

Perbandingan Algoritma *Machine Learning* dalam Analisis Sentimen Trans Jogja di Media Sosial

Comparison of Machine Learning Algorithms for Sentiment Analysis of Trans Jogja on Social Media

¹Putri Muryanti Setyowati*, ²Yoga Pristyanto, ³Arif Nur Rohman

^{1,2,3}Program Studi Sistem Informasi, Fakultas Ilmu Komputer, Universitas Amikom Yogyakarta, Yogyakarta, Indonesia

^{1,2,3}Jl. Ring Road Utara, Condong Catur, Depok, Sleman, Daerah Istimewa Yogyakarta, Indonesia

*e-mail: putrimuryanti@students.amikom.ac.id, yoga.pristyanto@amikom.ac.id,
arifrahman@amikom.ac.id

(received: 29 May 2026, revised: 23 July 2026, accepted: 24 July 2026)

Abstrak

Platform media sosial X dimanfaatkan masyarakat untuk menyampaikan pandangan dan pengalaman terkait layanan publik, termasuk Trans Jogja. Data dari platform tersebut dapat dianalisis untuk mengetahui persepsi pengguna terhadap kualitas layanan. Karena data berbentuk teks tidak terstruktur, penelitian ini menerapkan pendekatan klasifikasi sentimen berbasis machine learning. Penelitian bertujuan membandingkan performa sejumlah algoritma klasifikasi sentimen, yaitu algoritma Naive Bayes, Support Vector Machine (SVM), Random Forest, Neural Network, Logistic Regression, dan Decision Tree dalam mengklasifikasikan sentimen pengguna Trans Jogja di platform X. Pengumpulan data melalui proses crawling dengan bantuan Tweet Harvest menggunakan keyword yang berkaitan dengan Trans Jogja pada periode 1 Januari 2025 hingga 10 Juni 2026 sebanyak 3.035 tweet. Tahap preprocessing meliputi pembersihan data, case folding, tokenization, normalization, penghapusan kata umum (stopword), dan stemming. Representasi teks dilakukan menggunakan metode TF-IDF, kemudian dataset dibagi menjadi data training dan data testing dengan variasi rasio 90:10, 85:15, 80:20, 75:25, dan 70:30. Kinerja model dievaluasi melalui confusion matrix beserta metrik akurasi, presisi, recall, dan F1-score. Berdasarkan hasil pengujian, algoritma SVM menunjukkan performa yang paling optimal dan stabil dibandingkan metode lainnya. Pada rasio 85:15, algoritma memperoleh akurasi 91%, presisi 91%, recall 91%, dan f1-score 91%, sehingga dinilai paling efektif untuk klasifikasi sentimen pengguna layanan Trans Jogja.

Kata kunci: machine learning, analisis sentimen, trans jogja, media sosial.

Abstract

The social media platform X has become an important channel for the public to express opinions and share experiences regarding public services, including Trans Jogja. User-generated content from this platform provides valuable insights into public perceptions of service quality. However, because these data consist of unstructured text, sentiment classification techniques based on machine learning are required to analyze them effectively. This study aims to compare the performance of several machine learning algorithms for sentiment classification, including Naïve Bayes, Support Vector Machine (SVM), Random Forest, Neural Network, Logistic Regression, and Decision Tree, in classifying user sentiment toward Trans Jogja on the X platform. Data were collected through a web crawling process using Tweet Harvest with keywords related to Trans Jogja, covering the period from January 1, 2025, to June 10, 2026, resulting in a dataset of 3,035 tweets. The preprocessing stage included data cleaning, case folding, tokenization, normalization, stopwords removal, and stemming. Text representation was performed using the Term Frequency-Inverse Document Frequency (TF-IDF) method. The dataset was then divided into training and testing sets using five train-test split ratios: 90:10, 85:15, 80:20, 75:25, and 70:30. Model performance was evaluated using a confusion matrix and the corresponding accuracy, precision, recall, and F1-score metrics. The experimental results demonstrate that the Support Vector Machine (SVM) consistently outperformed the other algorithms

<http://sistemasi.ftik.unisi.ac.id>

across different data split ratios. At the 85:15 train–test split, the SVM achieved an accuracy of 91%, precision of 91%, recall of 91%, and an F1-score of 91%, indicating that it is the most effective algorithm for sentiment classification of Trans Jogja users on the X platform.

Keywords: *machine learning, sentiment analysis, trans jogja, social media.*

1 Pendahuluan

Dalam konteks mobilitas masyarakat perkotaan, khususnya di area padat penduduk seperti Yogyakarta, transportasi umum memegang peranan krusial. Layanan Trans Jogja dirancang dengan harapan dapat memitigasi kemacetan, mengoptimalkan efisiensi perjalanan, serta menyediakan opsi transportasi yang terjamin keamanan dan kenyamanannya bagi publik. Sejak beroperasi pada tahun 2008, Trans Jogja telah menjadi moda transportasi publik pilihan bagi berbagai segmen masyarakat, mulai dari pelajar, pekerja, hingga wisatawan, untuk menunjang aktivitas harian mereka. Lebih lanjut, pengembangan layanan Trans Jogja juga berorientasi pada peningkatan mutu pelayanan publik melalui implementasi sistem transportasi yang modern dan mudah diakses [1].

Kemajuan teknologi digital telah mendorong masyarakat untuk lebih proaktif dalam menyuarakan pandangan serta pengalaman mereka terkait layanan publik melalui platform media sosial, terutama Twitter atau X. Platform tersebut merupakan sumber data yang berharga, mengingat kemampuannya memfasilitasi pengguna untuk menyampaikan masukan, kritik, atau apresiasi secara langsung dan dalam format teks ringkas. Intensitas penggunaan media sosial yang tinggi memungkinkan pemanfaatan data opini publik untuk memperoleh pemahaman yang lebih cepat dan objektif mengenai persepsi masyarakat terhadap kualitas layanan transportasi [2].

Analisis sentimen adalah suatu pendekatan yang berfungsi untuk mengkategorikan opini publik berdasarkan arah polaritasnya, yakni sentimen positif, negatif, atau netral [3]. Pendekatan ini sering diaplikasikan untuk menunjang evaluasi layanan publik, dengan memanfaatkan data tekstual yang bersumber dari platform media sosial. Khususnya dalam sektor transportasi umum, analisis sentimen dapat digunakan untuk mengevaluasi tingkat kepuasan pengguna terkait fasilitas, jadwal, kenyamanan, dan aspek operasional lainnya. Studi-studi terdahulu mengindikasikan bahwa pandangan masyarakat mengenai transportasi umum di media sosial mampu merefleksikan kualitas layanan yang sesungguhnya dirasakan oleh pengguna [2], [4]. Selain itu, analisis sentimen juga dinilai efektif dalam menunjang pengambilan keputusan yang didasarkan pada data, baik untuk sektor transportasi publik maupun transportasi online [3].

Namun demikian, penelitian terkait analisis sentimen berbasis media sosial yang berfokus pada layanan Trans Jogja masih tergolong terbatas. Penelitian sejenis pada moda transportasi lain memang sudah cukup banyak dilakukan, namun masing-masing umumnya hanya menerapkan satu hingga dua algoritma klasifikasi tanpa perbandingan yang komprehensif dengan pendekatan lain, misalnya perbandingan Naive Bayes dan SVM pada TransJakarta [2], penerapan tunggal Random Forest pada Trans Semarang [4], atau perbandingan Naive Bayes dan Decision Tree pada JakLingko [5]. Yang lebih penting, algoritma yang teruji unggul di satu layanan transportasi ternyata tidak selalu menjadi yang terbaik di layanan lainnya. Telah banyak dilakukan penelitian serupa terkait moda transportasi lain, namun umumnya masing-masing hanya menggunakan satu hingga dua algoritma klasifikasi tanpa perbandingan yang menyeluruh dengan pendekatan lainnya, misalnya perbandingan Naive Bayes dan SVM di TransJakarta [2], penerapan tunggal Random Forest di Trans Semarang [4], atau perbandingan Naive Bayes dan Decision Tree pada JakLingko [5]. Algoritma yang teruji unggul di satu layanan transportasi ternyata tidak selalu menjadi yang terbaik di layanan lainnya, seperti SVM menunjukkan hasil yang lebih baik pada data TransJakarta [2], Random Forest di Trans Semarang [4], sedangkan Naive Bayes sedikit lebih unggul dibandingkan Decision Tree di JakLingko [5], dan Logistic Regression lebih unggul daripada MultinomialNB, SVM, serta KNN pada ulasan Gojek [6]. Variasi dalam hasil ini menunjukkan bahwa kinerja algoritma sentimen sangat dipengaruhi oleh karakteristik data.

