

KLASIFIKASI SENTIMEN PENGUNJUNG KAWASAN WISATA PANTAI TANJUNG LESUNG MENGGUNAKAN ALGORITMA DECISION TREE DAN XGBOOST

Muhammad Habibi^{1*}, Ahmad Hafizh Nazmudin², Muhammad Akbar Deedat³

^{1,2,3}Sistem Informasi Kelautan Universitas Pendidikan Indonesia Kampus di Serang; Jl. Ciracas No.38, Serang, Kec. Serang, Kota Serang, Banten 42116

Keywords:

Analisis Sentimen, *Decision Tree*, *XGBoost*, Pantai Tanjung Lesung, *Google Maps Review*

Correspondent Email:

m.habibi0@upi.edu



Copyright © 2026 The Author(s), Published by Jurnal Informatika dan Teknik Elektro Terapan. This is an open-access article distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0).

Abstrak. Penelitian ini bertujuan mengklasifikasikan sentimen ulasan wisata Pantai Tanjung Lesung pada Google Maps untuk memperoleh informasi mengenai persepsi dan tingkat kepuasan pengunjung. Dua arsitektur machine learning diimplementasikan: Decision Tree sebagai model dasar dan eXtreme Gradient Boosting (XGBoost) sebagai model optimasi. Dataset diperoleh melalui web scraping dan diproses dengan tahapan preprocessing teks, pendekatan leksikon dengan penanganan negasi, serta pembobotan fitur TF-IDF berbasis unigram dan bigram. Untuk mengatasi ketidakseimbangan kelas digunakan Synthetic Minority Over-sampling Technique (SMOTE), dan optimasi hyperparameter dilakukan menggunakan GridSearchCV. Hasil eksperimental menunjukkan XGBoost mencapai akurasi terbaik sebesar 90,76%, lebih tinggi dibanding Decision Tree sebesar 87,95%. Visualisasi WordCloud mengindikasikan ulasan positif didominasi aspek keindahan alam dan kenyamanan. Temuan ini menegaskan bahwa kombinasi preprocessing teks, SMOTE, dan penalaan hyperparameter dapat meningkatkan performa klasifikasi sentimen pada ulasan wisata berbasis Google Maps.

Abstract. This study aims to classify the sentiment of tourist reviews for Tanjung Lesung Beach on Google Maps to obtain insights into visitor perception and satisfaction. Two machine learning architectures were implemented: Decision Tree as a baseline model and eXtreme Gradient Boosting (XGBoost) as an optimized model. The dataset was obtained via web scraping and processed through text preprocessing, a lexicon-based approach with negation handling, and TF-IDF feature weighting using unigrams and bigrams. To address class imbalance, Synthetic Minority Over-sampling Technique (SMOTE) was applied, and hyperparameter optimization was performed using GridSearchCV. Experimental results show that XGBoost achieved the best accuracy of 90.76%, higher than Decision Tree at 87.95%. WordCloud visualizations indicate that positive reviews are dominated by mentions of natural beauty and visitor comfort. These findings demonstrate that combining text preprocessing, SMOTE, and hyperparameter tuning can improve sentiment classification performance for Google Maps-based tourist reviews.

1. PENDAHULUAN

Perkembangan platform digital telah mengubah cara wisatawan membagikan pengalaman perjalanan melalui ulasan daring secara real-time. Salah satu platform ulasan destinasi yang paling banyak digunakan adalah Google Maps Reviews. Data ulasan tersebut

mengandung informasi penting mengenai persepsi, tingkat kepuasan, hingga keluhan wisatawan terhadap suatu destinasi wisata [1]. Informasi ini memiliki nilai strategis, terutama bagi destinasi pariwisata bahari seperti Pantai Tanjung Lesung di Provinsi Banten yang telah ditetapkan sebagai Kawasan Ekonomi Khusus

(KEK). Pemanfaatan teknik Natural Language Processing (NLP) dan Machine Learning memungkinkan proses analisis ulasan dilakukan secara otomatis untuk mendukung pengambilan keputusan berbasis data dalam pengelolaan pariwisata [2].

Analisis semantik pada teks ulasan mampu memberikan pemahaman yang lebih mendalam mengenai kualitas layanan, fasilitas, dan infrastruktur dibandingkan penilaian berbasis peringkat bintang yang cenderung ambigu [3]. Dalam sektor pariwisata, pendekatan analisis sentimen telah digunakan untuk mengevaluasi preferensi wisatawan dan kualitas layanan destinasi secara lebih objektif [4], [5]. Pendekatan ini dinilai lebih representatif karena mampu menangkap opini dan pengalaman pengguna secara langsung melalui isi komentar yang diberikan.

Penelitian analisis sentimen berbasis ulasan wisata telah dilakukan menggunakan berbagai metode, seperti Naïve Bayes [6], [7], Support Vector Machine (SVM) [8], [9], dan Long Short-Term Memory (LSTM) [10], [11]. Selain itu, algoritma berbasis pohon keputusan seperti Decision Tree juga banyak digunakan karena mudah diinterpretasikan dan efisien dalam proses klasifikasi teks [12]. Namun, model pohon keputusan tunggal memiliki kelemahan berupa risiko overfitting dan sensitivitas yang tinggi terhadap data teks berdimensi besar [13]. Untuk meningkatkan performa klasifikasi dan mengurangi kesalahan prediksi, penelitian ini menerapkan algoritma eXtreme Gradient Boosting (XGBoost) sebagai metode ensemble learning berbasis boosting [14]. Algoritma ini mampu memperbaiki kesalahan prediksi secara iteratif dan memiliki mekanisme regularisasi untuk membantu mengurangi overfitting [15].

Selain pemilihan algoritma, klasifikasi sentimen pada ulasan wisata juga menghadapi tantangan berupa ketidakseimbangan data (class imbalance), di mana jumlah ulasan positif umumnya jauh lebih banyak dibandingkan ulasan negatif [2]. Kondisi tersebut dapat menyebabkan model cenderung berpihak pada kelas mayoritas dan kurang optimal dalam mengenali kelas minoritas. Tantangan lain muncul dari karakteristik bahasa Indonesia yang banyak menggunakan kata negasi, seperti “tidak bagus” atau “kurang nyaman”, yang dapat memengaruhi polaritas sentimen apabila tidak ditangani dengan baik. Oleh karena itu,

penelitian ini menerapkan pendekatan lexicon-based dengan negation handling untuk menjaga konsistensi polaritas sentimen, serta menggunakan Synthetic Minority Over-sampling Technique (SMOTE) untuk menyeimbangkan distribusi data latih sebelum proses klasifikasi dilakukan [7].

Penelitian ini bertujuan untuk membandingkan performa model Decision Tree dan XGBoost pada klasifikasi sentimen ulasan Pantai Tanjung Lesung menggunakan kombinasi preprocessing teks, pembobotan TF-IDF berbasis n-gram, SMOTE, dan hyperparameter tuning. Evaluasi dilakukan berdasarkan analisis semantik teks ulasan tanpa menggunakan rating bintang guna mengurangi potensi bias penilaian [3]. Selain melakukan klasifikasi sentimen, penelitian ini juga memvisualisasikan kata kunci dominan menggunakan WordCloud sebagai bahan evaluasi pengelolaan destinasi wisata berbasis data.

2. TINJAUAN PUSTAKA

Bagian ini memaparkan landasan teori dan tinjauan literatur yang mendasari pelaksanaan penelitian. Pembahasan diarahkan secara sistematis mulai dari konsep dasar analisis sentimen pariwisata, relevansi algoritma *Machine Learning* dalam domain kelautan dan pesisir, teknik prapemrosesan teks, hingga mekanisme kerja dari model klasifikasi yang diimplementasikan.

2.1. Analisis Sentimen pada Ulasan Wisata

Sebagai salah satu turunan dari disiplin Natural Language Processing (NLP), analisis sentimen difungsikan guna memetakan dan mengekstraksi polaritas opini publik ke dalam klasifikasi yang spesifik secara otomatis [1]. Dalam sektor pariwisata, analisis sentimen dimanfaatkan untuk memahami persepsi wisatawan berdasarkan ulasan yang diberikan pada platform digital seperti Google Maps Reviews.

Ulasan wisatawan mengandung informasi penting terkait kualitas pelayanan, kebersihan lingkungan, aksesibilitas, fasilitas umum, serta tingkat kepuasan pengunjung terhadap suatu destinasi wisata. Dibandingkan hanya menggunakan rating bintang, analisis teks ulasan dinilai mampu memberikan informasi

yang lebih mendalam mengenai pengalaman wisatawan [3].

