



Analisis Sentimen Kepuasan Publik terhadap Pembaruan Alun-Alun Jember Menggunakan Metode Random Forest

Mochamad Yoga Kurniawan^{1*}, Deni Arifianto², Wiwik Suharso³

¹⁻³ Universitas Muhammadiyah Jember, Indonesia

Email: yoga.kurniawan2323@gmail.com^{1*}, deniarifianto@unmuhjember.ac.id²,
wiwiksuharso@unmuhjember.ac.id³

*Penulis Korespondensi: yoga.kurniawan2323@email.com

Abstract. The renovation of Jember Alun-Alun has sparked a variety of responses from the public, many of which have been expressed on social media, necessitating an automated method to efficiently analyze these public opinions. This study aims to analyze public sentiment toward the renovation of Jember Alun-Alun using the Random Forest algorithm and to identify the distribution of sentiment into three categories: positive, negative, and neutral. Data was collected via web scraping techniques using APIs from the X and TikTok platforms with the keywords #jembernusantara and #alunalunjember during the period from June 6, 2024, to February 20, 2025, yielding a total of 1,420 comments, which were then processed through the stages of cleaning, case folding, tokenization, normalization, stopword removal, stemming, and feature weighting using TF-IDF, as well as class balancing with SMOTE and model evaluation using K-Fold Cross Validation (K=2 to K=10). The test results showed that the Random Forest model was able to classify sentiment with an accuracy ranging from 89% to 96%, with the highest accuracy of 96% at K=6, K=8, and K=10, and the best F1-Score for the negative class reaching 100%. It is therefore concluded that the combination of the Random Forest and SMOTE methods is effective for analyzing public sentiment in social media text data.

Keywords: Jember Square; K-Fold Cross-Validation; Random Forest; Sentiment Analysis; SMOTE.

Abstrak. Pembaruan Alun-Alun Jember menimbulkan beragam respons dari masyarakat yang banyak diekspresikan melalui media sosial, sehingga diperlukan metode otomatis untuk menganalisis opini publik tersebut secara efisien. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap pembaruan Alun-Alun Jember menggunakan algoritma Random Forest serta mengidentifikasi distribusi sentimen ke dalam tiga kategori: positif, negatif, dan netral. Data dikumpulkan melalui teknik web scraping menggunakan API dari platform X dan TikTok dengan kata kunci #jembernusantara dan #alunalunjember pada periode 6 Juni 2024 hingga 20 Februari 2025, memperoleh sebanyak 1.420 komentar yang kemudian diproses melalui tahapan cleaning, case folding, tokenizing, normalization, stopword removal, stemming, dan pembobotan fitur menggunakan TF-IDF, serta penyeimbangan kelas dengan SMOTE dan evaluasi model menggunakan K-Fold Cross Validation (K=2 hingga K=10). Hasil pengujian menunjukkan bahwa model Random Forest mampu mengklasifikasikan sentimen dengan akurasi berkisar antara 89% hingga 96%, dengan nilai tertinggi 96% pada K=6, K=8, dan K=10, serta F1-Score terbaik pada kelas negatif mencapai 100%, sehingga disimpulkan bahwa kombinasi metode Random Forest dan SMOTE efektif digunakan untuk analisis sentimen opini publik pada data teks media sosial.

Kata Kunci: Alun-Alun Jember; Analisis Sentimen; K-Fold Cross Validation; Random Forest; SMOTE.

1. LATAR BELAKANG

Perkembangan teknologi informasi menjadikan media sosial sebagai sarana utama masyarakat dalam menyampaikan opini yang dapat diakses secara terbuka dan dalam jumlah besar (Prihanum & Fadillah, 2024). Data dari media sosial memiliki potensi besar untuk dianalisis guna mengetahui kecenderungan opini publik terhadap suatu kebijakan atau program pembangunan, salah satunya melalui analisis sentimen yang mengklasifikasikan opini ke dalam kategori positif, negatif, dan netral.

Algoritma Random Forest merupakan salah satu metode pembelajaran mesin berbasis ensemble yang terbukti efektif dalam analisis sentimen karena mampu menangani data berukuran besar dan menghasilkan akurasi klasifikasi yang tinggi (Indrayanto et al., 2023). Berbagai penelitian terdahulu menunjukkan keberhasilan *Random Forest* dalam analisis sentimen, seperti pada ulasan pengguna aplikasi Dana (Larasati et al., 2022) serta ulasan aplikasi e-commerce Shopee, Tokopedia, Lazada, dan Blibli dengan kombinasi TF-IDF (Syah et al., 2024). Penelitian lain juga menunjukkan bahwa *Random Forest*, baik secara mandiri maupun dalam pendekatan ensemble, mampu mencapai performa tinggi dan stabil, termasuk pada data tidak seimbang (Danuraga et al., 2025; Jlifi et al., 2024; Khan et al., 2024).

Berdasarkan temuan tersebut, *Random Forest* dinilai relevan untuk digunakan dalam analisis sentimen opini masyarakat terhadap pembangunan Alun-Alun Jember, karena mampu memberikan gambaran tingkat kepuasan publik yang dapat dimanfaatkan dalam evaluasi kebijakan. Dalam penelitian ini, Python dipilih sebagai bahasa pemrograman karena didukung berbagai pustaka analisis data dan pemrosesan teks serta umum digunakan dalam penelitian data mining (Pratama, 2019). Pengumpulan data dilakukan menggunakan teknik web scraping dengan bantuan pustaka seperti BeautifulSoup dan Selenium untuk memperoleh data komentar dan ulasan secara terstruktur. Pendekatan ini telah banyak diterapkan dalam penelitian analisis sentimen di Indonesia dan terbukti mampu menghasilkan model klasifikasi dengan tingkat akurasi yang baik (Saputra & Rahim, 2025).

Meskipun demikian, sebagian besar penelitian sebelumnya masih berfokus pada konteks aplikasi komersial dan e-commerce, sementara kajian analisis sentimen terhadap opini masyarakat mengenai pembangunan fasilitas publik, khususnya di tingkat daerah, masih relatif terbatas. Selain itu, permasalahan ketidakseimbangan data sentimen sering kali menjadi tantangan dalam analisis opini publik, yang dapat memengaruhi kinerja model klasifikasi. Oleh karena itu, diperlukan pendekatan yang mengombinasikan algoritma *Random Forest* dengan teknik penyeimbangan data seperti Synthetic Minority Oversampling Technique (SMOTE) serta evaluasi model yang komprehensif menggunakan K-FOLD Cross Validation untuk menghasilkan model yang lebih optimal dan tidak bias.

