

Sentiment Analysis on Ferizy Application Reviews Using Support Vector Machine Method

[Analisis Sentimen Pada Ulasan Aplikasi Ferizy Menggunakan Metode Support Vector Machine]

Rizky Septia Putra¹⁾, Rohman Dijaya²⁾

¹⁾ Program Studi Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

²⁾ Program Studi Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

*Email Penulis Korespondensi: rohman.dijaya@umsida.ac.id

Abstract. Ferizy is a sea transportation ticket booking service that can be accessed by users through the official website and the Ferizy application on the Google Play Store. To improve the quality of the application, Ferizy users can provide reviews of application functions and performance through the Google Play Store. Reviews provided by users can be identified as positive or negative sentiments through sentiment analysis with the help of machine learning. In this situation, researchers compared Support Vector Machine, Naive Bayes, and LSTM methods by covering a dataset of 1,500 reviews. Based on the test, the RBF kernel in this study produced the highest accuracy at a training and test data ratio of 90: 10 with accuracy reaching 90.71%. Evaluation of SVM accuracy results using the Confusion Matrix method produces a percentage of 91%. Comparison with Naive Bayes and LSTM methods with the same dataset ratio only produces an accuracy of 83.88% and 85.33%.

Keywords - Ferizy; Google Play Store; Sentiment Analysis; Support Vector Machine

Abstrak. Ferizy merupakan layanan pemesanan tiket transportasi laut yang dapat diakses oleh pengguna melalui situs web resmi dan aplikasi Ferizy di Google Play Store. Untuk meningkatkan kualitas aplikasi, pengguna Ferizy dapat memberikan ulasan terhadap fungsi maupun kinerja aplikasi melalui Google Play Store. Ulasan yang diberikan pengguna dapat diidentifikasi sebagai sentimen positif atau negatif melalui analisis sentimen dengan bantuan machine learning. Dalam situasi ini peneliti membandingkan metode Support Vector Machine, Naive Bayes, dan LSTM dengan mencakup jumlah dataset sebesar 1.500 ulasan. Berdasarkan pengujian, kernel RBF pada penelitian ini menghasilkan akurasi tertinggi pada pembagian rasio data latih dan data uji 90 : 10 dengan akurasi mencapai 90,71%. Evaluasi hasil akurasi SVM menggunakan metode Confusion Matrix menghasilkan presentase sebesar 91%. Perbandingan dengan metode Naive Bayes dan LSTM dengan rasio dataset yang sama hanya menghasilkan akurasi sebesar 83,88% dan 85,33%.

Kata Kunci - Ferizy; Google Play Store; Analisis Sentimen; Support Vector Machine

I. PENDAHULUAN

Dalam upaya untuk meningkatkan kualitas pelayanan kepada pelanggan di sektor transportasi laut, PT. ASDP Indonesia Ferry telah mengambil langkah penting dengan menerapkan transformasi digital pada proses pemesanan dan pembelian tiket feri secara daring melalui peluncuran sistem aplikasi dan situs web resmi bernama Ferizy pada bulan Mei 2020. Platform aplikasi dan situs web Ferizy dibuat dengan maksud untuk menyederhanakan proses pemesanan dan pembelian tiket feri [1]. Saat ini aplikasi Ferizy telah tersedia di Google Play Store dan dapat diunduh serta pengguna dapat memberikan ulasan, kritik, dan saran terhadap fungsi maupun kinerja dari keseluruhan aplikasi.

Menurut informasi yang diperoleh dari Google Play Store pada bulan Agustus 2023, aplikasi Ferizy telah diunduh oleh lebih dari 1 juta pengguna dan menerima rating sebesar 3,4. Ulasan yang tersedia di Google Play Store adalah sumber informasi berharga yang berjumlah besar dan bersifat tidak terstruktur sehingga diperlukan suatu teknik untuk mengetahui bagaimana ulasan pengguna terhadap aplikasi tersebut [2]. Analisis sentimen dapat digunakan untuk mengidentifikasi ulasan pengguna menjadi ulasan sentimen positif dan sentimen negatif. Salah satu metode untuk mengklasifikasikan data ulasan di bidang text mining yaitu Support Vector Machine (SVM). Support Vector Machine (SVM) adalah algoritma yang terkenal karena kemampuannya menghasilkan solusi optimal dalam konteks klasifikasi. Algoritma ini dikembangkan oleh Vapnik sebagai model machine learning yang menggunakan kernel untuk tugas klasifikasi dan regresi [3].