Berdasarkan celah tersebut, penelitian ini menjadi krusial karena dua alasan utama. Pertama, dari perspektif objek penelitian, belum ada penelitian yang secara khusus memanfaatkan data opini pengguna Trans Jogja dari platform X, sementara karakteristik layanan transportasi publik di Yogyakarta memiliki potensi untuk menghasilkan pola opini yang berbeda dibandingkan dengan TransJakarta, Trans Semarang, JakLingko, maupun Gojek. Kedua, dari aspek metodologi, mengingat

<http://sistemasi.ftik.unisi.ac.id>

bahwa performa algoritma terbukti tidak konsisten di berbagai domain transportasi seperti yang telah dijelaskan, penting untuk menguji algoritma secara langsung pada dataset Trans Jogja itu sendiri dan bukan dengan mengandalkan hasil dari penelitian mengenai layanan lain. Penelitian ini mengisi kekosongan tersebut dengan membandingkan enam algoritma sekaligus, yaitu Naive Bayes, Support Vector Machine, Random Forest, Neural Network, Logistic Regression, dan Decision Tree.

Oleh karena itu, penelitian ini bertujuan menganalisis kinerja beberapa algoritma dalam mengkategorikan sentimen pengguna Trans Jogja di platform media sosial X, sekaligus mengidentifikasi model yang paling optimal. Hasil penelitian diharapkan dapat menjadi acuan bagi pengelola Trans Jogja dalam meningkatkan kualitas pelayanan transportasi publik sesuai dengan kebutuhan dan persepsi masyarakat.

2 Tinjauan Literatur

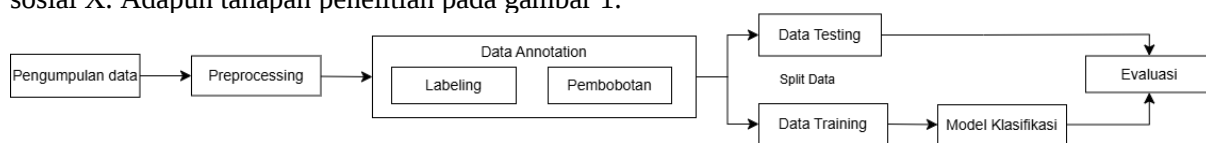
Beberapa penelitian telah menganalisis sentimen pengguna transportasi publik menggunakan algoritma machine learning. Ramdhania, Hidayat, dan Salkiawati [2] membandingkan Naive Bayes dan SVM pada data Twitter TransJakarta, di mana SVM memperoleh akurasi tertinggi sebesar 84,95% sementara Naive Bayes hanya mencapai 76,43%. Dalam penelitian mengenai layanan Trans Semarang, Fahreza dan Handoko [4] menerapkan algoritma Random Forest dan berhasil mencapai akurasi 81%, yang sekaligus menunjukkan bahwa algoritma tersebut mampu meningkatkan stabilitas dalam klasifikasi dan mengurangi risiko overfitting.

Penelitian pada layanan lain menunjukkan hasil yang beragam. Purnamasari dan Agoestanto [5] menganalisis sentimen terkait JakLingko melalui platform media sosial X dengan menggunakan Naive Bayes dan Decision Tree, yang menghasilkan akurasi masing-masing sebesar 84,9% dan 84,2%. Di sisi lain, dalam sektor transportasi berbasis aplikasi, Maulana et al [6] membandingkan Logistic Regression, MultinomialNB, SVM, dan KNN berdasarkan ulasan Gojek di Google Play Store, di mana Logistic Regression menunjukkan keunggulan dengan akurasi 82,45%. Ramadhani [7] juga menunjukkan bahwa Neural Network mampu menghasilkan klasifikasi yang baik dalam mendeteksi sentimen pengguna platform Gojek.

Berdasarkan penelitian terdahulu, terlihat bahwa algoritma yang menunjukkan hasil terbaik bervariasi di setiap layanan yaitu SVM untuk TransJakarta [2], Random Forest untuk Trans Semarang [4], Naive Bayes untuk JakLingko [5], dan Logistic Regression untuk Gojek [6]. Ketidakkonsistenan ini menekankan bahwa performa algoritma sentimen tergantung pada karakteristik dataset dan layanan tertentu, sehingga tidak bisa diasumsikan berlaku secara umum di berbagai moda transportasi. Selain itu, penelitian-penelitian tersebut hanya menguji maksimal empat algoritma. Penelitian yang secara spesifik mempelajari sentimen masyarakat terhadap layanan Trans Jogja menggunakan data dari media sosial X juga masih sangat terbatas. Oleh karena itu, penelitian ini membandingkan enam algoritma, yaitu Naive Bayes, Support Vector Machine, Random Forest, Neural Network, Logistic Regression, dan Decision Tree dengan menggunakan representasi TF-IDF serta beberapa skenario pembagian data.

3 Metode Penelitian

Penelitian ini memanfaatkan metode kuantitatif dengan penerapan analisis sentimen berbasis algoritma machine learning untuk mengklasifikasikan pendapat pengguna Trans Jogja pada media sosial X. Adapun tahapan penelitian pada gambar 1.



Gambar 1 Alur penelitian

3.1 Pengumpulan Data

Data dikumpulkan melalui teknik web scraping menggunakan library Python Tweet Harvest [5] dengan keyword ‘transjogja’, ‘#transjogja’, ‘#TransJogja’, ‘Trans Jogja’, ‘bus trans jogja’, dan ‘halte trans jogja’. Pengambilan data dilakukan dalam periode 1 Januari 2025 hingga 10 Juni 2026. Seluruh data hasil scraping kemudian disimpan dalam format csv.

3.2 Preprocessing

Dalam penelitian ini, preprocessing digunakan sebagai tahap awal untuk memperbaiki kualitas data teks sebelum memasuki tahap analisis berikutnya. Tahapan ini dilakukan guna meminimalkan noise pada data, menyeragamkan bentuk penulisan, serta menghasilkan data yang lebih relevan untuk diolah model klasifikasi [8]. Dalam penelitian ini, tahap preprocessing mencakup data cleaning, case folding, tokenization, normalization, stopword removal, dan stemming.

3.3 Labeling

Proses penentuan label data dilakukan secara otomatis dengan metode lexicon, yaitu kamus kata yang telah dikategorikan ke dalam sentimen positif, negatif, dan netral [9], [10]. Setiap kata pada teks dicocokkan dengan kamus untuk memperoleh nilai sentimen, kemudian dijumlahkan guna menentukan kecenderungan dokumen. Skor akhir yang bernilai tinggi (melebihi batas atas) menunjukkan dominasi kata positif sehingga dianggap sebagai sentimen positif, sementara skor yang rendah (di bawah batas bawah) mencerminkan dominasi kata negatif dan dikategorikan sebagai sentimen negatif, sedangkan skor di antara kedua batas tersebut dianggap sebagai sentimen netral. Ambang batas ditetapkan secara adaptif berdasarkan distribusi data untuk memastikan proporsi antar kelas lebih terjaga dan seimbang. Pendekatan ini dinilai efektif dan banyak digunakan untuk labeling data teks media sosial [11]. Untuk memastikan akurasi dan validitas hasil pelabelan, proses ini selanjutnya divalidasi oleh seorang ahli bahasa Indonesia, mengacu pada pendekatan yang dilakukan oleh Fitriyani, Jajuli, dan Umaidah [9]. Annotator meninjau label yang sudah ada dan memperbaiki yang dinilai tidak sesuai konteks pada 100 tweet yang dipilih secara acak, dengan fokus pada kasus ambigu, sarkasme, dan makna implisit yang sulit tertangkap kamus lexicon.

3.4 Pembobotan

Data hasil preprocessing yang berbentuk kalimat dikonversi menjadi representasi numerik. Konversi ini dilakukan melalui proses pembobotan kata guna menentukan nilai bobot yang berfungsi sebagai fitur. Penelitian ini menerapkan metode TF-IDF (Term Frequency-Inverse Document Frequency) untuk merepresentasikan teks dengan memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam dokumen. Kata yang jarang muncul akan memperoleh nilai lebih besar, sehingga metode ini efektif dalam merepresentasikan teks untuk analisis sentimen [4].

3.5 Split Data

Data yang telah dikonversi ke dalam format numerik selanjutnya dipisahkan menjadi dua segmen utama, yaitu data training dan data testing. Data training terdiri dari data historis yang merepresentasikan kondisi aktual, sedangkan data testing untuk mengevaluasi efektivitas pengklasifikasi dalam melaksanakan tugas klasifikasi secara akurat [12]. Pembagian data dilakukan dengan beberapa skenario perbandingan, yaitu 90:10, 85:15, 80:20, 75:25, dan 70:30 untuk mengetahui pengaruh proporsi data terhadap performa model.