Penelitian ini tidak hanya berfokus pada identifikasi distribusi sentimen masyarakat terhadap pembaruan Alun-Alun Jember, tetapi juga pada pengukuran performa algoritma *Random Forest* dalam melakukan klasifikasi sentimen. Evaluasi model dilakukan menggunakan parameter akurasi, precision, recall, dan F1-score untuk mengetahui kemampuan model dalam mengklasifikasikan opini publik secara tepat dan konsisten.

Adapun ruang lingkup penelitian dibatasi pada komentar berbahasa Indonesia yang diperoleh dari platform X dan TikTok selama periode 6 Juni 2024 hingga 20 Februari 2025. Data dikumpulkan menggunakan kata kunci dan tagar #jembernusantara dan #alunalunjember dengan total sebanyak 1.420 data komentar yang kemudian diklasifikasikan ke dalam tiga kelas sentimen, yaitu positif, negatif, dan netral. Pembatasan ini dilakukan agar penelitian lebih terarah serta hasil analisis yang diperoleh lebih relevan terhadap konteks pembaruan Alun-Alun Jember.

Berdasarkan uraian tersebut, penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap pembaruan Alun-Alun Jember menggunakan algoritma Random Forest berbasis TF-IDF dengan penerapan SMOTE dan validasi silang K-FOLD. Hasil penelitian diharapkan dapat memberikan gambaran yang akurat mengenai persepsi publik serta menjadi masukan yang bermanfaat bagi evaluasi dan pengambilan keputusan dalam perencanaan dan pengelolaan kebijakan pembangunan fasilitas publik.

2. KAJIAN TEORITIS

Analisis Sentimen

Analisis sentimen merupakan metode yang digunakan untuk memahami, mengukur, dan menginterpretasikan sikap, opini, atau perasaan yang terkandung dalam data teks. Teknik ini banyak diterapkan pada sumber data seperti media sosial, ulasan produk, dan percakapan publik untuk mengetahui opini masyarakat terhadap suatu isu. Secara teoritis, analisis sentimen mengelompokkan teks ke dalam kategori positif, negatif, atau netral melalui pendekatan linguistik maupun algoritma pembelajaran mesin (Suratnoaji et al., 2019). Perkembangan teknologi pemrosesan bahasa alami menjadikan analisis sentimen semakin relevan dalam penelitian sosial dan berbagai bidang lainnya.

Dalam praktiknya, analisis sentimen digunakan untuk mengungkap pola perilaku masyarakat, menganalisis respons terhadap kebijakan publik, serta memahami dinamika opini di media sosial. Penelitian di Indonesia menunjukkan bahwa analisis sentimen banyak dimanfaatkan dalam kajian kebijakan publik, citra politik, dan preferensi konsumen (Rahmawati, 2025). Selain itu, analisis sentimen juga digunakan untuk mendeteksi isu viral dan mendukung pengambilan keputusan berbasis data secara real-time (Prihanum & Fadillah, 2024). Dengan dukungan algoritma modern, analisis sentimen menjadi alat yang efektif untuk memahami persepsi masyarakat di era digital (Zulkifli, 2025).

Random Forest

Random Forest merupakan algoritma pembelajaran mesin berbasis ensemble yang banyak digunakan dalam analisis sentimen untuk mengklasifikasikan opini pada teks seperti ulasan produk dan media sosial. Penelitian (Manullang et al., 2023) menunjukkan bahwa Random Forest mampu menghasilkan klasifikasi yang akurat setelah melalui tahapan preprocessing teks, seperti tokenizing, stemming, dan stopword removal. Penelitian lain juga membuktikan bahwa Random Forest memiliki akurasi yang kompetitif dibandingkan algoritma lain seperti Support Vector Machine dan Naïve Bayes dalam analisis sentimen e-commerce (Mustofa et al., 2024).

Meskipun demikian, penerapan Random Forest menghadapi tantangan dalam pemilihan fitur dan optimasi parameter. Kinerja algoritma ini sangat dipengaruhi oleh kualitas preprocessing dan pengaturan parameter model. Random Forest juga terbukti efektif dalam menangani data tidak seimbang melalui strategi sampling dan pendekatan pembelajaran sensitif biaya (Ishak et al., 2024). Secara keseluruhan, Random Forest merupakan algoritma yang andal, tahan terhadap overfitting, serta mampu meningkatkan akurasi prediksi melalui mekanisme bootstrap dan voting mayoritas (Manullang et al., 2023).

Metode Random Forest meliputi tiga tahapan utama, yaitu pembentukan pohon keputusan menggunakan kriteria Gini Impurity atau Entropy, proses pemisahan node berdasarkan information gain, serta penggabungan hasil prediksi melalui mekanisme ensembling dan voting mayoritas (Firdaus et al., 2025).

Tahap Preprocessing

Tahap preprocessing merupakan langkah penting dalam mempersiapkan data sebelum dilakukan pemodelan menggunakan algoritma machine learning. Tujuan preprocessing adalah meningkatkan kualitas data sehingga hasil klasifikasi menjadi lebih optimal (Berutu et al., 2023). Tahapan preprocessing meliputi data cleaning untuk menghapus data duplikat dan karakter tidak relevan, case folding untuk menyeragamkan huruf, tokenizing untuk memecah teks menjadi kata, stopword removal untuk menghilangkan kata umum, stemming untuk mengubah kata ke bentuk dasar, serta normalisasi untuk menyamakan kata tidak baku. Setelah itu, data direpresentasikan dalam bentuk numerik menggunakan pembobotan kata dan dibagi menjadi data latih dan data uji.

TF-IDF (Term Frequency–Inverse Document Frequency)

TF-IDF merupakan metode pembobotan kata yang digunakan untuk menilai relevansi suatu kata dalam dokumen terhadap keseluruhan korpus. Nilai TF-IDF diperoleh dari hasil perkalian antara Term Frequency (TF) dan Inverse Document Frequency (IDF), sehingga kata

yang sering muncul dalam dokumen tertentu namun jarang muncul dalam dokumen lain akan memiliki bobot lebih tinggi (Adhiatma & Qoiriah, 2022). Metode ini efektif dalam mengurangi pengaruh kata umum yang memiliki bobot informasi rendah.