Beberapa penelitian sebelumnya telah menerapkan analisis sentimen pada sektor pariwisata. Penelitian oleh Khofifah et al. [1] menggunakan algoritma Naive Bayes untuk menganalisis ulasan tempat wisata pantai di Kabupaten Karawang berdasarkan Google Maps Reviews. Penelitian lain oleh Atmadja [6] melakukan analisis sentimen wisata di Kabupaten Sukabumi menggunakan metode Naive Bayes dan menunjukkan bahwa pendekatan klasifikasi teks mampu mengidentifikasi opini publik terhadap destinasi wisata secara efektif.

Selain itu, penelitian oleh Ipmawati et al. [8] menerapkan algoritma Support Vector Machine (SVM) pada ulasan tempat wisata berbasis Google Maps dan memperoleh performa klasifikasi yang baik pada data ulasan wisata. Penelitian Ramdani dan Cahyana [9] juga membandingkan metode SVM dan Long Short-Term Memory (LSTM) pada analisis sentimen wisata budaya dan menunjukkan bahwa pendekatan *Machine Learning* mampu meningkatkan akurasi klasifikasi opini wisatawan.

Pada domain wisata pantai, Yusuf et al. [10] menerapkan algoritma LSTM untuk menganalisis sentimen pengunjung terhadap fasilitas dan layanan di Pantai Pudak. Sementara itu, Oktaviana et al. [11] menggunakan metode Long Short-Term Memory pada ulasan objek wisata pantai berbasis Google Maps dan menunjukkan bahwa data ulasan digital dapat dimanfaatkan sebagai sumber evaluasi kualitas destinasi wisata.

Berdasarkan penelitian terdahulu tersebut, dapat disimpulkan bahwa analisis sentimen pada ulasan wisata memiliki peran penting dalam membantu pengelola destinasi memahami kebutuhan dan persepsi pengunjung secara lebih objektif.

2.2. Penerapan Machine Learning pada Bidang Pariwisata dan Kelautan

Perkembangan *Machine Learning* telah memberikan kontribusi dalam pengolahan data digital pada berbagai bidang, termasuk pariwisata dan kelautan. Dalam sektor pariwisata, *Machine Learning* digunakan untuk klasifikasi sentimen, prediksi perilaku

wisatawan, serta pengambilan keputusan berbasis data ulasan digital.

Zalukhu dan Lubis [15] menerapkan algoritma XGBoost untuk mengklasifikasikan wisatawan berdasarkan perilaku konsumen dan menunjukkan bahwa metode ensemble learning mampu menghasilkan performa klasifikasi yang baik. Selain itu, penelitian Christanto et al. [16] membandingkan algoritma Decision Tree, SVM, dan XGBoost dalam klasifikasi review hotel TripAdvisor. Hasil penelitian tersebut menunjukkan bahwa algoritma berbasis *boosting* memiliki kemampuan yang kompetitif dalam menangani data teks berdimensi tinggi. Pada bidang kelautan dan pesisir, *Machine Learning* juga telah diterapkan dalam berbagai studi, seperti identifikasi spesies mangrove menggunakan Random Forest [17], prediksi pasang surut air laut menggunakan Support Vector Regression [12], serta klasifikasi kualitas air budidaya ikan menggunakan Support Vector Machine [18]. Penerapan tersebut menunjukkan bahwa algoritma *Machine Learning* memiliki fleksibilitas tinggi dalam mengolah data kompleks pada domain kelautan maupun pariwisata.

2.3. Prapemrosesan Teks (*Text Preprocessing*)

Data ulasan yang diperoleh dari media digital umumnya masih mengandung banyak noise, seperti simbol, angka, emoticon, tanda baca berlebih, serta penggunaan bahasa tidak baku. Oleh karena itu, diperlukan tahapan prapemrosesan teks agar data siap digunakan dalam proses klasifikasi.

Tahapan *preprocessing* teks yang umum digunakan meliputi case folding, cleansing, tokenization, stopword removal, dan stemming. Case folding dilakukan untuk mengubah seluruh huruf menjadi huruf kecil. Cleansing digunakan untuk menghapus karakter yang tidak diperlukan, sedangkan tokenization bertujuan memecah kalimat menjadi kumpulan kata. Stopword removal digunakan untuk menghapus kata-kata umum yang tidak memiliki makna signifikan, sementara stemming bertujuan mengubah kata berimbuhan menjadi bentuk dasar.

Dalam analisis sentimen berbahasa Indonesia, salah satu tantangan utama adalah keberadaan kata negasi seperti “tidak”, “kurang”, atau “bukan”. Jika tidak ditangani

dengan baik, model klasifikasi dapat salah memahami makna suatu kalimat. Sebagai contoh, frasa “tidak bagus” dapat tetap terbaca sebagai sentimen positif apabila sistem hanya mendeteksi kata “bagus”. Oleh karena itu, mekanisme negation handling diperlukan untuk mempertahankan makna kontekstual dari teks ulasan sebelum masuk ke tahap klasifikasi.

2.4. Pembobotan TF-IDF dan Ekstraksi Fitur N-Gram

Agar data teks dapat diproses oleh algoritma *Machine Learning*, dokumen teks perlu diubah menjadi representasi numerik menggunakan teknik ekstraksi fitur. Salah satu metode yang umum digunakan adalah Term Frequency-Inverse Document Frequency (TF-IDF).

TF-IDF digunakan untuk mengukur tingkat kepentingan suatu kata dalam dokumen berdasarkan frekuensi kemunculan kata pada suatu dokumen dan keseluruhan korpus [20]. Semakin sering suatu kata muncul pada dokumen tertentu tetapi jarang muncul pada dokumen lain, maka bobot kata tersebut akan semakin tinggi.

Selain menggunakan unigram, penelitian klasifikasi sentimen juga sering memanfaatkan fitur n-gram, khususnya bigram, untuk memahami konteks dua kata yang berurutan. Penggunaan bigram memungkinkan sistem mengenali frasa seperti “pantai kotor”, “jalan rusak”, atau “pelayanan buruk” yang memiliki makna sentimen lebih jelas dibandingkan kata tunggal.

Namun, penggunaan TF-IDF dan n-gram menghasilkan matriks fitur berdimensi tinggi dan bersifat sparse matrix. Oleh karena itu, pembatasan jumlah fitur menggunakan parameter `max_features` diperlukan agar proses pelatihan model menjadi lebih efisien.

2.5. Penanganan Ketidakseimbangan Data Menggunakan SMOTE

Dalam klasifikasi sentimen, distribusi jumlah data antar kelas sering kali tidak seimbang. Pada data ulasan wisata, jumlah sentimen positif biasanya jauh lebih dominan dibandingkan sentimen negatif maupun netral [2]. Ketidakseimbangan data dapat menyebabkan model klasifikasi cenderung bias terhadap kelas mayoritas dan mengabaikan kelas minoritas.

Guna memitigasi kendala imbalanced data, penelitian ini mengaplikasikan pendekatan sintesis sampel buatan melalui metode Synthetic Minority Over-sampling Technique (SMOTE). SMOTE bekerja dengan menghasilkan data sintetis baru pada kelas minoritas berdasarkan kedekatan antar sampel menggunakan metode k-nearest neighbors.

Dengan penerapan SMOTE, distribusi data latih menjadi lebih seimbang sehingga model klasifikasi dapat mengenali pola pada seluruh kelas sentimen secara lebih optimal. Teknik ini juga membantu meningkatkan performa evaluasi, terutama pada nilai recall dan f1-score kelas minoritas.

2.6. Algoritma Decision Tree

Decision Tree merupakan algoritma klasifikasi berbasis struktur pohon yang membagi data ke dalam beberapa cabang berdasarkan aturan tertentu. Algoritma ini bekerja dengan memilih atribut terbaik sebagai node keputusan untuk memisahkan data hingga menghasilkan kelas akhir.

Keunggulan utama Decision Tree adalah kemampuannya menghasilkan model yang mudah dipahami dan diinterpretasikan. Struktur pohon keputusan memungkinkan pengguna mengetahui proses pengambilan keputusan model secara transparan [19].