3.6 Model Klasifikasi

Selanjutnya, model yang telah dibangun akan diuji dengan menggunakan data testing untuk menilai seberapa efektif model dalam mengenali sentimen dengan tepat.

3.6.1 Naïve Bayes

Metode Naïve Bayes adalah teknik klasifikasi yang berlandaskan pada teorema Bayes. Pendekatan ini mengkombinasikan prior probability dan conditional probability dalam membangun model probabilistik [13], dengan asumsi bahwa setiap fitur memiliki sifat independen antar fitur lainnya. Dalam konteks klasifikasi teks, fitur TF-IDF digunakan untuk meningkatkan relevansi suatu kata terhadap masing-masing kelas. Hasil prediksi kemudian ditentukan berdasarkan kelas yang memiliki nilai probabilitas tertinggi.

3.6.2 Support Vector Machine (SVM)

Support Vector Machine merupakan metode klasifikasi yang berfungsi dengan menemukan hyperplane terbaik untuk memisahkan berbagai kelas data dengan seefektif mungkin. Tujuannya untuk menemukan batas pemisah yang dapat mengklasifikasikan data secara maksimal dalam ruang

<http://sistemasi.ftik.unisi.ac.id>

berdimensi tinggi [14] sehingga efektif digunakan pada analisis sentimen teks [2]. Konfigurasi yang digunakan yaitu C (cost) dengan nilai sebesar 1, r (coefficient) dengan nilai sebesar 1, dan d (degree) sebesar 3 dengan kernel polynomial yang berfungsi untuk mencari hyperplane yang efektif untuk membedakan data dalam ruang yang sudah dimodifikasi. Parameter C berfungsi mengontrol keseimbangan antara peningkatan margin dan pengurangan kesalahan klasifikasi, parameter r berfungsi sebagai koefisien independen, dan parameter d menentukan derajat kernel polynomial yang digunakan saat memproses data ke dalam dimensi yang lebih tinggi.

3.6.3 Random Forest

Random Forest merupakan metode yang menciptakan sejumlah pohon keputusan untuk melakukan proses klasifikasi maupun regresi. Metode ini dirancang dengan memilih fitur secara acak dan memanfaatkan pendekatan CART (Classification and Regression Tree) [15] dengan kriteria pemisahan node berdasarkan gini impurity untuk memaksimalkan homogenitas data pada setiap cabang. Pada penelitian ini, model Random Forest dikonfigurasi dengan jumlah pohon sebanyak 300, *random_state* sebesar 42 agar hasil pelatihan tetap konsisten, dan *n_jobs* sebesar -1 yang memungkinkan proses komputasi berjalan secara bersamaan menggunakan seluruh jumlah core prosesor yang tersedia, sehingga mempercepat pelatihan model.

3.6.4 Neural Network

Neural Network didefinisikan sebagai suatu algoritma komputasi yang mengadopsi prinsip kerja otak biologis manusia dalam memproses informasi [16]. Dalam konteks penelitian ini, metode Multilayer Perceptron (MLP) diimplementasikan, yang strukturnya mencakup satu lapisan masukan, sejumlah lapisan tersembunyi, dan satu lapisan keluaran. Secara spesifik, lapisan masukan berfungsi untuk mengakomodasi penerimaan vektor fitur yang merepresentasikan data tekstual. Setelah itu, setiap neuron dalam lapisan tersembunyi akan melaksanakan perhitungan berbobot yang diiringi fungsi aktivasi, sebelum hasilnya diteruskan ke lapisan selanjutnya. Melalui serangkaian proses tersebut, prediksi sentimen akan terbentuk pada lapisan output [17]. Arsitektur MLP dikonfigurasi dengan *hidden_layer_size* sebesar 100 dan 50 neuron menggunakan fungsi *activation* Rectified Linear Unit (ReLU) di setiap lapisan tersembunyi. Proses optimasi bobot jaringan dilakukan dengan menggunakan algoritma Adam. Selain itu, *max_iter* ditetapkan sebesar 300 untuk memastikan bahwa pelatihan mencapai konvergensi yang paling optimal, sedangkan *random state* ditetapkan sebesar 42.

3.6.5 Logistic Regression

Logistic Regression adalah teknik klasifikasi yang digunakan untuk menganalisis keterkaitan antara variabel bebas, baik tunggal maupun jamak, dengan variabel terikat yang bersifat kategoris. Pada penelitian ini, pendekatan Multinomial Logistic Regression digunakan untuk mengelompokkan sentimen ke dalam tiga kelas, yaitu positif, negatif, dan netral [18].

3.6.6 Decision Tree

Decision Tree merupakan algoritma klasifikasi yang membagi data secara bertahap berdasarkan nilai fitur tertentu untuk memaksimalkan pemisahan antar kelas hingga mencapai kriteria penghentian, seperti kedalaman maksimum pohon atau kondisi homogen pada simpul daun [14]. Setelah model terbentuk, algoritma digunakan untuk memprediksi data testing, kemudian hasilnya dievaluasi dengan metrik akurasi, presisi, recall, dan f1-score untuk menilai kemampuan klasifikasi sentimen.

3.7 Evaluasi

Evaluasi terhadap model yang telah menjalani proses pelatihan dieksekusi melalui pemanfaatan confusion matrix. Metode ini merepresentasikan instrumen evaluasi klasifikasi yang disajikan dalam format tabel matriks, bertujuan untuk membandingkan luaran prediksi dengan data observasi aktual [19]. Representasi visual metrik evaluasi, yang meliputi akurasi, presisi, recall, serta f1-score, dihasilkan melalui pemanfaatan pustaka *seaborn* dan *scikit-learn*. Akurasi menunjukkan tingkat ketepatan model dalam menghasilkan prediksi yang tepat dibandingkan dengan total data testing yang dipakai, sementara itu presisi mengukur tingkat ketepatan model dalam mengidentifikasi suatu kategori spesifik. Di sisi lain, recall mengevaluasi kapabilitas model dalam mendeteksi seluruh data

<http://sistemasi.ftik.unisi.ac.id>

yang relevan. Sebagai metrik penyeimbang antara presisi dan recall, f1-score dimanfaatkan untuk menyajikan gambaran holistik mengenai kualitas performa model secara menyeluruh.

4 Hasil dan Pembahasan

Seluruh hasil penelitian disajikan dalam bentuk analisis untuk mengetahui kinerja model pada klasifikasi sentimen pengguna Trans Jogja menggunakan data dari platform X. Selanjutnya, interpretasi hasil dilakukan melalui pembahasan berdasarkan metrik evaluasi yang telah ditentukan.

4.1 Pengumpulan Data

Data dalam penelitian ini dikumpulkan dari platform media sosial X dengan memanfaatkan tools Tweet Harvest. Proses pengumpulan dilakukan menggunakan kata kunci seperti “transjogja”, “#transjogja”, “#TransJogja”, “Trans Jogja”, “bus trans jogja”, dan “halte trans jogja” agar data yang terkumpul relevan dengan fokus penelitian. Rentang waktu pengumpulan berlangsung dari 1 Januari 2025 hingga 10 Juni 2026. Melalui tahapan tersebut, berhasil dikumpulkan 3.035 tweet yang kemudian dijadikan dataset penelitian. Data yang diperoleh berupa teks berisi pendapat, pengalaman, serta tanggapan pengguna terhadap layanan Trans Jogja. Data mentah hasil crawling ditunjukkan oleh tabel 1 berikut:

Tabel 1 Data mentah

Created_at	Full_text
Tue Apr 14 00:16:35 +0000 2026	@adhie_tholha @Outstandjing Kemarin ke PWT liat bagus-bagus rapi linenya juga lbh bagus dr. Trans Jogja wqqq BRT Barlingmascakeb

Tabel di atas menunjukkan data mentah hasil crawling media sosial X, terdiri dari kolom created_at yang berisi waktu unggahan tweet dan kolom full_text yang berisi tweet pengguna terkait pelayanan Trans Jogja. Data tersebut masih mengandung elemen tidak terstruktur seperti mention, singkatan, dan kata tidak baku sehingga perlu dilakukan tahap preprocessing sebelum digunakan dalam analisis.

4.2 Preprocessing

Data preprocessing dilakukan dengan beberapa tahap antara lain:

4.2.1 Data Cleaning

Tahap cleaning dilakukan dengan menghilangkan elemen yang tidak diperlukan, seperti hashtag, nama pengguna, link, simbol, angka, maupun karakter khusus yang dapat mengganggu analisis [4]. Proses ini dilakukan untuk meminimalkan noise pada data teks sehingga informasi yang dihasilkan lebih merepresentasikan opini pengguna secara akurat. Hasil cleaning dapat dilihat pada tabel 2 berikut.