Berdasarkan penelitian sebelumnya telah melibatkan penggunaan metode text mining oleh para peneliti untuk menganalisis sentimen. Dalam penelitian sebelumnya oleh Herlinawati dkk mengenai analisis sentimen Aplikasi Zoom Cloud Meetings terhadap ulasan pengguna menggunakan metode Naive Bayes dan Support Vector Machine. Hasil penelitian menunjukkan bahwa Naive Bayes memiliki tingkat akurasi sebesar 74,37% dengan nilai AUC sekitar 0,659. Sementara itu, metode Support Vector Machine mencapai tingkat akurasi 81,22% dengan nilai AUC sebesar

0,886 [4]. Penelitian ini bertujuan untuk mengukur tingkat akurasi yang dihasilkan menggunakan metode Support Vector Machine (SVM) dalam mengklasifikasikan ulasan kedalam kategori kelas sentimen yaitu sentimen positif dan negatif.

Text Mining

Text Mining merupakan metode yang digunakan untuk mengeksplorasi informasi dari dokumen teks. Teknik ini bertujuan untuk mengidentifikasi pola-pola dalam kumpulan data teks. Lebih lanjut, text mining dapat dijelaskan sebagai sebuah proses yang bertujuan untuk mendapatkan informasi atau pola yang belum pernah teridentifikasi sebelumnya dengan menganalisis dataset dalam skala yang luas [5].

Analisis Sentimen

Analisis sentimen adalah sebuah bidang studi yang memadukan penggunaan pemrosesan bahasa alami, kecerdasan buatan, dan teknik text mining. Analisis sentimen dapat didefinisikan sebagai proses yang terdiri dari pemahaman, ekstraksi, dan pengolahan data teks secara otomatis dengan tujuan menghasilkan informasi yang bernilai [6].

Text Preprocessing

Text preprocessing adalah langkah awal dalam pengolahan teks yang bertujuan untuk mempersiapkan teks agar siap untuk diproses lebih lanjut menjadi data. Tahap text preprocessing melibatkan serangkaian tindakan rutin dan proses yang diperlukan untuk mempersiapkan data sebelum digunakan dalam operasi knowledge discovery sistem text mining. Terdapat beberapa tahapan yang digunakan pada text preprocessing yaitu case folding, tokenizing, normalisasi kata, filtering, dan stemming [7].

1. Case Folding adalah tahap dalam pemrosesan teks yang bertujuan untuk mengonversi seluruh huruf dalam teks menjadi huruf kecil [8].
2. Tokenizing merupakan proses untuk memisahkan kata, simbol, frase, dan entitas penting lainnya atau token dari suatu dokumen [9].
3. Normalisasi kata adalah proses mengubah frasa yang tidak baku menjadi baku dan juga menggantikan singkatan dengan kata-kata aslinya [10].
4. Filtering Stopword adalah tahapan yang bertujuan untuk mengekstraksi kata-kata penting dari hasil token. Pada tahap filtering stopwords melibatkan tindakan penyaringan kata-kata yang kurang memiliki relevansi penting sehingga kemudian dapat dihapus dari teks [11].
5. Stemming merupakan proses untuk mengubah sebuah kata menjadi bentuk kata dasarnya dengan menghilangkan imbuhan yang ada di awalan dan akhiran kata [12].

Term Frequency-Inverse Document Frequency (TF-IDF)

TF-IDF merupakan suatu algoritma yang menggabungkan antara Term Frequency (TF) dengan Inverse Document Frequency (IDF). TF merujuk pada seberapa sering sebuah kata muncul dalam suatu dokumen, sementara IDF digunakan untuk mengurangi pengaruh kata-kata yang umum dan sering muncul di banyak dokumen, dengan mempertimbangkan kebalikan dari frekuensi dokumen yang mengandung kata tersebut [13].

Persamaan yang digunakan untuk menghitung nilai masing – masing dokumen terhadap kata kunci di persamaan (1) sampai dengan persamaan (3) [14].

$$W_{d,t} = tf_{d,t} \times IDF \quad (1)$$

Keterangan :

W = nilai dokumen ke -n, d = dokumen, t = kata kunci

tf = term frequency (jumlah kemunculan kata)

IDF = Inverse Document Frequency

Nilai tf diperoleh dari :

$$tf_d = \frac{\text{Jumlah munculnya kata } t \text{ dalam dokumen}}{\text{Total jumlah seluruh kata dalam dokumen}} \quad (2)$$

Nilai IDF didapatkan dari :

$$IDF = \log \left(\frac{D}{df} \right) \quad (3)$$

Keterangan :

D = total kalimat yang tersedia,

df = jumlah dokumen yang mencakup kata kunci.