Penelitian menunjukkan bahwa penggunaan TF-IDF mampu meningkatkan akurasi klasifikasi teks berita hingga 87,5% dibandingkan metode bag-of-words (Soyusiawaty & Putra, 2023). Namun, TF-IDF memiliki keterbatasan dalam menangkap konteks semantik karena bersifat statistik. Oleh karena itu, beberapa penelitian menggabungkan TF-IDF dengan metode pembelajaran mendalam untuk menghasilkan representasi teks yang lebih kontekstual (Masuzzahra et al., 2025). Meskipun demikian, TF-IDF tetap menjadi metode yang efisien dan banyak digunakan dalam analisis teks.

K-FOLD Cross Validation dan Confusion Matrix

Confusion matrix merupakan alat evaluasi yang digunakan untuk mengukur kinerja model klasifikasi dengan membandingkan label aktual dan label prediksi. Confusion matrix terdiri dari True Positive, True Negative, False Positive, dan False Negative yang menunjukkan tingkat keberhasilan dan kesalahan model dalam melakukan klasifikasi (Nuraeni et al., 2024). Evaluasi menggunakan confusion matrix sangat penting untuk menilai akurasi dan keandalan model, terutama dalam analisis sentimen multi-kelas (Nuraeni et al., 2024).

Synthetic Minority Oversampling Technique (SMOTE)

Synthetic Minority Oversampling Technique (SMOTE) merupakan teknik oversampling yang digunakan untuk mengatasi ketidakseimbangan data pada pemodelan machine learning dengan cara menghasilkan data sintesis pada kelas minoritas agar distribusi kelas menjadi lebih seimbang. Metode ini terbukti mampu meningkatkan performa model klasifikasi pada dataset dengan distribusi kelas tidak merata (Syukron et al., 2020).

Berbeda dengan oversampling konvensional yang berisiko menimbulkan overfitting akibat duplikasi data, SMOTE menghasilkan data baru melalui interpolasi antar sampel kelas minoritas menggunakan konsep k-nearest neighbors (k-NN), sehingga data sintesis yang dihasilkan lebih variatif (Syukron et al., 2020). Namun, SMOTE memiliki keterbatasan berupa meningkatnya kompleksitas komputasi, potensi ketidaksesuaian data sintesis dengan distribusi asli, serta sensitivitas terhadap noise dan outlier pada dataset awal (Anjani et al., 2023).

Confusion Matrix

Confusion matrix merupakan tabel evaluasi kinerja model klasifikasi yang membandingkan nilai aktual dengan hasil prediksi model. Confusion matrix terdiri dari True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN), yang menggambarkan prediksi benar dan salah dari model klasifikasi (Nuraeni et al., 2024). Dalam

klasifikasi multikelas, confusion matrix digunakan untuk mengevaluasi kinerja prediksi pada kelas negatif, netral, dan positif (Nuraeni et al., 2024). Nilai dari confusion matrix selanjutnya digunakan untuk menghitung metrik evaluasi seperti accuracy, precision, recall, dan F1-score (Argina, 2020). Accuracy mengukur proporsi prediksi benar terhadap keseluruhan data, precision mengukur ketepatan prediksi positif, recall mengukur kemampuan model menangkap data positif, sedangkan F1-score merupakan rata-rata harmonis antara precision dan recall yang efektif untuk data tidak seimbang.

Text Mining

Text mining adalah proses otomatis untuk mengekstraksi informasi dari data teks tidak terstruktur dengan memanfaatkan Natural Language Processing (NLP), pembelajaran mesin, dan analisis statistik. Teknik ini banyak digunakan dalam analisis sentimen, klasifikasi teks, dan pengelompokan dokumen, terutama pada data digital yang terus meningkat (Agung & Santara, 2024). Dalam klasifikasi teks multilabel, text mining memungkinkan satu dokumen memiliki lebih dari satu label, misalnya pada survei kepuasan yang mencakup berbagai aspek layanan. Salah satu metode yang sering digunakan adalah Support Vector Machine (SVM), yang terbukti efektif dalam menangani kompleksitas data teks (Sulaiman et al., 2025).

Google Colaboratory

Google Colaboratory (Colab) merupakan platform berbasis cloud dari Google Research yang memungkinkan pengguna menjalankan kode Python langsung melalui peramban tanpa instalasi tambahan. Platform ini menyediakan akses gratis ke GPU dan TPU sehingga mendukung proses analisis data dan pembelajaran mesin secara efisien (Guntara, 2023). Meskipun memiliki keterbatasan dalam stabilitas dan kapasitas sumber daya, Google menyediakan layanan Colab Pro untuk performa yang lebih optimal. Google Colab dinilai memiliki potensi besar dalam meningkatkan aksesibilitas pembelajaran pemrograman karena mudah digunakan, mendukung kolaborasi daring, dan terintegrasi dengan teknologi cloud, khususnya dalam konteks pendidikan dan pengembangan model machine learning (Guntara, 2023).

Apify

Apify merupakan platform web scraping yang memungkinkan pengambilan data media sosial seperti Facebook, TikTok, dan Instagram secara otomatis, termasuk data like, komentar, hashtag, dan engagement rate untuk keperluan analisis sentimen dan bisnis (Kilay & Radianto, 2024). Pemanfaatan Apify mendukung pengumpulan data besar (big data) secara terstruktur, sehingga memudahkan analisis perilaku pengguna dan pengambilan keputusan berbasis data. Penelitian sebelumnya menunjukkan bahwa otomatisasi data menggunakan Apify mampu

meningkatkan efektivitas evaluasi kinerja, inovasi produk, serta adaptasi strategi bisnis di era digital (Jufri & Sonni, 2025).

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode analisis sentimen untuk mengkaji opini masyarakat terhadap pembaruan Alun-Alun Jember. Data penelitian diperoleh dari media sosial melalui teknik web scraping dengan bantuan bahasa pemrograman Python. Data yang telah dikumpulkan kemudian melalui tahapan preprocessing yang meliputi data cleaning, case folding, tokenizing, stopword removal, stemming, serta normalisasi. Tahapan ini dilakukan untuk meningkatkan kualitas data teks sehingga siap digunakan dalam proses pemodelan. Selanjutnya, data teks diubah ke dalam bentuk numerik menggunakan metode pembobotan kata TF-IDF agar dapat diproses oleh algoritma pembelajaran mesin.

