatan kualitas layanan digital perlu menjadi perhatian yang sama pentingnya dengan peningkatan kualitas produk.

Dari sisi pemodelan, Random Forest mampu memanfaatkan fitur hasil transformasi TF-IDF secara efektif melalui mekanisme ensemble yang meningkatkan stabilitas prediksi dan mengurangi risiko overfitting. Namun demikian, model masih menunjukkan keterbatasan dalam mengenali sebagian ulasan negatif, yang tercermin dari nilai recall kelas negatif sebesar 63,10%. Kondisi ini mengindikasikan bahwa ulasan negatif cenderung memiliki struktur kalimat yang lebih kompleks, ambigu, atau mengandung campuran sentimen yang menyulitkan proses klasifikasi.

Untuk penelitian selanjutnya, performa model dapat ditingkatkan melalui optimasi hyperparameter, penambahan jumlah dataset, penerapan teknik penyeimbangan data seperti SMOTE, maupun perbandingan dengan algoritma lain seperti SVM) Naïve Bayes, dan Logistic Regression. Jika digunakan bersama-sama, TF-IDF dan Random Forest menyediakan alat yang ampuh untuk menyelidiki sentimen konsumen dan memberikan wawasan yang dapat ditindaklanjuti untuk meningkatkan kualitas layanan dan kepuasan melalui pengambilan keputusan berbasis data.

KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa metode Random Forest berbasis text mining mampu diterapkan secara efektif untuk menganalisis sentimen ulasan pelanggan terhadap Tomoro Coffee. Pelanggan umumnya memberikan evaluasi yang baik terhadap barang dan jasa yang tersedia, karena sebagian besar dari 927 ulasan yang ditinjau termasuk dalam kategori suasana hati yang positif. Tahapan preprocessing yang terdiri atas *text cleaning*, *stopword removal*, *tokenizing*, dan *stemming* berhasil mengubah data teks mentah menjadi data yang lebih terstruktur sehingga dapat diproses secara optimal oleh model klasifikasi. Selain itu, transformasi menggunakan TF-IDF mampu menghasilkan representasi numerik yang efektif dalam menonjolkan kata-kata penting yang berkontribusi terhadap pembentukan sentimen. Algoritma Random Forest menunjukkan tingkat akurasi, recall, dan F1-score yang cukup baik sebesar 80,65% dalam hasil pengujian. Namun, algoritma ini lebih akurat dalam mengidentifikasi emosi positif daripada negatif. Selain itu, analisis word cloud mengungkapkan bahwa elemen operasional seperti aplikasi, kupon, promosi, dan metode pembayaran lebih terkait dengan sikap negatif, sementara fitur seperti layanan yang menyenangkan, kenyamanan, dan kualitas rasa mendominasi sentimen positif. Akibatnya, Tomoro Coffee dapat menggunakan temuan studi ini untuk menilai dan meningkatkan kualitas layanan mereka, khususnya di bidang operasional dan layanan digital bidang di mana pelanggan masih memiliki keluhan terbanyak.