Penelitian oleh Tukino dan Hakim [19] menggunakan Decision Tree untuk menganalisis sentimen objek wisata pada Google Maps dan menunjukkan bahwa algoritma ini mampu melakukan klasifikasi opini publik dengan baik. Namun, Decision Tree memiliki kelemahan berupa kecenderungan mengalami overfitting ketika digunakan pada data berdimensi tinggi atau data yang kompleks [20].

2.7. Algoritma XGBoost

Diperkenalkan pertama kali oleh Tianqi Chen dan Carlos Guestrin, metode eXtreme Gradient Boosting (XGBoost) beroperasi sebagai arsitektur ensemble yang mengandalkan teknik komputasi boosting sekuensial [14]. XGBoost bekerja dengan membangun sejumlah pohon keputusan secara bertahap, di mana setiap pohon baru bertujuan memperbaiki kesalahan prediksi dari pohon sebelumnya.

XGBoost memiliki beberapa keunggulan, seperti kemampuan menangani data berdimensi tinggi, proses komputasi yang efisien, serta fitur regularisasi yang membantu mengurangi risiko *overfitting*. Penelitian oleh Setiawan dan Nastiti [21] menunjukkan bahwa *XGBoost* memiliki performa yang kompetitif dibandingkan *SVM* dan *Extra Trees* dalam analisis sentimen aplikasi digital.

Selain itu, penelitian Christanto et al. [16] juga menunjukkan bahwa *XGBoost* mampu memberikan hasil klasifikasi yang baik pada data *review* hotel dibandingkan algoritma klasifikasi lainnya. Performa *XGBoost* dapat ditingkatkan melalui proses hyperparameter tuning untuk memperoleh kombinasi parameter terbaik [22].

2.8. Penelitian Terdahulu dan Research Gap

Berbagai penelitian sebelumnya telah menerapkan analisis sentimen pada sektor pariwisata menggunakan algoritma seperti *Naive Bayes*, *SVM*, *LSTM*, dan *Decision Tree* [1], [3], [6], [8], [9], [10], [11]. Namun, sebagian besar penelitian masih berfokus pada objek wisata umum, aplikasi digital, maupun *review* hotel.

Penelitian yang secara khusus membandingkan performa *Decision Tree* dan *XGBoost* pada data ulasan wisata pantai berbasis *Google Maps* masih relatif terbatas, terutama pada destinasi wisata bahari di Indonesia. Selain itu, penerapan kombinasi *preprocessing* teks, *negation handling*, TF-IDF, *n-gram*, dan *SMOTE* pada klasifikasi sentimen wisata pantai juga belum banyak dibahas secara komprehensif.

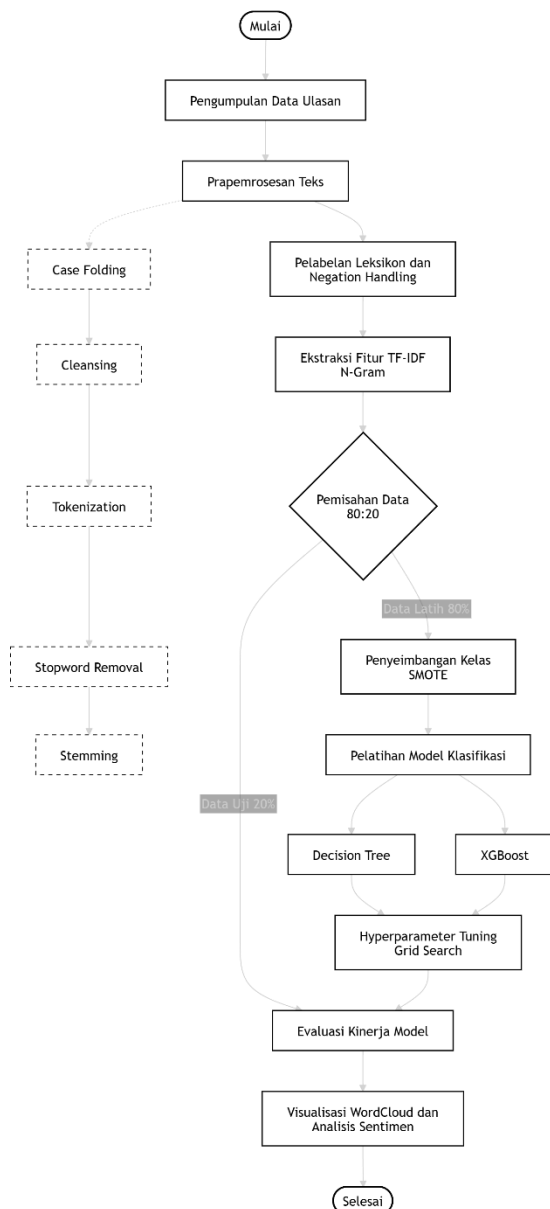
Oleh karena itu, penelitian ini dilakukan untuk menganalisis performa algoritma *Decision Tree* dan *XGBoost* dalam mengklasifikasikan sentimen pengunjung Pantai Tanjung Lesung berdasarkan ulasan *Google Maps Reviews*, dengan memanfaatkan tahapan *preprocessing* teks, ekstraksi fitur TF-IDF, *n-gram*, serta penanganan ketidakseimbangan data menggunakan *SMOTE*.

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif berbasis *Machine Learning* untuk melakukan klasifikasi sentimen terhadap ulasan

pengunjung kawasan wisata Pantai Tanjung Lesung. Tahapan penelitian dilakukan secara sistematis mulai dari pengumpulan data, prapemrosesan teks, ekstraksi fitur, penyeimbangan data, pemodelan klasifikasi, hyperparameter tuning, hingga evaluasi performa model.

Implementasi penelitian dilakukan menggunakan bahasa pemrograman Python pada platform Google Colab dengan dukungan beberapa pustaka seperti Pandas, NumPy, Scikit-learn, NLTK, Sastrawi, Imbalanced-learn, dan XGBoost. Secara keseluruhan, tahapan riset yang dilakukan dalam penelitian ini digambarkan melalui diagram alur pada Gambar 1.



Gambar 1. Diagram Alur Penelitian

3.1. Pengumpulan Data (Web Scraping)

Data yang digunakan dalam penelitian ini merupakan data sekunder berupa ulasan tekstual pengunjung yang diperoleh dari platform Google Maps Reviews kawasan wisata Pantai Tanjung Lesung. Proses pengumpulan data dilakukan menggunakan teknik web scraping berbasis Python untuk mengekstraksi data ulasan secara otomatis.

Data hasil scraping selanjutnya melalui proses seleksi awal dengan menghapus data kosong (missing value), data duplikat, dan data yang tidak relevan agar kualitas dataset tetap terjaga. Penelitian ini berfokus pada analisis isi

teks ulasan tanpa menggunakan rating bintang untuk mengurangi potensi bias penilaian [3], [23].

3.2. Prapemrosesan Teks (Text Processing)

Data tekstual hasil scraping umumnya masih mengandung derau (noise) seperti simbol, angka, tanda baca, karakter khusus, dan penggunaan bahasa tidak baku. Oleh karena itu, dilakukan tahapan prapemrosesan teks untuk membersihkan dan menormalisasi data sebelum digunakan pada proses klasifikasi.

Tahapan prapemrosesan teks yang dilakukan meliputi:

1. **Case Folding:** Mengubah seluruh huruf dalam teks ulasan menjadi huruf kecil (lowercase) serta menghilangkan karakter angka, tanda baca, dan simbol khusus yang tidak memiliki nilai semantik.
2. **Tokenization:** Memotong rangkaian kalimat dalam ulasan menjadi potongan kata-kata tunggal (token) agar dapat dianalisis sebagai entitas terpisah.
3. **Filtering atau Stopword Removal:** Menghapus kata-kata umum yang sering muncul dalam bahasa Indonesia namun tidak memiliki pengaruh terhadap bobot sentimen (seperti "yang", "dan", "di", "dari", "ke"). Proses ini memanfaatkan daftar kata baku (stopword list) standar serta kamus tambahan yang disesuaikan dengan karakteristik ulasan pariwisata.
4. **Stemming:** *Stemming* dilakukan untuk mengubah kata berimbuhan menjadi kata dasar menggunakan library Sastrawi.

Selain *preprocessing teks*, penelitian ini menerapkan pendekatan lexicon-based untuk proses pelabelan sentimen. Negation handling digunakan untuk mendeteksi kata negasi seperti “tidak”, “bukan”, dan “kurang” agar perubahan polaritas sentimen dapat dikenali dengan lebih baik [10]. Hasil pelabelan sentimen kemudian dikelompokkan ke dalam tiga kategori, yaitu positif, netral, dan negatif.