Tahap pemodelan dilakukan dengan menerapkan algoritma Random Forest yang dikombinasikan dengan teknik Synthetic Minority Oversampling Technique (SMOTE) untuk mengatasi permasalahan ketidakseimbangan data. Pada penelitian ini, proses pembentukan model Random Forest dilakukan dengan membangun sejumlah pohon keputusan (*decision tree*) yang dibentuk secara acak melalui mekanisme bootstrap sampling dan pemilihan subset fitur secara acak pada setiap node. Prediksi akhir diperoleh melalui mekanisme majority voting dari seluruh pohon keputusan sehingga menghasilkan klasifikasi sentimen yang lebih stabil dan tahan terhadap overfitting.

Sebelum proses pelatihan model dilakukan, dataset terlebih dahulu diseimbangkan menggunakan metode SMOTE untuk meningkatkan representasi kelas minoritas. Data sintetis dibangkitkan berdasarkan kedekatan antar sampel menggunakan pendekatan k-nearest neighbors sehingga distribusi kelas menjadi lebih proporsional dan model tidak cenderung bias terhadap kelas mayoritas.

Evaluasi kinerja model dilakukan menggunakan metode K-FOLD Cross Validation dengan variasi nilai K guna memperoleh hasil yang lebih stabil dan objektif. Pada setiap iterasi K-FOLD, dataset dibagi menjadi K bagian (*fold*) dengan satu fold digunakan sebagai data pengujian dan K-1 fold lainnya digunakan sebagai data pelatihan. Proses ini dilakukan secara berulang hingga seluruh fold memperoleh kesempatan menjadi data uji sehingga seluruh data dapat berperan sebagai data pelatihan maupun pengujian secara bergantian.

Penelitian ini menggunakan variasi nilai K mulai dari K=2 hingga K=10 untuk mengetahui pengaruh jumlah fold terhadap stabilitas performa model dan kemampuan generalisasi algoritma Random Forest. Penggunaan rentang tersebut juga bertujuan untuk

memperoleh konfigurasi validasi yang menghasilkan performa paling optimal berdasarkan nilai accuracy, precision, recall, dan F1-score pada masing-masing kelas sentimen. Berdasarkan hasil pengujian, peningkatan nilai K cenderung menghasilkan performa model yang lebih baik dan stabil, dengan akurasi terbaik sebesar 96% yang diperoleh pada K=6, K=8, dan K=10.

Kinerja model dievaluasi menggunakan confusion matrix serta metrik accuracy, precision, recall, dan F1-score pada setiap kelas sentimen. Hasil evaluasi tersebut digunakan untuk menilai efektivitas model dalam mengklasifikasikan sentimen masyarakat secara akurat dan konsisten.

4. HASIL DAN PEMBAHASAN

Pengambilan Data dengan Scraping

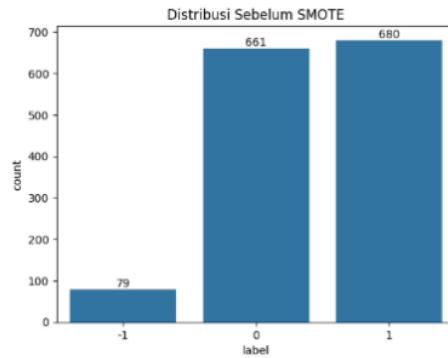
Pengumpulan data dalam penelitian ini dilakukan menggunakan teknik scraping melalui Google Colab dengan bantuan API dari platform X dan TikTok. Tujuan scraping adalah memperoleh data opini publik berupa komentar terkait pembaruan Alun-Alun Jember. Kata kunci yang digunakan berupa tagar #jembernusantara dan #alunalunjember, dengan rentang waktu pengambilan data dari 6 Juni 2024 hingga 20 Februari 2025. Proses scraping dilakukan secara otomatis untuk menjangkau data yang relevan dan kontekstual. Dari proses tersebut diperoleh 1.420 data komentar, yang selanjutnya disimpan dalam format CSV dan digunakan sebagai dasar pada tahap preprocessing.



Gambar 1. Scraping Data.

Pelabelan Data

Pada penelitian ini, pelabelan data dilakukan secara manual oleh Bapak Agus Widi Priyono, S.Pd., guru di SDN Pekalangan 1 Bondowoso diklasifikasikan ke dalam tiga kelas sentimen, yaitu positif sebanyak 680 data, negatif sebanyak 79 data, dan netral sebanyak 661 data.



Gambar 2. Label Sebelum Menggunakan SMOTE.

Text Preprocessing

Tahap text preprocessing bertujuan membersihkan dan menyiapkan data teks agar siap digunakan dalam proses analisis. Tahapan yang dilakukan meliputi cleaning, case folding, tokenizing, normalization, stopwords removal, dan stemming.

Cleaning

Untuk membuat konten menjadi lebih rapi, prosedur ini menghapus karakter-karakter yang tidak perlu, termasuk tanda baca, angka, URL, sebutan, tagar, dan emoji.

Case Folding

Pada proses ini Case Folding mengubah huruf dalam teks menjadi huruf kecil untuk menyeragamkan format kata.

Tokenizing

Tahap ini memecah kalimat atau teks menjadi bagian-bagian kata atau token agar lebih mudah diproses.

Normalization

Proses mengubah kata tidak baku, singkatan, atau bahasa tidak formal menjadi bentuk baku sesuai kaidah bahasa.

Stopword Removal

Selanjutnya menghapus kata-kata umum yang tidak memiliki makna penting dalam analisis, seperti “dan”, “di”, atau “yang”.

Stemming

Tahap terakhir Stemming mengubah kata turunan menjadi bentuk dasarnya untuk menyeragamkan berbagai variasi kata yang memiliki makna sama..

Penerapan TF-IDF

Ekstraksi fitur dilakukan menggunakan metode TF-IDF (Term Frequency–Inverse Document Frequency) untuk mengubah data teks menjadi representasi numerik yang dapat diproses oleh algoritma klasifikasi.

Term Frequency (TF)

TF mengukur frekuensi kemunculan suatu kata dalam dokumen. Semakin sering kata muncul, semakin tinggi nilai TF-nya.

Tabel TF:

	a	abi	abis	absen	ac	acara	ad	ada	adain	adem	...	yarabal	yasile	yaudah	ye	ying	yo	yok	yuk	yummy	yup
D1	0	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0	0	0	0	0	0
D2	0	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0	0	0	0	0	0
D3	0	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0	0	0	0	0	0
D4	0	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0	0	0	0	0	0
D5	0	0	0	0	0	0	0	0	0	0	...	0	0	0	0	0	0	0	0	0	0

5 rows x 1991 columns

Gambar 3. Hasil term Frequency.