REFERENSI

- Amalia, Ghina, Duren Sawit, And Kota Jakarta Timur. N.D. "INSTAGRAM MENGGUNAKAN ALGORITMA RANDOM FOREST." 13(3).
- Anggina, Sarah, Nanang Yudi Setiawan, And Fitra A. Bachtiar. 2022. "Analisis Ulasan Pelanggan Menggunakan Multinomial Naïve Bayes Classifier Dengan Lexicon-Based Dan TF-IDF Pada Formaggio Coffee And Resto." 7:76–90.
- Caesar, Yustitio, Dwijananda Galih Prameswara, Teknik Informatika, Fakultas Teknik, Universitas Nusantara, And Pgri Kediri. 2026. "Analisis Sentimen Ulasan Pelanggan Tomoro Coffee Menggunakan Algoritma Support Vector Machine (SVM)." 5:57–62.
- Dharma, Abdi, Windy Candra, And Josua Presen Turnip. 2023. "Implementation Of Random Forest Algoritihm On Sales Data To Predict Potential Churn Of Products." 7(2):866–72.
- Firdaus, Raihan Zahran, Satrio Hadi Wijoyo, And Welly Purnomo. 2025. "Analisis Sentimen Berbasis Aspek Ulasan Pengguna Aplikasi Alfagift Menggunakan Metode Random Forest Dan Pemodelan Topik Latent Dirichlet Allocation." 9(2):1–10.
- Isnayni, Berliana Nur, Nurirwan Saputra, And Tri Hastono. 2024. "SENTIMENT ANALYSIS OF COFFEE SHOP REVIEWS USING RANDOM FOREST CLASSIFIER METHOD." 1(4):233–44.
- Julians, Adhe Ronny, Daniel Herman, And Fredy Manongga. 2024. "Analisis Konten Budaya Kolaboratif Berbasis Grounded Theory Menggunakan Text Mining." 21(2):230–50.
- Jurnal, Jtik, And Teknologi Informasi. 2025. "ANALISIS SENTIMEN ULASAN KONSUMEN MENGGUNAKAN ALGORITMA TF-ID UNTUK (STUDI KASUS : GUNTHER PREMIUM COFFEE)." 16(2):8–16.
- Larasati, Fanka Angelina, Dian Eka Ratnawati, And Buce Trias Hanggara. 2022. "Analisis Sentimen Ulasan Aplikasi Dana Dengan Metode Random Forest." 6(9):4305–13.
- Morama, Hana Chyntia, Dian Eka Ratnawati, And Issa Arwani. 2022. "Analisis Sentimen Berbasis Aspek Terhadap Ulasan Hotel Tentrem Yogyakarta Menggunakan Algoritma Random Forest Classifier." 6(4):1702–8.
- No, Vol, Gina Purnama Insany, And Ivana Lucia Kharisma. 2024. "Edumatic : Jurnal Pendidikan Informatika Penerapan Algoritma Random Forest Untuk Menganalisis Ulasan Aplikasi Spotify Pada Google Play." 8(2):369–78. Doi: 10.29408/Edumatic.V8i2.26394.
- Pradhana, Faisal Reza, Eko Prasetyo Widhi, Niken Ratna Sari, Muhammad Amrico, Putra Nurcahyo, Program Studi, Teknik Informatika, Fakultas Sains, And Universitas Darussalam Gontor. 2025. "KLASIFIKASI SENTIMEN ULASAN COFFEE SHOP DI PONOROGO BERBASIS NAÏVE BAYES." 4(2):739–49. Doi: 10.70247/Jumistik.V4i2.239.
- Rania, Naura Zainaty, Rama Dian Syah, Fakultas Teknologi, Industri Universitas, Informasi Universitas Gunadarma, Jawa Barat, Random Forest, And Aplikasi Gojek. N.D. "ANALISIS SENTIMEN TERHADAP APLIKASI GOJEK PADA PLAY STORE MENGGUNAKAN METODE RANDOM FOREST CLASSIFIER." 29(2):144–53.