3.3. Ekstraksi Fitur Menggunakan TF-IDF dan N-Gram

$$TF-IDF(t, d) = TF(t, d) \times \log\left(\frac{N}{DF(t)}\right)$$

Persamaan 1. Rumus TF-IDF

TF-IDF digunakan untuk memberikan bobot pada kata berdasarkan frekuensi kemunculannya dalam dokumen dan keseluruhan korpus.

Untuk meningkatkan kemampuan model dalam memahami konteks kalimat, penelitian ini menggunakan kombinasi fitur unigram dan bigram melalui parameter *n-gram range* (1,2). Penggunaan bigram memungkinkan model mengenali hubungan antar kata yang membentuk makna sentimen tertentu, seperti “pantai kotor”, “jalan rusak”, atau “pelayanan buruk”.

Selain itu, dilakukan pembatasan jumlah fitur menggunakan parameter *max_features* untuk mengurangi dimensi matriks fitur yang bersifat sparse serta meningkatkan efisiensi proses komputasi model.

3.4. Pemisahan Dataset dan Penyeimbangan Data

Dataset dibagi menjadi data latih (training data) dan data uji (testing data) menggunakan metode stratified split dengan rasio 80:20. Pendekatan stratified digunakan untuk mempertahankan proporsi distribusi kelas sentimen pada data latih dan data uji.

Permasalahan ketidakseimbangan kelas (class imbalance) ditangani menggunakan metode Synthetic Minority Over-sampling Technique (SMOTE). Teknik SMOTE diterapkan secara khusus pada data latih dengan menghasilkan sampel sintetis baru pada kelas minoritas berdasarkan pendekatan k-nearest neighbors [11].

SMOTE digunakan untuk mengurangi ketidakseimbangan kelas pada data latih sehingga model dapat mengenali kelas minoritas dengan lebih baik.

3.5. Pemodelan Klasifikasi (*Machine Learning*)

Proses pemodelan klasifikasi dilakukan melalui pendekatan komparatif dua tahap untuk membuktikan efektivitas Teknik algoritma *boosting* dalam meningkatkan performa sistem.

1. Model Dasar (*Decision Tree*)

Decision Tree merupakan algoritma klasifikasi berbasis pohon keputusan yang membagi data berdasarkan atribut tertentu menggunakan nilai impurity seperti Gini Index atau Information Gain. Algoritma ini dipilih

karena memiliki interpretabilitas yang baik dan mudah diterapkan pada proses klasifikasi teks.

2. Model Optimasi (*eXtreme Gradient Boosting*)

XGBoost merupakan algoritma *ensemble learning* berbasis *gradient boosting* yang dikembangkan oleh Tianqi Chen dan Carlos Guestrin [1]. Algoritma ini dipilih karena memiliki performa yang baik pada data berdimensi tinggi dan dilengkapi mekanisme regularisasi untuk mengurangi overfitting.

3.6. Hyperparameter Tuning Menggunakan *GridSearchCV*

Untuk memperoleh performa model yang optimal, penelitian ini menerapkan proses hyperparameter tuning menggunakan metode *GridSearchCV*. Metode ini bekerja dengan menguji berbagai kombinasi parameter secara sistematis untuk memperoleh konfigurasi model terbaik berdasarkan hasil evaluasi *cross-validation*.

Pada algoritma *Decision Tree*, parameter yang diuji meliputi:

- *criterion*,
- *max_depth*,
- dan *min_samples_split*.

Sementara itu, pada algoritma *XGBoost* parameter yang diuji meliputi:

- *n_estimators*,
- *max_depth*,
- *learning_rate*,
- dan *subsample*.

Hyperparameter tuning dilakukan menggunakan *GridSearchCV* untuk memperoleh kombinasi parameter terbaik berdasarkan hasil *cross-validation*.

3.7. Evaluasi Model

Evaluasi performa model dilakukan menggunakan data uji yang tidak terlibat dalam proses pelatihan maupun proses SMOTE untuk menjaga objektivitas hasil evaluasi. Pengukuran performa klasifikasi dilakukan menggunakan Confusion Matrix dengan beberapa metrik evaluasi, yaitu accuracy, precision, recall, dan f1-score [16], [23].

- **Akurasi (Accuracy):** Mengukur persentase ketepatan model dalam mengklasifikasikan seluruh ulasan dengan benar.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + EN}$$

Persamaan 2. Rumus Accuracy

- **Precision:** Mengukur rasio ulasan yang benar-benar positif dibandingkan dengan total ulasan yang diprediksi positif oleh model.

$$Precision = \frac{TP}{TP + FP}$$

Persamaan 3. Rumus Precision

- **Recall:** Mengukur kemampuan model dalam menemukan kembali seluruh ulasan yang berlabel positif asli dari dataset.

$$Recall = \frac{TP}{TP + FN}$$

Persamaan 4. Rumus Recall

- **F1-Score:** Nilai rata-rata harmonis yang menyeimbangkan antara komponen precision dan recall.

$$F-1 \text{ Score} = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

Persamaan 5. Rumus F-1 Score

Hasil evaluasi dari kedua algoritma kemudian dibandingkan untuk menentukan model klasifikasi sentimen yang memiliki performa terbaik pada data ulasan pengunjung Pantai Tanjung Lesung.

4. HASIL DAN PEMBAHASAN

Bagian ini menguraikan hasil eksperimen klasifikasi sentimen terhadap ulasan pengunjung Pantai Tanjung Lesung, mulai dari tahapan prapemrosesan, analisis visual sebaran kata, penanganan ketidakseimbangan data, hingga komparasi mendalam antara performa model dasar dan model algoritma *boosting*.

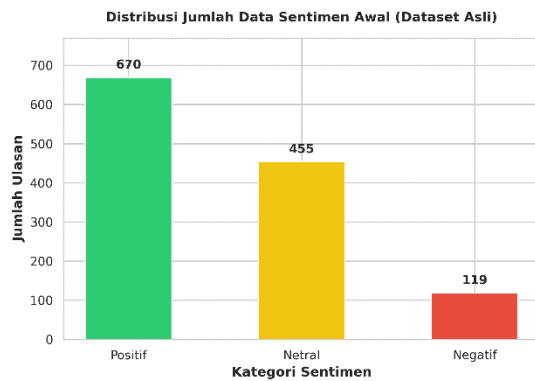
4.1. Pengumpulan Data, Prapemrosesan Teks, dan Pelabelan Leksikon

Tahap awal penelitian ini melibatkan pengumpulan dataset berupa ulasan teks dari pengunjung Kawasan Ekonomi Khusus (KEK) Pantai Tanjung Lesung. Data tersebut diekstraksi secara otomatis dari platform

Google Maps Reviews menggunakan teknik *web scraping*. Dari hasil penarikan data mentah, dilakukan proses penyaringan data (*filtering*) secara ketat, di mana data bernilai kosong (NA) atau ulasan yang hanya memberikan peringkat bintang (*star rating*) tanpa teks deskriptif dihapus dari korpus. Keputusan untuk hanya berfokus pada analisis semantik teks ini selaras dengan temuan literatur terdahulu [3], [23], yang menegaskan bahwa ekstraksi opini murni dari teks memberikan representasi kepuasan maupun keluhan pengunjung yang jauh lebih komprehensif, spesifik, dan objektif dibandingkan penilaian peringkat bintang yang kerap bersifat ambigu. Melalui proses pembersihan awal ini, diperoleh sebanyak 1.244 baris ulasan teks yang valid.

Data teks mentah hasil *scraping* umumnya memiliki tingkat derau (*noise*) yang tinggi akibat struktur kalimat informal, *emoticon*, dan tanda baca yang tidak relevan secara semantik. Oleh karena itu, dataset diproses melalui tahapan prapemrosesan yang meliputi *case folding*, *cleansing*, *tokenization*, *stopword removal*, dan *stemming*. Implementasi tahapan ini penting dalam klasifikasi teks berbahasa Indonesia karena, sebagaimana dibuktikan dalam berbagai studi [1], [19], reduksi derau dan standarisasi kata ke bentuk dasar (*stem*) secara langsung akan menurunkan kompleksitas dimensi komputasi dan secara signifikan meningkatkan kemampuan algoritma dalam mengenali pola teks.