IDF (Inverse Document Frequency)

IDF digunakan untuk mengukur seberapa jarang kata tersebut muncul di seluruh dokumen. Semakin jarang sebuah kata muncul di dokumen lain, semakin tinggi nilai IDFnya.

Tabel IDF:

	IDF
a	2.240383
abi	3.416474
abis	2.571376
absen	1.884995
ac	2.462232

Gambar 4. Hasil Inverse Document Frequency.

TF-IDF (Term Frequency-Inverse Document Frequency)

Teknik pembobotan kata TF-IDF (Term Frequency–Inverse Document Frequency) digunakan untuk mengubah data teks menjadi representasi numerik berdasarkan tingkat kepentingannya dalam dokumen. Nilai TF menunjukkan frekuensi kemunculan suatu kata dalam dokumen, sedangkan IDF mengukur tingkat kelangkaan kata tersebut di seluruh dokumen. Kata yang sering muncul dalam satu dokumen namun jarang muncul pada dokumen lain akan memiliki bobot TF-IDF yang lebih tinggi, yang diperoleh dari hasil perkalian TF dan

IDF. Bobot ini kemudian direpresentasikan sebagai vektor numerik dan digunakan sebagai masukan pada algoritma klasifikasi untuk mengidentifikasi kategori data, seperti sentimen.

TF-IDF Non Zero:		TFIDF_nonzero
Document		
0	D1	[0.584323, 0.326471, 0.484838, 0.567587]
1	D2	[0.078503, 0.125212, 0.263832, 0.142525, 0.193...
2	D3	[0.256902, 0.134837, 0.214895, 0.056457, 0.214...
3	D4	[0.087226, 0.142904, 0.152787, 0.167554, 0.205...
4	D5	[0.34377, 0.778861, 0.575633, 0.778861]

Gambar 5. Hasil Term Frequency-Inverse Document Frequency.

Implementasi SMOTE

Pada tahap ini, data yang telah diseimbangkan menggunakan teknik *SMOTE* digunakan dalam penerapan algoritma *Random Forest*. *SMOTE* bertujuan untuk menyeimbangkan distribusi jumlah sampel pada setiap kelas dengan menghasilkan data sintesis dari kelas minoritas. Proses ini dilakukan secara sistematis pada setiap iterasi fold, dimulai dari pembagian data, penerapan *SMOTE* pada data pelatihan, hingga proses pelatihan model *Random Forest*. Selanjutnya, evaluasi kinerja model dilakukan menggunakan data uji untuk mengukur akurasi dan ketahanan klasifikasi.

```
# 6. SMOTE (SEBELUM K-FOLD)
# =====
smote = SMOTE(random_state=42)
X_res, y_res = smote.fit_resample(X, y)

print("\nDistribusi Setelah SMOTE:")
print(Counter(y_res))

# =====
# 7. K-FOLD + RANDOM FOREST
# =====
k_values = [2, 3, 4, 5, 6, 8, 10]

for k in k_values:
    print("\n=====")
    print(f"K-FOLD = {k}")
    print("\n=====")

    kf = KFold(n_splits=k, shuffle=True, random_state=42)

    acc_scores = []

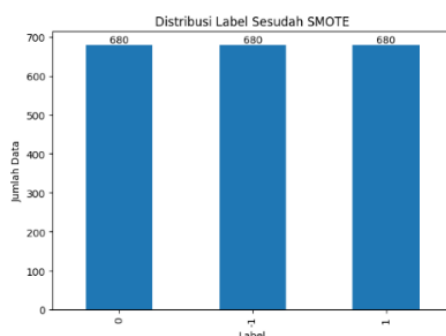
    for fold, (train_idx, test_idx) in enumerate(kf.split(X_res)):
        X_train, X_test = X_res[train_idx], X_res[test_idx]
        y_train, y_test = y_res[train_idx], y_res[test_idx]
```

Gambar 6. Script proses SMOTE dan K-FOLD.

Penyeimbangan data menggunakan SMOTE dalam skema K-FOLD Cross Validation dengan nilai K = 2, 3, 4, 5, 6, 8, dan 10. Pada setiap fold, data sintesis dihasilkan untuk meningkatkan representasi kelas minoritas pada data pelatihan. Pendekatan ini bertujuan mengatasi ketidakseimbangan data agar model tidak bias terhadap kelas mayoritas dan mampu menghasilkan performa klasifikasi yang lebih optimal. Berdasarkan hasil evaluasi pada gambar 5 model Random Forest menunjukkan performa yang baik dengan nilai akurasi berada pada rentang 89% hingga 96%. Nilai akurasi tertinggi sebesar 96% diperoleh pada nilai K=6, K=8, dan K=10. Evaluasi per kelas menunjukkan bahwa kelas negatif (-1) memiliki nilai F1-Score

tertinggi hingga mencapai 100%, yang mengindikasikan bahwa model mampu mengenali pola sentimen negatif secara sangat konsisten.

Sementara itu, kelas netral (0) dan positif (1) menunjukkan nilai precision dan recall yang relatif lebih rendah dibandingkan kelas negatif. Hal ini disebabkan oleh adanya kemiripan konteks bahasa dan ambiguitas makna dalam komentar media sosial, sehingga memengaruhi ketepatan klasifikasi model pada kedua kelas tersebut.



Gambar 7. Distribusi Setelah SMOTE.

Distribusi label setelah SMOTE menunjukkan bahwa jumlah data pada setiap kelas telah seimbang, sehingga proses pengujian selanjutnya dapat dilakukan menggunakan skema K-FOLD Cross Validation dengan variasi nilai K.

Tabel 1. Data SMOTE F1.