Tabel 2 Hasil cleaning

Full_text	Clean
@buchenuppak Jadi nanti kalau mau ke Prambanan bisa naik krl atau transjogja (tp ini lama dan memusingkan bgttt). Rekomen naik krl aja (bisa dri tugulempuyangan 8rb tarifnya) trs turun st brambanan. Udh deh bisa jalan kaki kalau kuat wkkw atau ojek online	Jadi nanti kalau mau ke Prambanan bisa naik krl atau transjogja tp ini lama dan memusingkan bgttt Rekomen naik krl aja bisa dri tugulempuyangan rb tarifnya trs turun st brambanan Udh deh bisa jalan kaki kalau kuat wkkw atau ojek online

Hasil dari tahap cleaning membuat data lebih terstruktur dan mempermudah proses analisis pada tahapan berikutnya.

4.2.2 Case Folding

Setiap huruf kapital dalam dokumen dikonversi menjadi huruf kecil [9] dengan memanfaatkan fungsi `text.lower()` yang terdapat dalam library Pandas. Tahapan ini dilakukan untuk menyeragamkan format penulisan agar perbedaan huruf kapital dan huruf kecil tidak mempengaruhi proses analisis, sehingga kata dengan bentuk berbeda diperlakukan sebagai entitas yang sama. Hasil case folding ditunjukkan pada tabel 3 berikut.

Tabel 3 Hasil case folding

Clean	Case_folding
Jadi nanti kalau mau ke Prambanan bisa naik krl atau transjogja tp ini lama dan memusingkan bgttt Rekomen naik krl aja bisa dri tugulempuyangan rb tarifnya trs turun st brambanan Udh deh bisa jalan kaki kalau kuat wkkw atau ojek online	jadi nanti kalau mau ke prambanan bisa naik krl atau transjogja tp ini lama dan memusingkan bgttt rekomen naik krl aja bisa dri tugulempuyangan rb tarifnya trs turun st brambanan udh deh bisa jalan kaki kalau kuat wkkw atau ojek online

Pada tabel di atas, kolom `case_folding` berhasil menampilkan transformasi teks ke dalam bentuk huruf kecil.

4.2.3 Tokenization

Setelah proses case folding dilakukan, data kemudian dipisahkan menjadi token dalam proses tokenization [9]. Dalam penelitian ini, fungsi `word_tokenize` pada library NLP digunakan dalam proses tokenization untuk memisahkan teks menjadi token-token tertentu. Hasil tokenization ditunjukkan pada tabel 4 berikut.

Tabel 4 Hasil tokenization

Case_folding	Token
jadi nanti kalau mau ke prambanan bisa naik krl atau transjogja tp ini lama dan memusingkan bgttt rekomen naik krl aja bisa dri tugulempuyangan rb tarifnya trs turun st brambanan udh deh bisa jalan kaki kalau kuat wkkw atau ojek online	['jadi', 'nanti', 'kalau', 'mau', 'ke', 'prambanan', 'bisa', 'naik', 'krl', 'atau', 'transjogja', 'tp', 'ini', 'lama', 'dan', 'memusingkan', 'bgttt', 'rekomen', 'naik', 'krl', 'aja', 'bisa', 'dri', 'tugulempuyangan', 'rb', 'tarifnya', 'trs', 'turun', 'st', 'brambanan', 'udh', 'deh', 'bisa', 'jalan', 'kaki', 'kalau', 'kuat', 'wkkw', 'atau', 'ojek', 'online']

Berdasarkan hasil yang didapat, setiap dokumen teks berhasil direpresentasikan dalam bentuk daftar token.

4.2.4 Normalization

Selanjutnya, data dari platform media sosial sering kali mengandung istilah informal, bahasa gaul, bahkan kata-kata yang memiliki makna kasar. Oleh karena itu, diperlukan proses normalization untuk menyamakan kata tidak baku dan istilah gaul yang memiliki arti serupa [9]. Dalam penelitian ini, proses normalization memanfaatkan library IndoNLP guna mengonversi kosakata non-standar menjadi format baku dalam bahasa Indonesia. Hasil normalization ditunjukkan pada tabel 5 berikut.

Tabel 5 Hasil normalization

Token	Normalization
['jadi', 'nanti', 'kalau', 'mau', 'ke', 'prambanan', 'bisa', 'naik', 'krl', 'atau', 'transjogja', 'tp', 'ini', 'lama', 'dan', 'memusingkan', 'bgttt', 'rekomen', 'naik', 'krl', 'aja', 'bisa', 'dri', 'tugulempuyangan', 'rb', 'tarifnya', 'trs', 'turun', 'st', 'brambanan', 'udh', 'deh', 'bisa', 'jalan', 'kaki', 'kalau', 'kuat', 'wkkw', 'atau', 'ojek', 'online']	['jadi', 'nanti', 'kalau', 'mau', 'ke', 'prambanan', 'bisa', 'naik', 'krl', 'atau', 'transjogja', 'tapi', 'ini', 'lama', 'dan', 'memusingkan', 'banget', 'rekomen', 'naik', 'krl', 'saja', 'bisa', 'dari', 'tugulempuyangan', 'ribu', 'tarifnya', 'terus', 'turun', 'st', 'brambanan', 'sudah', 'deh', 'bisa', 'jalan', 'kaki', 'kalau', 'kuat', 'wkkw', 'atau',

<http://sistemasi.ftik.unisi.ac.id>

'ojek', 'online']

Kolom normalization pada tabel di atas merupakan hasil perubahan kata tidak baku, singkatan, dan bahasa gaul menjadi kata baku.

4.2.5 Stopword

Kemudian, tahap stopwords removal dilakukan untuk menghilangkan kata-kata yang sering muncul tetapi dianggap tidak memberikan kontribusi makna yang signifikan dalam proses analisis [13] seperti “dan”, “yang”, “di”, dan kata lain yang serupa. Hasil stopwords ditunjukkan pada tabel 6 berikut.

Tabel 6 Hasil stopwords removal

Normalization	Stopword
['jadi', 'nanti', 'kalau', 'mau', 'ke', 'prambanan', 'bisa', 'naik', 'krl', 'atau', 'transjogja', 'tapi', 'ini', 'lama', 'dan', 'memusingkan', 'banget', 'rekomen', 'naik', 'krl', 'saja', 'bisa', 'dari', 'tugulempuyangan', 'ribu', 'tarifnya', 'terus', 'turun', 'st', 'brambanan', 'sudah', 'deh', 'bisa', 'jalan', 'kaki', 'kalau', 'kuat', 'wkkw', 'atau', 'ojek', 'online']	['prambanan', 'krl', 'transjogja', 'tapi', 'memusingkan', 'banget', 'rekomen', 'krl', 'tugulempuyangan', 'ribu', 'tarifnya', 'turun', 'st', 'brambanan', 'deh', 'jalan', 'kaki', 'kuat', 'wkkw', 'ojek', 'online']

Berdasarkan tabel di atas, kata-kata umum yang tidak memiliki kontribusi besar terhadap analisis sentimen telah dihilangkan sehingga teks lebih berisi kata-kata penting dan relevan.

4.2.6 Stemming

Tahapan terakhir dalam preprocessing adalah stemming, yaitu proses mengubah kata-kata yang memiliki afiks menjadi bentuk kata dasar melalui penghapusan prefiks, sufiks, dan infiks. Pada penelitian ini, stemming dilakukan menggunakan library Sastrawi melalui StemmerFactory terhadap setiap token hasil normalization. Hasil stemming ditunjukkan pada tabel 7 berikut.

Tabel 7 Hasil stemming

Stemming	Text_data
['prambanan', 'krl', 'transjogja', 'tapi', 'pusing', 'banget', 'rekomen', 'krl', 'tugulempuyangan', 'ribu', 'tarif', 'turun', 'st', 'brambanan', 'deh', 'jalan', 'kaki', 'kuat', 'wkkw', 'ojek', 'online']	prambanan krl transjogja tapi pusing banget rekomen krl tugulempuyangan ribu tarif turun st brambanan deh jalan kaki kuat wkkw ojek online

Pada tabel di atas menunjukkan hasil perubahan kata berimbuhan menjadi bentuk dasar. Setelah itu, hasil stemming digabung kembali menjadi teks lengkap pada kolom text_data.