Setelah teks dibakukan, tahapan selanjutnya adalah pelabelan kelas sentimen menggunakan metode berbasis leksikon (*Lexicon-based*). Proses ini bekerja dengan mengalkulasi kemunculan kata kunci (*keywords*) spesifik yang merepresentasikan polaritas positif atau negatif, sementara sisa ulasan yang tidak memenuhi ambang batas diklasifikasikan sebagai kategori Netral. Berdasarkan hasil pelabelan leksikon tersebut, diperoleh sebaran distribusi sentimen pengunjung seperti yang divisualisasikan pada Gambar 2.



Gambar 2. Distribusi Kelas Sentimen Hasil Pelabelan Leksikon

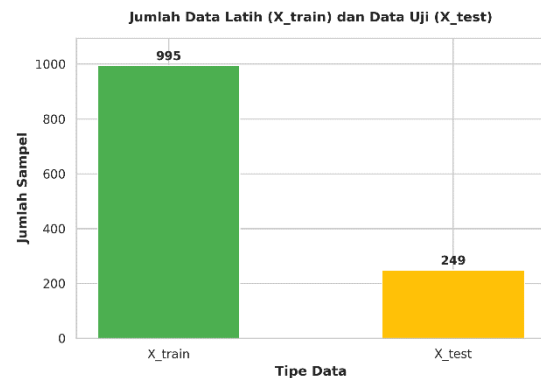
Visualisasi pada Gambar 2 menunjukkan bahwa sebaran kelas didominasi oleh sentimen Positif sebanyak 673 ulasan, diikuti oleh sentimen Netral sebanyak 456 ulasan, sedangkan sentimen Negatif berada pada proporsi terendah dengan hanya 115 ulasan. Kondisi ini secara empiris merepresentasikan tingginya kepuasan publik terhadap KEK Pantai Tanjung Lesung.

Kendati demikian, dari perspektif komputasi *Machine Learning*, distribusi yang terbentuk sangat timpang (*highly imbalanced*). Fenomena ketimpangan data ini sering kali dijumpai pada studi kasus analisis sentimen pariwisata [2]. Berbagai literatur eksperimental menegaskan bahwa dataset yang tidak seimbang akan memicu bias klasifikasi, di mana model akan mengalami *overfitting* terhadap kelas mayoritas (Positif) dan gagal mengenali fitur dari kelas minoritas (Negatif) [20], [21]. Mengingat kelas minoritas justru menyimpan informasi penting berupa keluhan dan titik kelemahan fasilitas wisata, fakta ketimpangan data ini menjadi justifikasi akademis yang kuat mengenai urgensi penerapan teknik *Synthetic Minority Over-sampling Technique* (SMOTE) pada tahapan berikutnya, guna memastikan model memiliki sensitivitas pengenalan yang adil di seluruh tingkatan kelas.

4.2. Pemisahan Data, Ekstraksi Fitur, dan Penyeimbangan Kelas (SMOTE)

Setelah tahapan prapemrosesan dan pelabelan selesai, dataset yang berjumlah 1.244 baris dibagi menjadi data latih (*training data*) dan data uji (*testing data*) menggunakan metode *stratified split* dengan rasio 80:20. Pendekatan

stratified ini digunakan secara ketat untuk menjaga agar proporsi distribusi kelas sentimen pada subset data latih dan data uji tetap konsisten dengan distribusi asli korpus [13]. Berdasarkan pembagian tersebut, diperoleh distribusi data latih sebanyak 995 sampel dan data uji sebanyak 249 sampel, sebagaimana divisualisasikan pada Gambar 3.



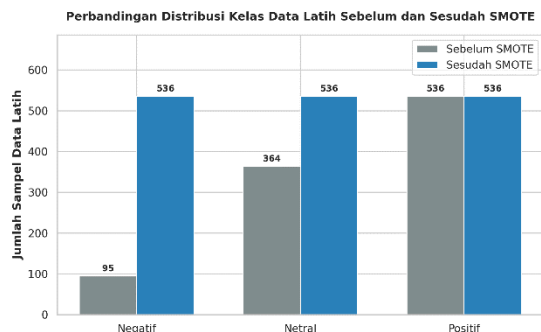
Gambar 3. Distribusi Data Uji dan Data Latih (Rasio 80:20)

Selanjutnya, data ulasan teks ditransformasikan menjadi representasi numerik menggunakan metode pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF). Untuk meningkatkan kemampuan model dalam memahami konteks dan hubungan semantik antarfrasa majemuk yang sering muncul pada ulasan pariwisata, penelitian ini menggunakan kombinasi fitur *unigram* dan *bigram* melalui pengaturan parameter *n-gram range* (1,2) [3], [16]. Penerapan kombinasi ini penting agar model dapat mengenali unit makna yang terdiri dari dua kata, seperti “pantai kotor”, “jalan rusak”, atau “harga mahal”. Sementara itu, pembatasan jumlah fitur maksimum melalui parameter *max_features* dibatasi pada 4.000 fitur terpenting untuk meminimalkan dimensi matriks yang bersifat renggang (*sparse*) serta menghemat beban komputasi.

Distribusi kelas pada data latih (995 sampel) pada awalnya mewarisi masalah ketidakseimbangan kelas (*class imbalance*) yang cukup signifikan, di mana terdapat 536 ulasan Positif, 364 ulasan Netral, dan hanya 95 ulasan Negatif. Berbagai literatur eksperimental menegaskan bahwa kondisi dataset yang tidak seimbang akan memicu bias klasifikasi, di mana algoritma cenderung mengalami *overfitting* pada kelas mayoritas dan mengabaikan karakteristik kelas minoritas [16],

JITET (Jurnal Informatika dan Teknik Elektro Terapan) pISSN: 2303-0577 eISSN: 2830-7062 Habibi dkk

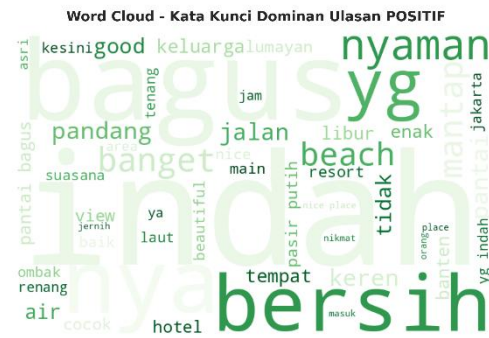
Algoritma SMOTE bekerja secara komputasi dengan mencari tetangga terdekat (*k-nearest neighbors*) pada ruang fitur kelas minoritas, lalu menggenerasi sampel baru secara sintesis di sepanjang garis batas keputusan (*decision boundary*) [4], [24]. Melalui proses penyesuaian ini, jumlah sampel pada kelas Negatif dan Netral ditingkatkan secara merata hingga menyamai batas volume kelas mayoritas (Positif) pada data latih, yaitu sebesar 536 sampel untuk masing-masing kategori sentimen. Hasil akhir penyeimbangan ini meningkatkan total dimensi data latih pasca-SMOTE secara optimal menjadi 1.608 baris data dengan proporsi yang sepenuhnya seimbang (1:1:1), seperti yang ditunjukkan pada Gambar 4.



Gambar 4. Perbandingan Distribusi Kelas Data Latih Sebelum dan Sesudah SMOTE

4.3. Ekstraksi Fitur Teks dan Analisis *Wordcloud*

Setelah data ulasan ditransformasikan ke dalam bentuk matriks fitur menggunakan metode TF-IDF dan melalui proses pelabelan sentimen, tahap selanjutnya adalah visualisasi distribusi kata menggunakan WordCloud. Visualisasi ini digunakan untuk menampilkan kata-kata yang paling sering muncul pada masing-masing kelas sentimen sehingga dapat memberikan gambaran umum mengenai topik yang dominan dalam ulasan pengunjung [3], [25].



Gambar 5. Visualisasi *WordCloud* Distribusi Kata Sentimen Positif

Berdasarkan visualisasi WordCloud pada Gambar 5, kata-kata yang paling dominan pada sentimen positif antara lain “pantai”, “indah”, “bagus”, dan “bersih”. Dominasi kata-kata tersebut menunjukkan bahwa sebagian besar pengunjung memberikan apresiasi terhadap keindahan alam dan kondisi lingkungan wisata Pantai Tanjung Lesung. Temuan ini mengindikasikan bahwa aspek estetika dan kenyamanan lingkungan menjadi faktor utama yang memengaruhi kepuasan wisatawan.