- Romadoni, Fajar, Yuyun Umaidah, And Betha Nurina Sari. 2020. "Text Mining Untuk Analisis Sentimen Pelanggan Terhadap Layanan Uang Elektronik Menggunakan Algoritma Support Vector Machine." 09:247–53.
- Wijayanto, Sena, Dedy Agung Prabowo, Daniel Yeri Kristiyanto, And M. Yoka Fathoni. 2023. "Analisis Sentimen Berbasis Aspek Pada Layanan Hotel Di Wilayah Kabupaten Banyumas Dengan Word2Vec Dan Random Forest." 8(1):1–3.
- Amalia, Ghina, Duren Sawit, And Kota Jakarta Timur. N.D. "INSTAGRAM MENGGUNAKAN ALGORITMA RANDOM FOREST." 13(3).
- Anggina, Sarah, Nanang Yudi Setiawan, And Fitra A. Bachtiar. 2022. "Analisis Ulasan Pelanggan Menggunakan Multinomial Naïve Bayes Classifier Dengan Lexicon-Based Dan TF-IDF Pada Formaggio Coffee And Resto." 7:76–90.
- Caesar, Yustitio, Dwijananda Galih Prameswara, Teknik Informatika, Fakultas Teknik, Universitas Nusantara, And Pgri Kediri. 2026. "Analisis Sentimen Ulasan Pelanggan Tomoro Coffee Menggunakan Algoritma Support Vector Machine (SVM)." 5:57–62.
- Dharma, Abdi, Windy Candra, And Josua Presen Turnip. 2023. "Implementation Of Random Forest Algoritim On Sales Data To Predict Potential Churn Of Products." 7(2):866–72.
- Firdaus, Raihan Zahran, Satrio Hadi Wijoyo, And Welly Purnomo. 2025. "Analisis Sentimen Berbasis Aspek Ulasan Pengguna Aplikasi Alfagift Menggunakan Metode Random Forest Dan Pemodelan Topik Latent Dirichlet Allocation." 9(2):1–10.
- Isnayni, Berliana Nur, Nurirwan Saputra, And Tri Hastono. 2024. "SENTIMENT ANALYSIS OF COFFEE SHOP REVIEWS USING RANDOM FOREST CLASSIFIER METHOD." 1(4):233–44.
- Julians, Adhe Ronny, Daniel Herman, And Fredy Manongga. 2024. "Analisis Konten Budaya Kolaboratif Berbasis Grounded Theory Menggunakan Text Mining." 21(2):230–50.
- Jurnal, Jtik, And Teknologi Informasi. 2025. "ANALISIS SENTIMEN ULASAN KONSUMEN MENGGUNAKAN ALGORITMA TF-ID UNTUK (STUDI KASUS: GUNTHERM PREMIUM COFFEE)." 16(2):8–16.
- Larasati, Fanka Angelina, Dian Eka Ratnawati, And Buce Trias Hanggara. 2022. "Analisis Sentimen Ulasan Aplikasi Dana Dengan Metode Random Forest." 6(9):4305–13.
- Morama, Hana Chyntia, Dian Eka Ratnawati, And Issa Arwani. 2022. "Analisis Sentimen Berbasis Aspek Terhadap Ulasan Hotel Tentrem Yogyakarta Menggunakan Algoritma Random Forest Classifier." 6(4):1702–8.
- No, Vol, Gina Purnama Insany, And Ivana Lucia Kharisma. 2024. "Edumatic : Jurnal Pendidikan Informatika Penerapan Algoritma Random Forest Untuk Menganalisis Ulasan Aplikasi Spotify Pada Google Play." 8(2):369–78. Doi: 10.29408/Edumatic.V8i2.26394.
- Pradhana, Faisal Reza, Eko Prasetyo Widhi, Niken Ratna Sari, Muhammad Amrico, Putra Nurcahyo, Program Studi, Teknik Informatika, Fakultas Sains, And Universitas Darussalam Gontor. 2025. "KLASIFIKASI SENTIMEN ULASAN COFFEE SHOP DI PONOROGO BERBASIS NAÏVE BAYES." 4(2):739–49. Doi: 10.70247/Jumistik.V4i2.239.
- Rania, Naura Zainaty, Rama Dian Syah, Fakultas Teknologi, Industri Universitas, Informasi Universitas Gunadarma, Jawa Barat, Random Forest, And Aplikasi Gojek. N.D. "ANALISIS SENTIMEN TERHADAP APLIKASI GOJEK PADA PLAY STORE MENGGUNAKAN METODE RANDOM FOREST CLASSIFIER." 29(2):144–53.
- Romadoni, Fajar, Yuyun Umaidah, And Betha Nurina Sari. 2020. "Text Mining Untuk Analisis Sentimen Pelanggan Terhadap Layanan Uang Elektronik Menggunakan Algoritma Support Vector Machine." 09:247–53.
- Wijayanto, Sena, Dedy Agung Prabowo, Daniel Yeri Kristiyanto, And M. Yoka Fathoni. 2023. "Analisis Sentimen Berbasis Aspek Pada Layanan Hotel Di Wilayah Kabupaten Banyumas Dengan Word2Vec Dan Random Forest." 8(1):1–3.