4.3 Labeling

Metode lexicon merupakan pendekatan yang diterapkan dalam penelitian ini. Metode ini mengintegrasikan pemanfaatan kamus yang mengkategorikan kata-kata berdasarkan nilai positif dan negatif, serta dilengkapi dengan bobot masing-masing [20]. Penentuan label dilakukan dengan menerapkan ambang batas yang ditentukan secara adaptif berdasarkan sebaran data, sehingga proporsi data pada masing-masing kelas dapat lebih terkontrol dan seimbang. Dokumen dengan skor di bawah batas bawah dikategorikan sebagai sentimen positif, di atas batas atas sebagai sentimen negatif, sedangkan skor di antara kedua batas tersebut termasuk sentimen netral. Hasil labeling ditunjukkan pada tabel 8 berikut.

Tabel 8 Hasil labeling

Kelas	Hasil	Persentase
-------	-------	------------

<http://sistemasi.ftik.unisi.ac.id>

Netral	1133	37,3%
Negatif	976	32,2%
Positif	926	30,5%

Berdasarkan hasil tabel di atas, memperlihatkan bahwa mayoritas pengguna memberikan opini yang informatif atau netral mengenai layanan Trans Jogja dengan persentase yang diperoleh 37,3%. Hasil pelabelan ini divalidasi oleh annotator terhadap 100 tweet yang dipilih secara acak. Hasil validasi ditunjukkan pada tabel 9 berikut.

Tabel 9 Hasil validasi

Text_data	Lexicon	Annotator
ancen adoh adoh mbak transjogja reliable warlok akan akomodasi turis doang	Positif	Negatif
realisasi transjogja date cerita cinta cinta anak pondok kota jogja	Positif	Netral
komentar kait lambat tuju banget sender crew trans jogja enggak profesional sengaja lambat biar pas operasional habis enggak poolgarasi	Netral	Negatif
tuju jogja tuh provinsi banget jadi wisata slow living warga taat atur ramah harga kuliner jangkau wisata krl trans	Netral	Positif
transjogja malam malam telinga dipasangin tws puter murphy radio beh vibesnya lo	Netral	Positif

Berdasarkan tabel 9, terlihat bahwa hasil pelabelan lexicon tidak sepenuhnya sama dengan validasi dari annotator. Perbedaan ini muncul karena teks komentar dari platform media sosial sering kali mengandung makna implisit, sarkasme, atau sentimen campuran dalam satu kalimat sehingga sulit ditangkap oleh kamus kata.

4.4 Pembobotan

Proses ini diimplementasikan menggunakan TfidfVectorizer pada library scikit-learn. Hasil TF-IDF dapat dilihat pada tabel 10 berikut.

Tabel 10 Hasil TF-IDF

Dokumen	Kata	Nilai TF-IDF
0	bus	0,151
0	bus stop	0,167
0	daerah	0,141
0	daerah istimewa	0,180
0	ev	0,157

Hasil TF-IDF di atas menunjukkan tingkat kepentingan kata atau frasa dalam dokumen. Semakin tinggi angkanya, semakin besar kontribusi kata tersebut dalam membedakan dokumen dari kumpulan dokumen lain.

4.5 Split Data

Tahap berikutnya adalah pembagian data, proses ini dilaksanakan dengan memanfaatkan fungsi `train_test_split` yang terdapat dalam library scikit-learn. Hasil Split Data dapat dilihat pada tabel 11 berikut.

Tabel 11 Perbandingan split

Rasio Perbandingan	Training Data	Testing Data
90:10	2731	304

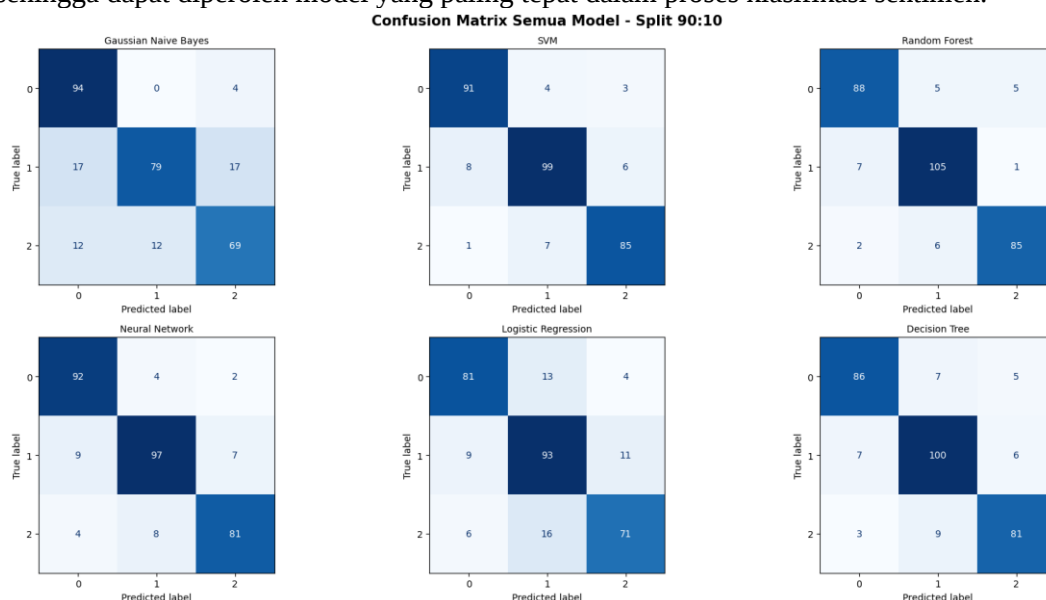
<http://sistemasi.ftik.unisi.ac.id>

85:15	2579	456
80:20	2428	607
75:25	2276	759
70:30	2124	911

Berdasarkan tabel di atas, semakin besar rasio data training yang digunakan maka jumlah data testing menjadi lebih sedikit. Perbedaan proporsi data di atas digunakan untuk melihat pengaruh jumlah data testing dan data training terhadap performa model klasifikasi.

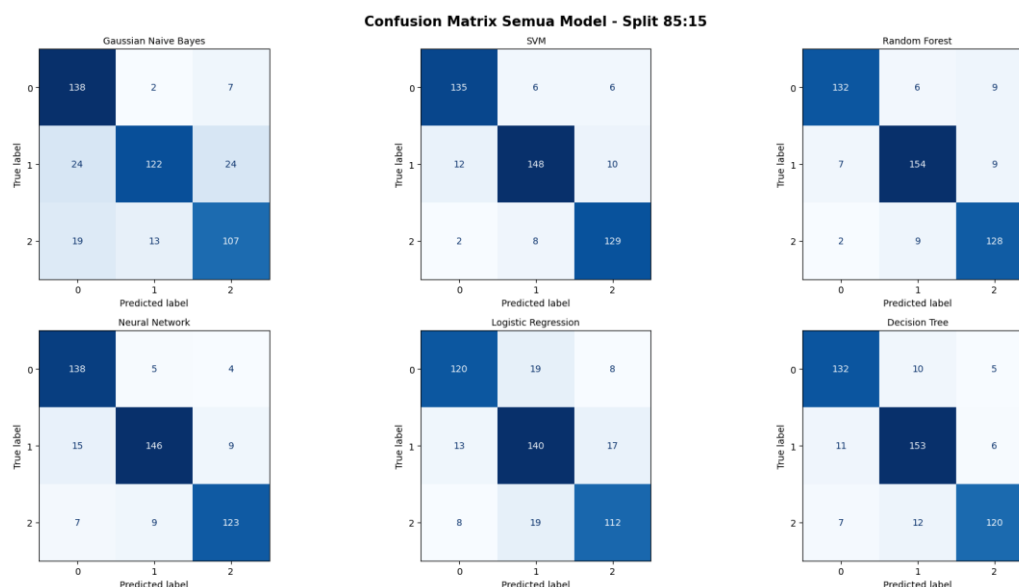
4.6 Model Klasifikasi dan Evaluasi

Klasifikasi dilakukan menggunakan algoritma Naive Bayes, SVM, Random Forest, Neural Network, Logistic Regression, dan Decision Tree. Pemilihan berbagai algoritma ini bertujuan untuk mengevaluasi dan membandingkan keefektifan setiap metode dalam menangani karakteristik data teks, sehingga dapat diperoleh model yang paling tepat dalam proses klasifikasi sentimen.