Gambar 6. Visualisasi *WordCloud* Distribusi Kata Sentimen Negatif

Pada visualisasi WordCloud sentimen negatif (Gambar 6), kata-kata seperti “mahal”, “jalan”, “akses”, dan “jauh” muncul dengan frekuensi yang cukup dominan. Kemunculan kata-kata tersebut menunjukkan bahwa beberapa pengunjung masih memberikan perhatian terhadap aspek aksesibilitas dan biaya wisata. Ulasan negatif umumnya berkaitan dengan kondisi infrastruktur jalan menuju lokasi wisata serta persepsi terhadap harga tiket atau fasilitas yang tersedia. Temuan ini dapat menjadi bahan evaluasi bagi pengelola kawasan wisata dalam meningkatkan kualitas layanan dan akses menuju destinasi.



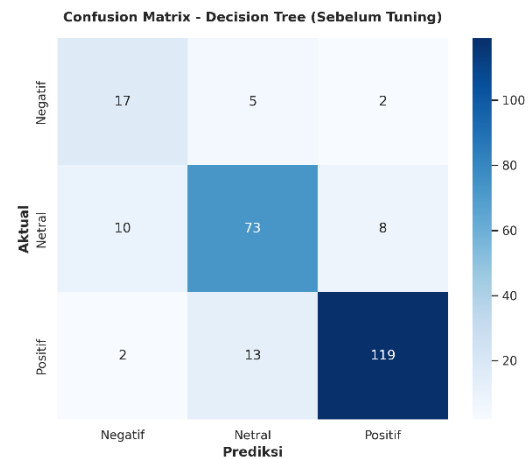
Gambar 7. Visualisasi *WordCloud* Distribusi Kata Sentimen Netral

Sementara itu, *WordCloud* pada kelas sentimen netral (Gambar 7) menunjukkan dominasi kata-kata yang bersifat informatif dan deskriptif, seperti “pantai”, “jalan”, dan “lumayan”. Ulasan pada kategori ini umumnya tidak mengandung ekspresi kepuasan maupun keluhan secara kuat, melainkan berisi informasi umum mengenai kondisi lokasi wisata dan pengalaman perjalanan pengunjung.

Secara keseluruhan, visualisasi *WordCloud* membantu memberikan gambaran awal mengenai persepsi pengunjung terhadap Pantai Tanjung Lesung berdasarkan frekuensi kemunculan kata dalam ulasan. Informasi ini dapat digunakan sebagai pendukung analisis sentimen untuk memahami aspek-aspek yang paling sering dibahas oleh wisatawan.

4.4. Evaluasi Kinerja Model Klasifikasi (Sebelum *Tuning*)

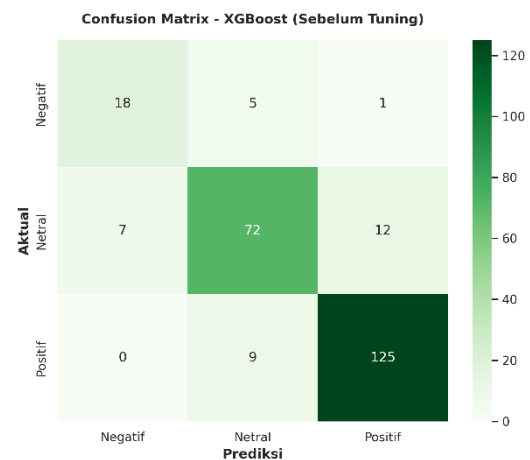
Setelah proses penyeimbangan data menggunakan SMOTE pada data latih selesai dilakukan, tahap berikutnya adalah pelatihan model klasifikasi menggunakan algoritma *Decision Tree* dan *XGBoost*. Pengujian awal dilakukan menggunakan parameter bawaan (default parameter) pada masing-masing algoritma terhadap 249 data uji. Evaluasi performa model dilakukan menggunakan confusion matrix untuk melihat kemampuan model dalam mengklasifikasikan setiap kelas sentimen secara lebih rinci [10].



Gambar 8. *Confusion Matrix* Model *Decision Tree* Sebelum *Tuning*

Berdasarkan confusion matrix pada Gambar 8, model *Decision Tree* mampu mengklasifikasikan 119 data sentimen positif, 73 data sentimen netral, dan 17 data sentimen negatif dengan benar. Namun demikian, masih terdapat beberapa kesalahan klasifikasi antar kelas sentimen, terutama pada kelas netral dan negatif. Beberapa data netral diprediksi sebagai negatif, sedangkan sebagian data positif diprediksi sebagai netral.

Hasil tersebut menunjukkan bahwa model *Decision Tree* masih memiliki keterbatasan dalam memisahkan pola antar kelas pada data teks berdimensi tinggi. Kondisi ini sejalan dengan penelitian sebelumnya yang menyatakan bahwa algoritma pohon keputusan tunggal cenderung sensitif terhadap kompleksitas fitur teks dan berpotensi mengalami *overfitting* pada data sparse.



Gambar 9. *Confusion Matrix* Model *XGBoost* Sebelum *Tuning*

Selanjutnya, proses klasifikasi dilakukan menggunakan algoritma XGBoost sebagai metode ensemble learning berbasis algoritma boosting. Berdasarkan confusion matrix pada Gambar 9, model XGBoost menunjukkan performa yang lebih baik dibandingkan Decision Tree. Model ini mampu mengklasifikasikan 125 data sentimen positif dengan benar dan meningkatkan ketepatan prediksi pada kelas negatif menjadi 18 data [5].

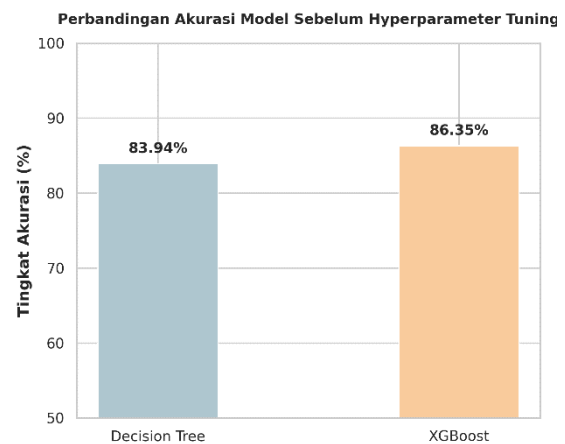
Penurunan jumlah misklasifikasi menunjukkan bahwa pendekatan algoritma boosting pada XGBoost mampu membantu model dalam mempelajari pola data secara lebih optimal. Selain itu, mekanisme iteratif pada XGBoost memungkinkan proses koreksi terhadap kesalahan prediksi dari pohon keputusan sebelumnya sehingga performa klasifikasi menjadi lebih stabil.

Tabel 1. Perbandingan Performa Model Sebelum Hyperparameter

Model	Accuracy	Precision	Recall	F1-Score
Decision Tree	83,94%	0,77	0,80	0,78
XGBoost	86,35%	0,82	0,82	0,82

Berdasarkan Tabel 1, model XGBoost menunjukkan performa yang lebih baik dibandingkan Decision Tree pada seluruh metrik evaluasi sebelum proses hyperparameter tuning dilakukan. Model XGBoost memperoleh accuracy sebesar 86,35% dengan nilai macro average F1-score sebesar 0,82, sedangkan Decision Tree memperoleh accuracy sebesar 83,94% dengan macro average F1-score sebesar 0,78.

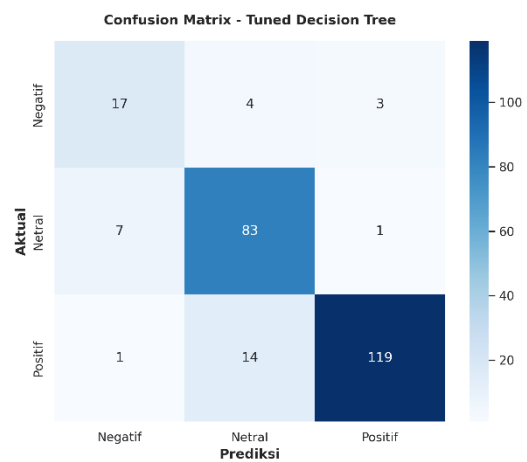
Peningkatan performa pada model XGBoost menunjukkan bahwa pendekatan algoritma boosting mampu membantu model dalam mengenali pola sentimen pada data teks secara lebih baik. Selain itu, nilai precision dan recall pada kelas negatif juga mengalami peningkatan dibandingkan Decision Tree, sehingga model menjadi lebih baik dalam mengklasifikasikan kelas minoritas.