Fold	Label	Distribusi Data	SMOTE F1- Score	Precision	Recal	Accuracy
K=2						
Uji 1	-1	1020	99%	99%	99%	91%
	0	1020	88%	80%	98%	
	1	1020	86%	97%	77%	
Uji 2	-1	1020	98%	99%	97%	90%
	0	1020	88%	82%	96%	
	1	1020	85%	94%	78%	
K=3						
Uji 1	-1	680	98%	98%	98%	90%
	0	680	87%	80%	95%	
	1	680	84%	94%	77%	
Uji 2	-1	680	100%	100%	100%	93%
	0	680	90%	84%	97%	
	1	680	89%	97%	82%	
Uji 3	-1	680	99%	100%	98%	93%
	0	680	91%	85%	97%	
	1	680	89%	96%	83%	
K=4						
Uji 1	-1	510	98%	99%	98%	92%
	0	510	90%	84%	96%	
	1	510	88%	95%	82%	

Fold	Label	Distribusi Data	SMOTE	F1- Score	Precision	Recal	Accuracy
Uji 2	-1	510		99%	99%	100%	
	0	510		88%	82%	96%	92%
	1	510		88%	97%	80%	
Uji 3	-1	510		100%	100%	99%	
	0	510		92%	86%	98%	94%
	1	510		91%	97%	85%	
Uji 4	-1	510		99%	100%	98%	
	0	510		92%	86%	99%	94%
	1	510		89%	97%	83%	
K=5							
Uji 1	-1	408		99%	99%	99%	
	0	408		89%	84%	95%	92%
	1	408		87%	94%	82%	
Uji 2	-1	408		98%	98%	98%	
	0	408		88%	81%	95%	91%
	1	408		87%	94%	81%	
Uji 3	-1	408		99%	99%	99%	
	0	408		91%	86%	98%	94%
	1	408		91%	98%	85%	
Uji 4	-1	408		99%	99%	99%	
	0	408		93%	88%	99%	94%
	1	408		90%	98%	84%	
Uji 5	-1	408		99%	100%	98%	
	0	408		90%	83%	98%	93%
	1	408		88%	96%	81%	
K=6							
Uji 1	-1	340		99%	99%	99%	
	0	340		88%	83%	94%	92%
	1	340		87%	94%	82%	
Uji 2	-1	340		98%	97%	98%	
	0	340		86%	81%	91%	89%
	1	340		82%	89%	77%	
Uji 3	-1	340		100%	100%	99%	
	0	340		95%	91%	98%	96%
	1	340		94%	97%	91%	
Uji 4	-1	340		100%	100%	100%	
	0	340		91%	85%	99%	93%
	1	340		89%	99%	82%	
Uji 5	-1	340		100%	100%	99%	
	0	340		94%	89%	99%	95%
	1	340		93%	98%	88%	
Uji 6	-1	340		99%	100%	98%	
	0	340		91%	85%	97%	94%
	1	340		90%	97%	85%	
K=8							
Uji 1	-1	255		99%	99%	99%	
	0	255		86%	79%	95%	90%
	1	255		84%	94%	76%	

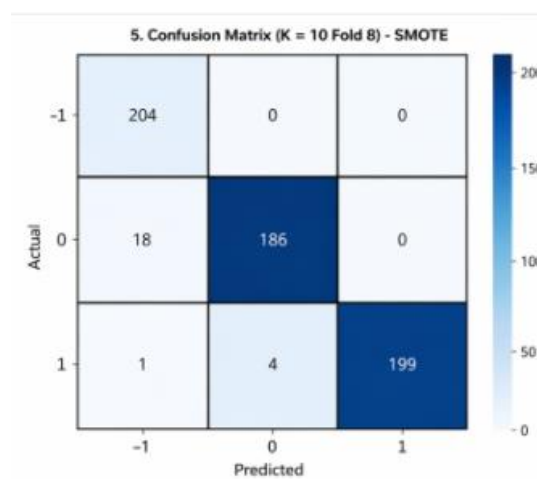
	Fold	Label	Distribusi Data	SMOTE	F1- Score	Precision	Recal	Accuracy
		-1	255		99%	100%	97%	
Uji 2	0	255			92%	89%	95%	94%
	1	255			91%	94%	89%	
	-1	255			98%	98%	99%	
Uji 3	0	255			90%	84%	97%	93%
	1	255			90%	97%	83%	
	-1	255			99%	100%	99%	
Uji 4	0	255			93%	89%	97%	95%
	1	255			92%	96%	88%	
	-1	255			100%	100%	100%	
Uji 5	0	255			91%	86%	97%	93%
	1	255			90%	96%	84%	
	-1	255			100%	100%	100%	
Uji 6	0	255			95%	91%	100%	96%
	1	255			95%	100%	90%	
	-1	255			99%	100%	98%	
Uji 7	0	255			92%	85%	100%	93%
	1	255			87%	97%	80%	
	-1	255			99%	100%	98%	
Uji 8	0	255			89%	83%	95%	92%
	1	255			88%	94%	83%	
K=10								
	-1	204			99%	99%	99%	
Uji 1	0	204			88%	81%	96%	91%
	1	204			84%	94%	77%	
	-1	204			100%	100%	100%	
Uji 2	0	204			92%	88%	95%	95%
	1	204			92%	95%	88%	
	-1	204			99%	100%	98%	
Uji 3	0	204			88%	81%	95%	90%
	1	204			86%	94%	80%	
	-1	204			99%	99%	99%	
Uji 4	0	204			92%	86%	98%	94%
	1	204			91%	97%	86%	
	-1	204			99%	100%	97%	
Uji 5	0	204			94%	88%	100%	95%
	1	204			93%	98%	88%	
	-1	204			100%	100%	100%	
Uji 6	0	204			91%	86%	96%	93%
	1	204			91%	96%	86%	
	-1	204			99%	100%	98%	
Uji 7	0	204			92%	85%	99%	94%
	1	204			91%	98%	84%	
	-1	204			99%	100%	98%	
Uji 8	0	204			95%	91%	100%	96%
	1	204			93%	98%	88%	
	-1	204			99%	100%	99%	
Uji 9	0	204			92%	86%	100%	95%

Fold	Label	Distribusi Data	SMOTE	F1- Score	Precision	Recal	Accuracy
	1	204		91%	100%	84%	
	-1	204		99%	100%	98%	
Uji 10	0	204		90%	85%	96%	93%
	1	204		89%	95%	84%	

Berdasarkan hasil pengujian, model menunjukkan performa terbaik pada kelas negatif (-1) dengan Nilai F1-Score mencapai 100%, pada berbagai nilai K yang menunjukkan bahwa model mampu mengenali komentar negatif dengan sangat baik karena pola kata yang lebih eksplisi, yang menandakan tingkat akurasi dan konsistensi yang sangat baik. Untuk kelas netral (0), performa meningkat seiring bertambahnya nilai K, dengan F1-Score tertinggi sebesar 95% pada K = 6, 8, dan 10. Sementara itu, kelas positif (1) juga menunjukkan kinerja yang baik dengan F1-Score maksimum sebesar 95% pada K = 8.