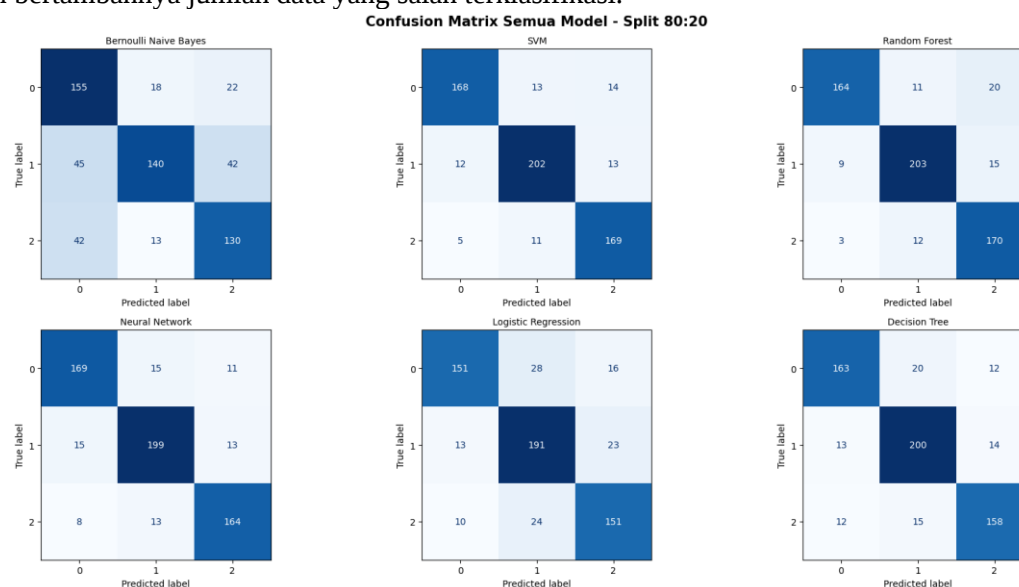
Gambar 2 Confusion matrix rasio 90:10

Berdasarkan confusion matrix, seluruh algoritma menunjukkan kinerja klasifikasi yang memuaskan, terlihat dari dominasi prediksi yang tepat di diagonal utama. Hal ini didukung oleh proporsi data training yang lebih banyak dibandingkan dengan data testing. Random Forest menunjukkan jumlah prediksi benar tertinggi dan kesalahan klasifikasi yang relatif paling sedikit, diikuti oleh SVM dan Neural Network yang juga menghasilkan distribusi prediksi yang sangat baik pada ketiga kelas. Decision Tree masih menunjukkan beberapa kesalahan antar kelas, sedangkan Logistic Regression dan Gaussian Naive Bayes menghasilkan jumlah kesalahan yang lebih besar, terutama pada kelas 1 dan kelas 2.



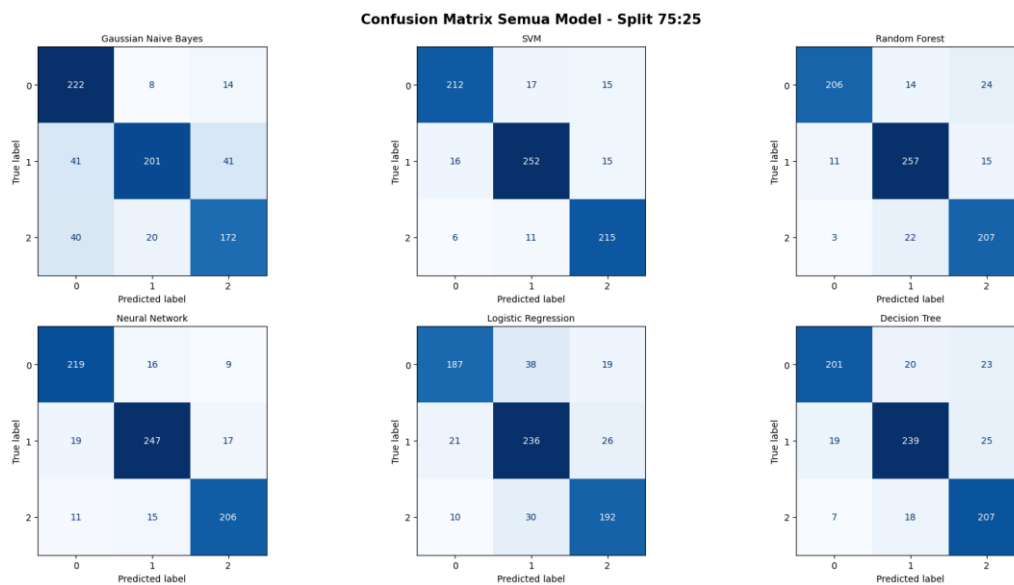
Gambar 3 Confusion matrix rasio 85:15

Berdasarkan gambar tersebut, meskipun terdapat peningkatan jumlah kesalahan klasifikasi dibandingkan dengan rasio 90:10, pola confusion matrix masih menunjukkan kinerja yang cukup baik di semua model. Random Forest dan SVM tetap menonjol sebagai model terbaik dengan dominasi prediksi yang benar pada diagonal utama. Decision Tree dan Neural Network masih menunjukkan stabilitas yang baik dengan dominasi prediksi benar pada diagonal utama. Meskipun mulai terlihat peningkatan kesalahan klasifikasi antar kelas, terutama antara kelas 0 dan kelas 1. Di sisi lain, Logistic Regression dan Naive Bayes menunjukkan penurunan performa yang lebih signifikan seiring dengan bertambahnya jumlah data yang salah terklasifikasi.



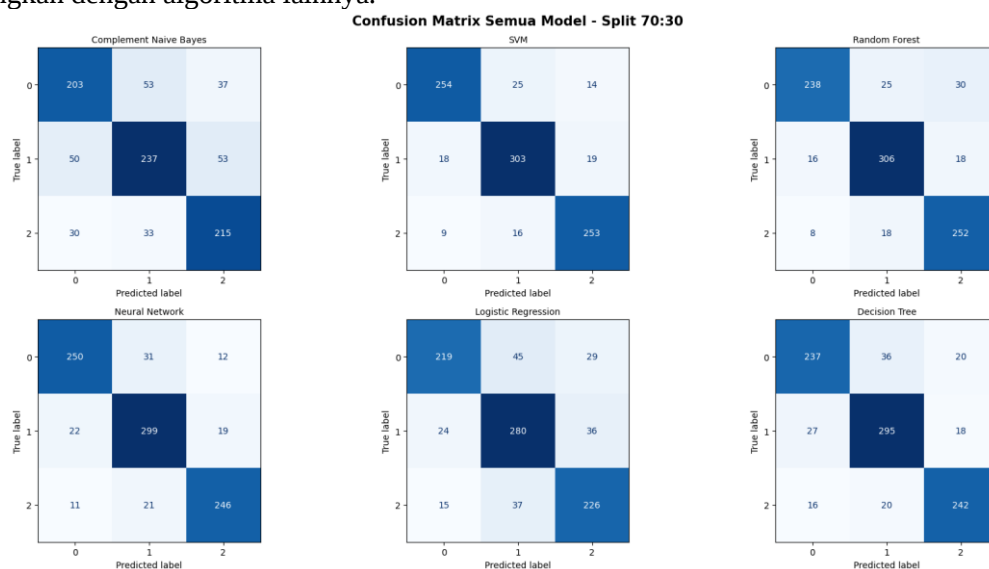
Gambar 4 Confusion matrix rasio 80:20

Berdasarkan confusion matrix, seluruh model masih menunjukkan kinerja klasifikasi yang cukup baik meskipun terdapat peningkatan kesalahan dibandingkan rasio sebelumnya. Random Forest dan SVM tetap menjadi algoritma yang paling unggul dengan dominasi prediksi yang benar pada diagonal utama. Decision Tree dan Neural Network relatif konsisten dengan dominasi prediksi benar pada ketiga kelas, meskipun mulai terlihat peningkatan kesalahan klasifikasi dibandingkan rasio sebelumnya. Sebaliknya, Logistic Regression dan terutama Naive Bayes mengalami penurunan performa yang lebih signifikan akibat meningkatnya jumlah data yang salah diklasifikasikan.



Gambar 5 Confusion matrix rasio 75:25

Hasil confusion matrix menunjukkan bahwa berkurangnya proporsi data training dibandingkan rasio sebelumnya menyebabkan peningkatan jumlah kesalahan klasifikasi yang ditunjukkan oleh bertambahnya nilai di luar diagonal utama. Meskipun mengalami peningkatan kesalahan klasifikasi, Random Forest dan SVM tetap menjadi model yang paling baik, khususnya pada kelas 2. Decision Tree dan Neural Network menunjukkan kestabilan yang relatif, namun kinerjanya menurun jika dibandingkan dengan rasio sebelumnya. Di sisi lain, Logistic Regression dan Naive Bayes mengalami peningkatan kesalahan prediksi yang signifikan sehingga memiliki kinerja yang lebih rendah dibandingkan dengan algoritma lainnya.