Gambar 10. Perbandingan Model Sebelum Hyperparameter Tuning

4.5. Evaluasi Klasifikasi Model Setelah Hyperparameter Tuning (Grid Search)

Untuk memperoleh performa model yang lebih optimal, penelitian ini menerapkan proses hyperparameter tuning menggunakan metode GridSearchCV. Metode ini bekerja dengan menguji berbagai kombinasi parameter secara sistematis untuk memperoleh konfigurasi terbaik berdasarkan hasil cross-validation [23]. Proses Hyperparameter Tuning diterapkan pada kedua model klasifikasi guna meningkatkan performa model dan mengurangi kesalahan prediksi yang masih terjadi pada tahap pengujian awal (baseline).

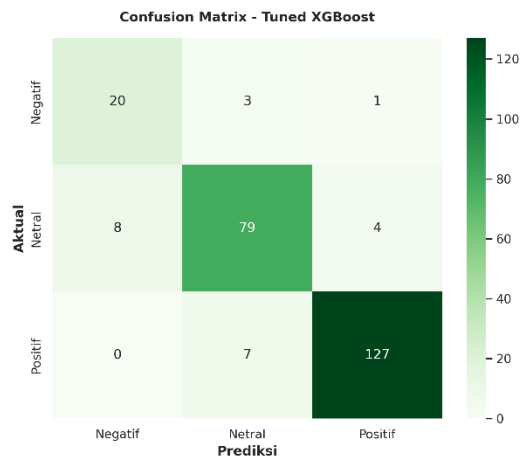


Gambar 11. Confusion Matrix Model Decision Tree Setelah Tuning

Pada algoritma Decision Tree, proses tuning dilakukan dengan mengoptimalkan parameter seperti `max_depth`, `min_samples_split`, dan `min_samples_leaf`. Berdasarkan confusion

matrix pada Gambar 11, proses Hyperparameter Tuning mampu meningkatkan kemampuan klasifikasi model, terutama pada kelas netral. Jumlah prediksi benar pada kelas netral meningkat dari 73 data pada model baseline menjadi 83 data setelah hyperparameter tuning dilakukan.

Selain itu, tingkat misklasifikasi pada kelas positif juga mengalami penurunan. Hasil tersebut menunjukkan bahwa Hyperparameter Tuning mampu membantu model dalam membentuk struktur pohon keputusan yang lebih optimal. Temuan ini sejalan dengan penelitian sebelumnya yang menyatakan bahwa pembatasan kompleksitas pohon keputusan dapat membantu mengurangi overfitting pada model Decision Tree [10].



Gambar 12. Confusion Matrix Model XGBoost Setelah Tuning

Pada model XGBoost, proses hyperparameter tuning dilakukan dengan mengoptimalkan parameter `learning_rate`, `max_depth`, `n_estimators`, dan `subsample`. Berdasarkan hasil evaluasi pada Gambar 12, model Tuned XGBoost menunjukkan peningkatan performa pada seluruh kelas sentimen dibandingkan model baseline.

Model berhasil mengklasifikasikan 127 data sentimen positif dengan benar serta meningkatkan kemampuan deteksi pada kelas negatif menjadi 20 data yang terprediksi dengan benar. Selain itu, jumlah prediksi benar pada kelas netral juga meningkat menjadi 79 data. Peningkatan tersebut menunjukkan bahwa kombinasi parameter yang optimal mampu membantu model dalam mengenali pola sentimen secara lebih baik pada data teks.

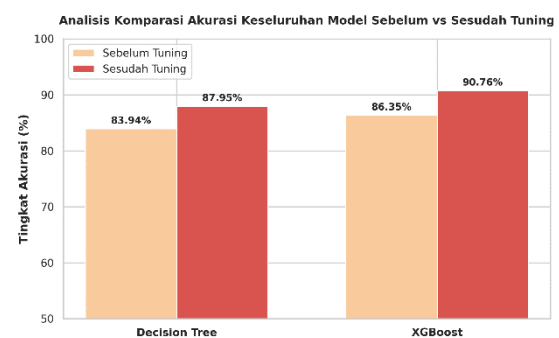
Hasil ini sejalan dengan penelitian sebelumnya yang menyebutkan bahwa XGBoost memiliki kemampuan yang baik dalam menangani data berdimensi tinggi serta dilengkapi mekanisme regularisasi untuk membantu mengurangi overfitting selama proses algoritma boosting [5].

Tabel 2. Perbandingan Performa Model Setelah Hyperparameter Tuning

Model	Accuracy	Precision	Recall	F1-Score
Decision Tree	87,95%	0,82	0,84	0,83
XGBoost	90,76%	0,85	0,88	0,87

Tabel 3. Kombinasi Hyperparameter Terbaik

Model	Parameter Terbaik
Decision Tree	<code>max_depth = 15</code> , <code>min_samples_split = 5</code> , <code>min_samples_leaf = 2</code>
XGBoost	<code>learning_rate = 0,08</code> , <code>max_depth = 12</code> , <code>n_estimators = 450</code> , <code>subsample = 0,9</code>



Gambar 13. Analisis Komparasi Akurasi Keseluruhan Model Sebelum vs Sesudah Tuning.

Berdasarkan Tabel 2 dan Gambar 13, performa kedua model mengalami peningkatan setelah dilakukan proses hyperparameter tuning. Model Tuned XGBoost memperoleh performa terbaik dengan accuracy sebesar 90,76%, sedangkan Tuned Decision Tree memperoleh accuracy sebesar 87,95%.

Peningkatan performa tersebut menunjukkan bahwa kombinasi preprocessing teks, TF-IDF, SMOTE, dan hyperparameter tuning mampu meningkatkan kemampuan

model dalam mengenali pola sentimen pada dataset ulasan wisata. Selain itu, nilai recall pada kelas negatif juga mengalami peningkatan setelah hyperparameter tuning dilakukan, yang menunjukkan bahwa model menjadi lebih baik dalam mengidentifikasi kelas minoritas [5], [23].

Penelitian ini masih memiliki beberapa keterbatasan, seperti jumlah dataset yang relatif terbatas dan penggunaan sumber data yang hanya berasal dari Google Maps. Selain itu, pendekatan lexicon-based masih memiliki keterbatasan dalam memahami konteks kalimat yang kompleks, seperti sarkasme atau opini ambigu. Oleh karena itu, penelitian selanjutnya dapat mengembangkan metode berbasis deep learning atau transformer untuk meningkatkan pemahaman konteks sentimen secara lebih mendalam.

5. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan dapat ditarik beberapa kesimpulan sebagai berikut:

- Algoritma XGBoost menunjukkan performa klasifikasi yang lebih baik dibandingkan Decision Tree pada analisis sentimen ulasan pengunjung Pantai Tanjung Lesung. Setelah dilakukan hyperparameter tuning menggunakan GridSearchCV, model XGBoost memperoleh akurasi sebesar 90,76%, sedangkan Decision Tree memperoleh akurasi sebesar 87,95%.
- Tahapan preprocessing teks, ekstraksi fitur TF-IDF dengan kombinasi unigram dan bigram, serta penerapan SMOTE pada data latih membantu meningkatkan performa model dalam mengenali pola sentimen pada dataset ulasan wisata.
- Penerapan SMOTE berhasil mengurangi ketidakseimbangan kelas pada data latih sehingga model menjadi lebih baik dalam mengklasifikasikan kelas minoritas, terutama pada sentimen negatif.
- Hasil analisis sentimen menunjukkan bahwa mayoritas pengunjung memberikan ulasan positif terhadap Pantai Tanjung Lesung, terutama terkait keindahan alam dan kondisi lingkungan wisata. Sementara itu, ulasan negatif umumnya berkaitan dengan akses jalan, jarak tempuh, dan biaya wisata.

- Visualisasi WordCloud membantu mengidentifikasi kata-kata yang sering muncul pada masing-masing kelas sentimen sehingga dapat memberikan gambaran umum mengenai persepsi pengunjung terhadap kawasan wisata Pantai Tanjung Lesung.