Secara keseluruhan, nilai K = 6, 8, dan 10 menghasilkan akurasi tertinggi sebesar 96%, yang mengindikasikan bahwa penggunaan jumlah fold yang lebih besar mampu meningkatkan kualitas pelatihan dan keandalan model dalam mengklasifikasikan sentimen.

Pengujian menggunakan algoritma Random Forest yang dikombinasikan dengan SMOTE dan K-FOLD Cross Validation menunjukkan bahwa kinerja model cenderung meningkat seiring bertambahnya nilai K. Akurasi rata-rata meningkat dari 90,83% pada K = 2 menjadi 93,48% pada K = 10. Meskipun kelas netral memiliki nilai recall yang cenderung lebih tinggi dibandingkan precision, secara umum nilai precision, recall, dan F1-score pada seluruh kelas berada pada tingkat yang tinggi dan seimbang, menandakan performa model yang stabil dan konsisten.



Gambar 8. Confusion Matrix K = 10.

Confusion matrix pada Gambar 8 menggambarkan perbandingan antara label aktual dan label prediksi model, sebagian besar data berhasil diklasifikasikan dengan benar oleh

model. Kelas negatif memiliki tingkat kesalahan klasifikasi yang sangat rendah, sedangkan beberapa kesalahan masih ditemukan pada kelas netral dan positif. Kondisi ini menunjukkan bahwa ekspresi sentimen negatif dalam komentar cenderung lebih eksplisit dan mudah dikenali oleh model dibandingkan sentimen netral dan positif. Rincian klasifikasi berdasarkan label aktual adalah sebagai berikut:

- a. Sebanyak 199 data positif berhasil diklasifikasikan dengan benar sebagai positif.
- b. Sebanyak 4 data positif salah diklasifikasikan sebagai netral.
- c. Sebanyak 1 data positif salah diklasifikasikan sebagai negatif.
- d. Sebanyak 186 data netral berhasil diklasifikasikan dengan benar sebagai netral.
- e. Sebanyak 0 data netral salah diklasifikasikan sebagai positif.
- f. Sebanyak 18 data netral salah diklasifikasikan sebagai negatif.
- g. Sebanyak 204 data negatif berhasil diklasifikasikan dengan benar sebagai negatif.
- h. Sebanyak 0 data negatif salah diklasifikasikan sebagai positif.
- i. Sebanyak 0 data negatif salah diklasifikasikan sebagai netral.

Hasil penelitian ini sejalan dengan penelitian sebelumnya yang menyatakan bahwa algoritma Random Forest memiliki performa yang kompetitif dibandingkan metode lain seperti Support Vector Machine (SVM) dan Naive Bayes dalam analisis sentimen teks media sosial. Meskipun penelitian ini tidak melakukan perbandingan langsung antar algoritma, hasil evaluasi menunjukkan bahwa Random Forest mampu memberikan hasil klasifikasi yang stabil, khususnya pada data yang tidak seimbang. Hasil analisis sentimen ini memberikan implikasi penting bagi pemerintah daerah Kabupaten Jember dalam mengevaluasi kebijakan renovasi Alun-Alun Jember. Distribusi sentimen yang dihasilkan dapat digunakan sebagai masukan dalam perbaikan fasilitas, peningkatan kenyamanan, serta perumusan kebijakan publik yang lebih responsif terhadap aspirasi masyarakat.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil pengujian menggunakan K-FOLD Cross Validation ($K = 2-10$), algoritma *Random Forest* menunjukkan performa yang baik dengan nilai akurasi berkisar antara 89%–96%, di mana akurasi tertinggi sebesar 96% diperoleh pada $K = 6$, $K = 8$, dan $K = 10$. Dari sisi per kelas, kelas negatif (-1) mencapai performa terbaik dengan F1-Score sebesar 100% pada beberapa fold, yang menunjukkan konsistensi model dalam mengklasifikasikan sentimen negatif. Pada kelas netral (0) dan kelas positif (1), nilai F1-Score tertinggi mencapai 95%, masing-masing diperoleh pada $K = 6$, 8, dan 10 untuk kelas netral serta $K = 8$ untuk kelas positif.

Secara umum, peningkatan nilai K cenderung menghasilkan performa model yang lebih optimal dan stabil. Namun demikian, performa antar kelas belum sepenuhnya merata, di mana kelas negatif menunjukkan hasil paling tinggi, sementara kelas netral dan positif masih memiliki keterbatasan. Hal ini tercermin dari nilai precision dan recall pada kelas netral yang berada pada rentang 79%–91%, serta nilai recall pada kelas positif yang berkisar antara 76%–91%. Selain itu, penggunaan SMOTE berpotensi mempengaruhi tingginya performa pada kelas tertentu, dan metode TF-IDF belum sepenuhnya mampu menangkap konteks semantik teks. Oleh karena itu, meskipun model menunjukkan kinerja yang baik dalam klasifikasi sentimen publik terhadap pembaruan Alun-Alun Jember, pengembangan lebih lanjut masih diperlukan untuk meningkatkan keseimbangan performa antar kelas dan kemampuan generalisasi model.

Didalam Penelitian ini masih terdapat keterbatasan yakni pada potensi noise data media sosial dan keterbatasan representasi platform. Oleh sebab itu, untuk penelitian selanjutnya disarankan menggunakan algoritma pembandingan seperti SVM, Naive Bayes, atau model deep learning.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada semua pihak yang telah memberikan bimbingan, dukungan, dan bantuan selama proses penyusunan penelitian ini. Semoga penelitian ini dapat memberikan manfaat bagi pengembangan ilmu pengetahuan.