Gambar 6 Confusion matrix rasio 70:30

Pada hasil di atas, hampir semua algoritma menunjukkan peningkatan kesalahan klasifikasi seiring dengan berkurangnya data pelatihan dan bertambahnya data training. Namun, Random Forest dan SVM tetap menunjukkan hasil unggul dengan prediksi yang benar pada diagonal utama. Neural Network masih menunjukkan performa yang sangat baik dan relatif stabil dengan dominasi prediksi benar pada ketiga kelas, sedangkan Decision Tree mulai menunjukkan peningkatan kesalahan klasifikasi dibandingkan rasio sebelumnya. Sedangkan Logistic Regression dan Naive Bayes memperlihatkan kesalahan klasifikasi tertinggi pada rasio ini.

Secara keseluruhan, SVM, Random Forest, dan Neural Network menunjukkan performa yang paling baik dan konsisten pada seluruh rasio pembagian data. Di antara ketiganya, Random Forest

<http://sistemasi.ftik.unisi.ac.id>

menghasilkan prediksi yang paling stabil dan seimbang. Hal ini karena mekanisme ensemble dan pemilihan subset fitur secara acak memungkinkan Random Forest memanfaatkan fitur-fitur informatif pada data TF-IDF yang berdimensi tinggi dan bersifat sparse, sekaligus mengurangi pengaruh noise dan risiko overfitting. Temuan ini diperkuat oleh nilai akurasi, presisi, recall, dan f1-score yang relatif tinggi dan konsisten pada berbagai rasio pembagian data.

Tabel 12 Perbandingan model

Rasio Perbandingan	Algoritma	Akurasi	Presisi	Recall	F1-score
90:10	SVM	90%	90%	90%	90%
	Neural Network	90%	90%	90%	90%
	Random Forest	89%	89%	89%	89%
	Decision Tree	88%	88%	88%	88%
	Logistic Regression	81%	81%	81%	81%
	Naive Bayes (Gaussian)	79%	80%	80%	80%
85:15	SVM	91%	91%	91%	91%
	Random Forest	90%	90%	90%	90%
	Neural Network	90%	90%	90%	90%
	Decision Tree	89%	89%	89%	89%
	Logistic Regression	82%	82%	82%	82%
	Naive Bayes (Gaussian)	80%	81%	80%	80%
80:20	SVM	89%	89%	89%	89%
	Random Forest	89%	89%	88%	88%
	Neural Network	88%	88%	89%	88%
	Decision Tree	88%	88%	86%	86%
	Logistic Regression	86%	81%	81%	81%
	Naive Bayes (Bernoulli)	70%	72%	70%	70%
75:25	Neural Network	89%	89%	89%	89%
	SVM	88%	88%	88%	88%
	Random Forest	88%	88%	87%	87%
	Decision Tree	85%	85%	85%	85%
	Logistic regression	81%	81%	81%	81%
	Naive Bayes (Gaussian)	78%	79%	78%	78%
70:30	Neural Network	88%	88%	88%	88%
	Random Forest	88%	88%	88%	88%
	SVM	87%	87%	87%	87%
	Decision Tree	85%	85%	85%	85%
	Logistic Regression	80%	80%	80%	80%
	Naive Bayes (Complement)	72%	72%	72%	72%

Berdasarkan hasil evaluasi, algoritma dengan proporsi prediksi benar yang dominan pada diagonal utama confusion matrix cenderung menghasilkan metrik akurasi, presisi, recall, dan f1-score yang lebih tinggi serta seimbang antar kelas. SVM menunjukkan performa yang sangat baik dan konsisten pada seluruh skenario pembagian data. SVM memperoleh akurasi tertinggi sebesar 91% pada rasio 85:15, kemudian mencapai 90% pada rasio 90:10, 89% pada rasio 80:20, serta tetap berada pada kisaran 87%-88% pada rasio lainnya. Rentang perbedaan performa yang relatif kecil menunjukkan bahwa SVM memiliki kemampuan generalisasi yang baik pada data sentimen pengguna Trans Jogja. Hasil ini sejalan dengan penelitian Ramdhania, Hidayat, dan Sakiawati [2] yang

menunjukkan SVM mampu menghasilkan performa klasifikasi yang tinggi pada analisis sentimen transportasi publik berbasis Twitter. Kemampuan SVM dalam menemukan hyperplane optimal pada data berdimensi tinggi, seperti representasi TF-IDF, menjadikannya efektif untuk klasifikasi teks dengan jumlah fitur yang besar. Random Forest juga menunjukkan performa yang tinggi dan relatif stabil di seluruh skenario pembagian data. Akurasi yang diperoleh berturut-turut sebesar 88% pada rasio 70:30, 88% pada rasio 75:25, 89% pada rasio 80:20, 90% pada rasio 85:15, dan 89% pada rasio 90:10. Meskipun tidak selalu menjadi algoritma dengan akurasi tertinggi pada setiap rasio pembagian data, Random Forest secara konsisten menghasilkan nilai akurasi, presisi, recall, dan f1-score yang seimbang. Temuan ini sejalan dengan penelitian Fahreza dan Handoko [4] mengenai analisis sentimen masyarakat terhadap Trans Semarang yang menunjukkan bahwa Random Forest mampu menghasilkan performa klasifikasi yang baik pada data opini transportasi publik. Stabilitas Random Forest dapat dikaitkan dengan mekanisme ensemble learning yang menggabungkan banyak pohon keputusan sehingga mampu mengurangi risiko overfitting dan meningkatkan kemampuan generalisasi model. Neural Network menunjukkan performa yang kompetitif dan relatif stabil dengan akurasi berkisar antara 88% hingga 90% pada seluruh rasio pembagian data. Performa terbaik diperoleh pada rasio 90:10 dan 85:15 dengan akurasi sebesar 90%, sedangkan pada rasio 80:20, 75:25, dan 70:30 akurasinya berada pada rentang 88%–89%. Hasil ini menunjukkan bahwa Neural Network mampu mempelajari pola non-linear dari representasi TF-IDF dengan baik sehingga tetap menghasilkan performa yang konsisten meskipun proporsi data latih dan data uji berubah. Algoritma Decision Tree menunjukkan performa yang cukup baik, dengan akurasi tertinggi sebesar 89% pada rasio 85:15, diikuti akurasi 88% pada rasio 90:10 dan 80:20, serta 85% pada rasio 75:25 dan 70:30. Namun, nilai recall pada rasio 80:20 mengalami penurunan menjadi 86%, menunjukkan bahwa kemampuan model dalam mengenali seluruh kelas belum sepenuhnya konsisten. Pola ini memperlihatkan karakteristik Decision Tree yang cenderung sensitif terhadap perubahan data pelatihan dan lebih rentan mengalami overfitting dibandingkan metode ensemble seperti Random Forest. Logistic Regression menghasilkan performa yang berada pada tingkat menengah dengan akurasi berkisar antara 80% hingga 86%. Performa terbaik dicapai pada rasio 80:20 dengan akurasi sebesar 86%, sedangkan pada rasio lainnya akurasi berada pada rentang 80%–82%. Hasil ini menunjukkan bahwa Logistic Regression masih mampu melakukan klasifikasi sentimen dengan cukup baik, namun kemampuannya dalam menangkap pola yang kompleks pada data teks berbasis TF-IDF masih berada di bawah SVM, Random Forest, dan Neural Network. Sementara itu, Naive Bayes menunjukkan performa terendah dibandingkan algoritma lainnya di seluruh skenario pengujian. Akurasi yang diperoleh berkisar antara 70% hingga 80%, dengan performa terbaik sebesar 80% pada rasio 85:15 dan performa terendah sebesar 70% pada rasio 80:20 ketika menggunakan Bernoulli Naive Bayes. Temuan ini berbeda dengan penelitian Purnamasari dan Agoestanto [5] pada analisis sentimen pengguna JakLingko yang menunjukkan bahwa Naive Bayes masih dapat menghasilkan performa yang kompetitif. Perbedaan tersebut kemungkinan dipengaruhi oleh karakteristik dataset, distribusi kelas, serta variasi kosakata pada opini pengguna Trans Jogja yang cenderung informal dan tidak selalu memenuhi asumsi independensi antar fitur yang menjadi dasar kerja algoritma Naive Bayes. Secara keseluruhan, performa model meningkat seiring bertambahnya proporsi data training, dengan rasio 90:10 menghasilkan hasil terbaik pada hampir seluruh algoritma. Berdasarkan seluruh metrik yang diuji, Random Forest dapat disimpulkan sebagai model paling unggul.