Sebagai pengembangan penelitian selanjutnya, metode klasifikasi dapat dikombinasikan dengan pendekatan deep learning seperti IndoBERT atau transformer-based model lainnya untuk memperoleh performa yang lebih optimal. Selain itu, penelitian berikutnya juga dapat memanfaatkan sumber data dari media sosial lain serta mengintegrasikan analisis spasial berbasis Sistem Informasi Geografis (SIG) untuk memetakan persebaran keluhan maupun kepuasan wisatawan secara lebih detail.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Program Studi Sistem Informasi Kelautan Universitas Pendidikan Indonesia Kampus Serang atas dukungan akademik dan fasilitas penelitian yang telah diberikan selama proses penelitian berlangsung.

DAFTAR PUSTAKA

- [1] W. Khofifah, D. N. Rahayu, and A. M. Yusuf, "Analisis Sentimen Menggunakan Naive Bayes Untuk Melihat Review Masyarakat Terhadap Tempat Wisata Pantai Di Kabupaten Karawang Pada Ulasan Google Maps," *J. Interkom J. Publ. Ilm. Bid. Teknol. Inf. dan Komun.*, vol. 16, 2022.
- [2] M. A. Anwar, H. Mulyo, and T. Tamrin, "Optimalisasi Algoritma Naive Bayes Dengan Teknik Ensemble Dalam Analisis Sentimen Twitter Pantai Kartini Jepara," *J. Minfo Polgan*, vol. 13, pp. 1331–1341, 2024.
- [3] H. Azizah, F. Syuhada, and Y. Sa'adati, "Sentimen Analisis Tempat Wisata Berdasarkan Ulasan Google Maps Menggunakan Metode Naive Bayes (Studi Kasus Bukit Merese)," *SainsTech Innov. J.*, vol. 7, no. November, pp. 467–476, 2024.
- [4] Y. W. Syaifudin and R. A. Irawan, "Implementasi Analisis Clustering dan Sentimen Data Twitter Pada Opini Wisata Pantai Menggunakan Metode K-Means," *J. Inform. Polinema*, vol. 4, pp. 189–194, 2018.
- [5] A. N. Izza, D. E. Ratnawati, W. Hayuhardhika, and W. H. N. Putra, "Analisis Sentimen Objek Wisata di Provinsi Sulawesi Selatan Berdasarkan Ulasan Pengunjung Menggunakan

- Metode Random Forest Classifier,” *J. Sist. Informasi, Teknol. Informasi, dan Edukasi Sist. Inf.*, vol. 3, no. 2, pp. 97–105, 2022.
- [6] B. R. Atmadja, “Analisis Sentimen Bahasa Indonesia Pada Tempat Wisata di Kabupaten Sukabumi Dengan Naïve Bayes,” *J. Ilm. Elektron. dan Komput.*, vol. 15, no. 2, pp. 371–382, 2022.
- [7] D. A. Fatah, E. M. S. Rochman, W. Setiawan, A. R. Aulia, F. I. Kamil, and A. Su’ud, “Sentiment Analysis of Public Opinion Towards Tourism in Bangkalan Regency Using Naïve Bayes Method,” in *E3S Web of Conferences*, 2024, pp. 1–8.
- [8] J. Ipmawati, Saifulloh, and Kusnawi, “Sentiment Analysis of Tourist Attractions Based on Reviews on Google Maps Using the Support Vector Machine Algorithm,” *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. April, pp. 247–256, 2024.
- [9] R. Ramdani and R. Cahyana, “Analisis Sentimen Ulasan Wisata Budaya Menggunakan Metode Support Vector Machine dan Long Short-Term Memory,” *J. Algoritma*, pp. 1188–1199, 2025, doi: 10.33364/algoritma/v.22-2.2585.
- [10] M. Z. A. Yusuf, S. Lestanti, and W. D. Puspitasari, “Analisis Sentimen Pengunjung Terhadap Fasilitas Dan Layanan Di Pantai Pudak Menggunakan Algoritma Lstm Muhammad,” *J. Ilm. Wahana Pendidik.*, vol. 12, pp. 128–141, 2026.
- [11] N. Oktaviana, H. C. Rustamaji, and H. Sofyan, “Sentiment Analysis On Reviews Of Beach Tourism Objects On Google Maps Using Long-Short Term Memory Method,” in *Seminar Nasional Informatika 2021 (SEMNASIF 2021)*, 2021, pp. 133–143.
- [12] W. A. Arifin, I. Ariawan, A. A. Rosalia, A. S. Sasongko, and M. R. Apriansyah, “Model Prediksi Pasang Surut Air Laut Pada Stasiun Pushidrosal Bakauheni Lampung Menggunakan Support Vector Regression,” *J. Kemaritiman Indones. J. Marit.*, vol. 2, no. 2, pp. 105–112, 2023.
- [13] N. S. Fitriyanti, D. R. Azhari, M. G. Shafa, and A. R. T. E. Ahmad, “Model Sistem Manajemen Pengetahuan di Program Studi Sistem Informasi Kelautan,” *J. Kemaritiman Indones. J. Marit.*, vol. 1, no. 2, pp. 89–100, 2023.
- [14] T. Chen and C. Guestrin, “XGBoost: A Scalable Tree Boosting System,” in *ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016.
- [15] J. Zalukhu and A. H. Lubis, “Klasifikasi Wisatawan Berdasarkan Perilaku Konsumen Dengan Menggunakan Algoritma XGBoost,” *JITE (Journal Informatics Telecommun. Eng.*, vol. 4, no. January, 2021.
- [16] H. Christanto *et al.*, “Analisis Perbandingan Decision Tree , Support Vector Machine , dan Xgboost dalam Mengklasifikasi Review Hotel Trip Advisor,” *J. Teknol. Inform. dan Komput. MH. Thamrin*, vol. 9, no. 1, pp. 306–319, 2023.
- [17] I. Ariawan, A. A. Rosalia, L. Anzani, L. O. A. Minsaris, and Lukman, “Identifikasi Spesies Mangrove Menggunakan Algoritme Random Forest,” *J. Kemaritiman Indones. J. Marit.*, vol. 2, no. 2, pp. 87–96, 2021.
- [18] T. K. Aulia, W. A. Arifin, and M. Rudi, “Klasifikasi Kualitas Air Budidaya Ikan Nila Menggunakan Support Vector Machine,” *J. Algoritma*, pp. 1842–1852, 2025, doi: 10.33364/algoritma/v.22-2.2356.
- [19] Tukino and A. R. Hakim, “Analisis Sentimen Objek Wisata di Google Maps Menggunakan Metode Decision Tree,” *Comput. Based Inf. Syst. J.*, vol. 01, pp. 122–130, 2024.
- [20] M. Guia, R. R. Silva, and J. Bernardio, “Comparison of Naïve Bayes , Support Vector Machine , Decision Trees and Random Forest on Sentiment Analysis,” in *International Joint Conference on Knowledge Discovery, Knowledge Engineering and Knowledge Management (IC3K2019)*, 2019, pp. 525–531. doi: 10.5220/0008364105250531.
- [21] M. J. Setiawan and V. R. S. Nastiti, “DANA App Sentiment Analysis: Comparison of XGBoost , SVM , and Extra Trees,” *J. SISFOKOM (Sistem Inf. dan Komputer)*, vol. 13, pp. 337–345, 2024.
- [22] L. Handyanto, W. A. Arifin, H. Syafri, U. P. Indonesia, and J. Barat, “Perbandingan Metode Hyperparameter Tuning Pada Model XGBoost Untuk Prediksi Multioutput,” *J. Rekayasa Teknol. Inf.*, vol. 10, no. 1, pp. 50–59, 2026.
- [23] R. T. Handayanto, Herlawati, P. D. Atika, F. N. Khasanah, A. Y. P. Yusuf, and D. Y. Septia, “Analisis Sentimen Pada Situs Google Review dengan Naïve Bayes dan Support Vector Machine,” *J. Komtika (Komputasi dan Inform.)*, vol. 5, no. 2, pp. 153–163, 2021.
- [24] I. Ariawan, Y. Herdiyeni, and I. Z. Siregar, “Geometry Feature Extraction of Shorea Leaf Venation Based on Digital Image and Classification Using Random Forest,” *International J. Comput. Digit. Syst.*, vol. 1, no. 1, 2022.
- [25] G. Rafianto, Hannie, and A. Mas’um, “Analisis Sentimen Berbasis Topik Ulasan Pengguna Aplikasi Roblox Menggunakan Integrasi LDA dan SVM,” *JITET J. Inform. dan Tek. Elektro Terap.*, vol. 14, no. 2, 2026.