DAFTAR REFERENSI

- Adhiatma, F. D., & Qoiriah, A. (2022). Penerapan metode TF-IDF dan deep neural network untuk analisa sentimen pada data ulasan hotel. *Journal of Informatics and Computer Science (JINACS)*, 4, 183–193.
- Agung, S., & Santara, A. (2024). Implementasi text mining untuk analisis review pada aplikasi crowdfunding LX dan ST menggunakan metode sentiment analysis. *LANCAH: Jurnal Inovasi dan Tren*, 2(1), 124–130.
- Anjani, A. F., Anggraeni, D., & Tirta, I. M. (2023). Implementasi *Random Forest* menggunakan SMOTE untuk analisis sentimen ulasan aplikasi Sister for Students UNEJ. *Jurnal Nasional Teknologi dan Sistem Informasi*, 9(2), 163–172.
- Argina, A. M. (2020). Penerapan metode klasifikasi *k-nearest neighbor* pada dataset penderita penyakit diabetes. *Indonesian Journal of Data and Science*, 1(2), 29–33.
- Berutu, S. S., Budiati, H., Jatmika, J., & Gulo, F. (2023). Data preprocessing approach for machine learning-based sentiment classification. *Jurnal Infotel*, 15(4), 317–325.
- Danuraga, I., Heriansyah, R., & Astuti, L. W. (2025). Analisis sentimen opini publik terhadap tren berita menggunakan algoritma *Random Forest*. *Universitas Indo Global Mandiri*.

- Firdaus, R. Z., Wijoyo, S. H., & Purnomo, W. (2025). Analisis sentimen berbasis aspek ulasan pengguna aplikasi Alfragift menggunakan metode *Random Forest* dan pemodelan topik *Latent Dirichlet Allocation*. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 9(2).
- Guntara, R. G. (2023). Visualisasi data laporan penjualan toko online melalui pendekatan data science menggunakan Google Colab. *ULIL ALBAB: Jurnal Ilmiah Multidisiplin*, 2(6), 2091–2100.
- Indrayanto, C. G., Ratnawati, D. E., & Rahayudi, B. (2023). Analisis sentimen data ulasan pengguna aplikasi MyPertamina di Indonesia pada Google Play Store menggunakan metode *Random Forest*. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 7(3), 1131–1139.
- Ishak, D. I., Rahman, Y., & Tupamahu, F. (2024). Early detection of mental health risk indicators in children using machine learning based on teacher questionnaires in Islamic early childhood education in Gorontalo. *JIP: Jurnal Informatika Polinema*. <https://doi.org/10.30603/au.v24i2.6439>
- Jlifi, B., Abidi, C., & Duvallet, C. (2024). Beyond the use of a novel ensemble-based Random Forest-BERT model (Ens-RF-BERT) for the sentiment analysis of the hashtag COVID-19 tweets. *Social Network Analysis and Mining*, 14(1), 88.
- Jufri, I., & Sonni, A. F. (2025). Analisis respon publik dan sentimen video berbagi Willy Salim yang viral di TikTok. *Jurnal Riset Jurnalistik dan Media Digital*, 77–88.
- Khan, T. A., Sadiq, R., Shahid, Z., Alam, M. M., & Su'ud, M. B. M. (2024). Sentiment analysis using support vector machine and *Random Forest*. *Journal of Informatics and Web Engineering*, 3(1), 67–75.
- Kilay, T. N., & Radianto, A. J. V. (2024). Peran *engagement rate* media sosial terhadap intensitas input media sosial dan nilai perusahaan. *Jurnal Akuntansi Neraca*, 2(3).
- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). Analisis sentimen ulasan aplikasi Dana dengan metode *Random Forest*. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 6(9), 4305–4313.
- Manullang, O., Prianto, C., & Harani, N. H. (2023). Analisis sentimen untuk memprediksi hasil calon pemilu presiden menggunakan *lexicon based* dan *Random Forest*. *Jurnal Ilmiah Informatika*, 11(2), 159–169.
- Masuzzahra, T. R., Umam, K., Mustofa, H., & Handayani, M. R. (2025). Hana: An AI chatbot for Islamic jurisprudence on menstruation using SBERT and TF-IDF. *Journal of Applied Informatics and Computing*, 9(3), 1013–1024.
- Mustofa, Y. A., Surya, I., & Idris, K. (2024). Pendekatan ensemble pada analisis sentimen ulasan aplikasi Google Play Store. 6, 181–188.
- Nuraeni, F., Kurniadi, D., & Diazki, M. H. (2024). Algoritma *k-nearest neighbor* pada kasus dataset *imbalanced* untuk klasifikasi kinerja karyawan perusahaan. *Jurnal Teknologi Informasi dan Ilmu Komputer*, 11(3), 557–568.
- Pratama, S. F. (2019). *Analisis sentimen Twitter debat calon presiden Indonesia menggunakan metode fine-grained sentiment analysis* (Skripsi). Universitas Narotama.
- Prihanum, A., & Fadillah, D. (2024). Analisis sentimen tanggapan publik di media sosial Instagram terhadap program CSR "ESG Existence (EXIST)" PT Telkom. *Jurnal Audiens*, 5(4), 565–580.

- Rahmawati, R. (2025). *Kebijakan publik: Analisis teori dan politik*. YPAD Penerbit.
- Saputra, I., & Rahim, R. (2025). Analisis sentimen kinerja lembaga legislatif di Indonesia menggunakan algoritma *Random Forest* berbasis data media sosial X. *TiN Terapan Informatika Nusantara*, 5(10), 613–624. <https://doi.org/10.47065/tin.v5i10.7133>
- Soyusiawaty, D., & Putra, F. G. (2023). Pengembangan chatbot untuk layanan Pimpinan Daerah Muhammadiyah Kota Yogyakarta menggunakan metode *rule-based*. *Jurnal Penerapan Sistem Informasi (Komputer Manajemen)*, 4(2), 354–363.
- Sulaiman, M. H., Muda, N., & Abdul Razak, F. (2025). Analyzing patient complaints in web-based reviews of private hospitals in Selangor, Malaysia, using large language model-assisted content analysis: Mixed methods study. *JMIR Formative Research*, 9, e69075.
- Suratnoaji, C., Nurhadi, N., & Candrasari, Y. (2019). *Metode analisis media sosial berbasis big data*. Sasanti Institute.
- Syah, A., Nurdiansyah, F., & Rahman, A. Y. (2024). Analisis sentimen aplikasi Shopee, Tokopedia, Lazada, dan Blibli menggunakan leksikon dan *Random Forest*. *Jurnal Informatika dan Teknik Elektro Terapan*, 12(3S1).
- Syukron, M., Santoso, R., & Widiharah, T. (2020). Perbandingan metode SMOTE-*Random Forest* dan SMOTE-XGBoost untuk klasifikasi tingkat penyakit hepatitis C pada *imbalanced class data*. *Jurnal Gaussian*, 9(3), 227–236.
- Zulkifli, R. (2025). Analisis sentimen real-time media sosial menggunakan *edge computing* dan Apache Kafka. *Bit-Tech*, 7(3), 1106–1117.