5 Kesimpulan

Penelitian berhasil melakukan klasifikasi sentimen pengguna layanan Trans Jogja di media sosial X dengan enam algoritma machine learning, yaitu Naive Bayes, SVM, Random Forest, Neural Network, Logistic Regression, dan Decision Tree. Tahapan preprocessing dan pembobotan TF-IDF menghasilkan data teks yang terstruktur dan siap diolah. Hasil pengujian menunjukkan bahwa proporsi data training memengaruhi kinerja model, dengan rasio 85:15 menghasilkan performa terbaik pada sebagian besar algoritma, sedangkan rasio 90:10 juga memberikan performa yang kompetitif pada beberapa model. Berdasarkan metrik evaluasi, SVM, Random Forest, dan Neural Network menunjukkan performa yang paling baik dan konsisten. SVM mencapai akurasi tertinggi sebesar 91% pada rasio 85:15 dan menunjukkan kemampuan generalisasi yang sangat baik, sedangkan Random Forest menghasilkan nilai evaluasi yang tinggi dan relatif seimbang antar kelas dengan distribusi

kesalahan klasifikasi yang lebih merata. Dengan demikian, SVM merupakan algoritma dengan performa tertinggi pada skenario tertentu, sedangkan Random Forest merupakan model yang paling stabil dan seimbang secara keseluruhan untuk klasifikasi sentimen pengguna layanan Trans Jogja.

Referensi

- [1] A. W. Nurfadillah, Kiki Rasmala Sani, “Kebijakan Transportasi Publik dalam meningkatkan Pelayanan terhadap Masyarakat: Studi Kasus Penggunaan Transjogja,” *Sawala : Jurnal Administrasi Negara*, Vol. 11, No. 1, pp. 54–66, Jun. 2023, DOI: 10.30656/SAWALA.V11I1.5834.
- [2] K. F. Ramdhania, D. F. Hidayat, and R. Salkiawati, “Implementasi Metode Naïve Bayes dan Support Vector Machine (SVM) untuk menganalisis Sentimen Pengguna Twitter terhadap Transjakarta,” *JMPM: Jurnal Matematika dan Pendidikan Matematika*, Vol. 9, No. 1, pp. 1–14, Mar. 2024, DOI: 10.26594/JMPM.V9I1.4494.
- [3] F. S. Nufus and W. Gata, “Sentiment Analysis of Public Opinion on Transportation Services in Indonesia using Machine Learning,” *Jurnal Techno Nusa Mandiri*, Vol. 22, No. 2, pp. 143–150, Sep. 2025, DOI: 10.33480/TECHNO.V22I2.6577.
- [4] R. Fachreza and W. T. Handoko, “Analisis Sentimen Masyarakat terhadap Kinerja Trans Semarang menggunakan Metode Random Forest,” *Progresif: Jurnal Ilmiah Komputer*, Vol. 20, No. 2, pp. 724–734, Aug. 2024, DOI: 10.35889/PROGRESIF.V20I2.2093.
- [5] A. Purnamasari and A. Agoestanto, “Modeling of Naïve Bayes and Decision Tree Algorithms to Analyze Sentiment Related to Jaklingko Public Transportation on Social Media X (Twitter),” Vol. 5, No. 1, pp. 67–78, 2024, DOI: 10.30598/ppcst.2024.knmxxii.67-78.
- [6] A. Maulana, I. K. Afifah, A. Mubarrak, K. R. Fauzan, A. Dwintara, and B. P. Zen, “Comparison of Logistic Regression, Multinomialnb, SVM, and K-NN Methods on Sentiment Analysis of Gojek App Reviews on the Google Play Store,” *Jurnal Teknik Informatika (Jutif)*, Vol. 4, No. 6, pp. 1487–1494, Dec. 2023, DOI: 10.52436/1.JUTIF.2023.4.6.863.
- [7] R. Ramadhani, “Analisis Sentimen Ulasan Platform Gojek menggunakan Artificial Neural Network (ANN),” Jun. 2025.
- [8] D. A. Ruriyanto, F. A. Raya, and R. Siswoyo, “Analisis Sentimen Ulasan Publik pada X menggunakan Naive Bayes untuk meningkatkan Kualitas Layanan Transjakarta,” *Journal of Research and Publication Innovation*, Vol. 4, No. 1, 2026, Accessed: May 16, 2026. [Online]. Available: <http://jurnal.portalpublikasi.id/index.php/JORAPI/article/view/2039>
- [9] R. N. Fitriyani, M. Jajuli, and Y. Umaidah, “Analisis Sentimen Pengguna KRL terhadap KAI Commuter Line melalui Sosial Media X dengan Algoritma Naive Bayes Classifier,” *Jurnal Informatika dan Teknik Elektro Terapan*, Vol. 13, No. 3S1, pp. 2830–7062, Oct. 2025, DOI: 10.23960/JITET.V13I3S1.7539.
- [10] S. Zaidan Ali Difyah, E. Arip Winanto, P. Alam Jusia, and F. Ilmu Komputer, “Analisis Sentimen terhadap Tagar Kabur Aja Dulu di Twitter menggunakan Metode Lexicon-based,” *Jurnal PROCESSOR*, Vol. 20, No. 2, Oct. 2025, DOI: 10.33998/PROCESSOR.2025.20.2.2542.
- [11] E. Febrianti, “Evaluasi Kepuasan Pengguna Transportasi Online di Indonesia melalui Analisis Sentimen Platform X dengan Pendekatan Lexicon-based,” *urnal Sistem Informasi Komputer (SIKOM)*, Vol. 1, No. 2, 2024, Accessed: May 16, 2026. [Online]. Available: <http://journal.beaninstitute.id/index.php/sikom/article/view/159>
- [12] B. N. Azmi, A. Hermawan, and D. Avianto, “Analisis Pengaruh Komposisi Data Training dan Data Testing pada Penggunaan PCA dan Algoritma Decision Tree untuk Klasifikasi Penderita Penyakit Liver,” *Jurnal Teknologi Informasi dan Multimedia*, Vol. 4, No. 4, pp. 281–290, Feb. 2023, DOI: 10.35746/JTIM.V4I4.298.
- [13] A. Zein and D. Kusumaningsih, “Analisis Sentimen Review Pengguna Aplikasi BLU BCA pada Play Store menggunakan Algoritma Naïve Bayes,” *senafti.budiluhur.ac.id*, Vol. 4, No. 2, 2025, Accessed: May 16, 2026. [Online]. Available: <https://senafti.budiluhur.ac.id/senafti/article/download/1703/895>
- [14] M. Fadhilla, R. Wandri, A. Hanafiah, P. R. Setiawan, Y. Arta, and S. Daulay, “Analisis Performa Algoritma Machine Learning untuk Identifikasi Depresi pada Mahasiswa,” *Journal*

- of Informatics Management and Information Technology*, Vol. 5, No. 1, pp. 40–47, Jan. 2025, DOI: 10.47065/JIMAT.V5I1.473.
- [15] P. K. Sari and R. R. Suryono, “Komparasi Algoritma *Support Vector Machine* dan *Random Forest* untuk Analisis Sentimen *Metaverse*,” *Jurnal Mnemonic*, Vol. 7, No. 1, pp. 31–39, Feb. 2024, DOI: 10.36040/MNEMONIC.V7I1.8977.
- [16] M. R. Adipratama and N. Safriadi, “Analisis Sentimen terhadap Rencana Penerapan *E-Voting* pada Pemilu di Indonesia,” *JLK*, Vol. 7, No. 1, 2024.
- [17] J. D. P. Beneng, “Analisis Sentimen Pengguna *Twitter* terhadap *AI-Generated Art* menggunakan Metode *Naive Bayes* dan *Neural Network Classifier*,” 2023. Accessed: May 16, 2026. [Online]. Available: <https://eprints.unmer.ac.id/id/eprint/5452/>
- [18] A. Sintawati, F. Amalya, A. Hidayat, and N. M. Supyan, “Evaluasi Kinerja *SVM* dan *Logistic Regression* pada Data *Multiclass* dalam Analisis Sentimen Film *Dirty Vote* dengan Metode Pelabelan *Lexicon based*,” *RIGGS: Journal of Artificial Intelligence and Digital Business*, Vol. 4, No. 2, pp. 821–832, May 2025, DOI: 10.31004/RIGGS.V4I2.576.
- [19] F. R. Valerian, M. Syarief, and D. A. Fatah, “Klasifikasi Tingkat Obesitas menggunakan Metode *GBM* dan *Confusion Matrix*,” *JATI (Jurnal Mahasiswa Teknik Informatika)*, Vol. 9, No. 2, pp. 2242–2249, Mar. 2025, DOI: 10.36040/JATI.V9I2.13062.
- [20] O. Manullang, P. Cahyo, and N. H. Harani, “Analisis Sentimen untuk memprediksi Hasil Calon Pemilu Presiden menggunakan *Lexicon based* dan *Random Forest*,” *Jurnal Ilmiah Informatika*, Vol. 11, No. 02, pp. 159–169, Sep. 2023, DOI: 10.33884/JIF.V11I02.7987.