Support Vector Machine

Support Vector Machine (SVM) adalah suatu algoritma klasifikasi yang terdapat pada machine learning termasuk dalam kategori supervised learning. SVM digunakan untuk memprediksi kategori atau kelas suatu data berdasarkan pola yang ditemukan selama proses training. Support Vector Machine melakukan klasifikasi dengan cara mencari garis pembatas (hyperlane) terbaik yang berfungsi sebagai pemisah dua kelas data yang dikembangkan oleh Vladimir Vapnik [15].

Pada SVM memiliki beberapa pilihan kernel yang dapat digunakan yaitu kernel Linear, Polynomial, Radial Basis Function (RBF), dan Sigmoid. Pada setiap kernel memiliki tingkat akurasi yang berbeda-beda dalam implementasinya [16].

Confusion Matrix

Confusion matrix adalah suatu teknik untuk mengevaluasi kinerja suatu sistem klasifikasi pada skenario dengan dua kelas. Dalam confusion matrix, terdapat dua istilah, yaitu aktual yang merepresentasikan label dari data asli, dan predicted yang merujuk pada label yang diberikan oleh proses klasifikasi [17].

Confusion matrix direpresentasikan pada tabel 1 dibawah :

Tabel 1. Confusion Matrix

	Aktual	
<i>Predicted</i>		
Positif	TP	FP
Negatif	FN	TN

Hasil confusion matrix selanjutnya akan dipergunakan untuk menghitung accuracy, precision, recall dan F1 score untuk menganalisa kinerja dari algoritma dalam melakukan klasifikasi dengan persamaan (4) sampai dengan persamaan (7) [18].

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (4)$$

$$Precision = \frac{TP}{TP+FP} \quad (5)$$

$$Recall = \frac{TP}{TP+FN} \quad (6)$$

$$F1 \text{ score} = \frac{2.(recall.precision)}{(recall+precision)} \quad (7)$$

Keterangan :

1. True Positive (TP) : Nilai yang diperoleh ketika hasil pengklasifikasian menunjukkan label positif dan data aslinya juga diberi label positif.
2. True Negative (TN) : Nilai yang diperoleh ketika hasil pengklasifikasian menunjukkan label negatif dan data aslinya juga diberi label negatif.
3. False Positive (FP) : Nilai yang diperoleh ketika hasil pengklasifikasian menunjukkan label positif dan data aslinya juga diberi label negatif.
4. False Negative (FN) : Nilai yang diperoleh ketika hasil pengklasifikasian menunjukkan label negatif dan data aslinya juga diberi label positif.

II. METODE

Pengumpulan Data

Dataset dalam penelitian ini diperoleh dengan teknik Web Scraping dari situs web Google Play Store, dengan menggunakan Google Colab yang dapat diakses di Browser. Dataset yang diperoleh pada tanggal 02 April 2022 sampai 24 Agustus 2023 berupa kumpulan ulasan untuk aplikasi Ferizy yang terdapat di Google Play Store. Jumlah data ulasan yang digunakan untuk dataset penelitian ini berjumlah 1.500 data ulasan. Dataset disimpan dalam format excel yang nantinya akan dimasukkan kedalam sistem. Dalam ilustrasi yang disajikan pada gambar 1 merupakan contoh data yang terdapat dalam dataset yang diperoleh melalui teknik web scraping.

no	username	score	at	content
1	Dsatia channel	1	2023-08-24 08:50:48	buka aplikasi loading memuat data terus gak selesai2.
2	itha firman	1	2023-08-23 13:28:54	Aplikasi buluk,nyusahin
3	MR Sutan Alam	5	2023-08-23 13:14:52	Sangat bagus
4	Saya Ganteng	2	2023-08-23 08:00:25	Minimal kalo buat database itu pake jaringan tersendiri
5	Dani Firmansyal	5	2023-08-22 11:47:44	Sangat membantu dalam perjalanan
6	bagas samiaji	1	2023-08-21 11:10:54	Aplikasi Lemott, Tidak Flexible
7	Muhid Abdillah	1	2023-08-21 10:29:36	Lemot parah.
8	Pra Mono	5	2023-08-21 03:03:42	Aplikasi yang sangat membantu penumpang kapal Fer
9	Syahrul Iwan	1	2023-08-20 07:20:13	Kenapa tiket udah pakai system online harus di print o
10	Rokhiman ST	1	2023-08-19 13:44:15	Aplikasi lemot banget kalau mau dibuka,, bukanya me
11	Vickran Audians	1	2023-08-19 08:48:36	Error apk gak di bener2in yg kompeten dong bikin apk

Gambar 1. Contoh Dataset

Pelabelan data dilakukan dengan berdasarkan rating ulasan sebagai acuan sentimen. Ulasan dengan rating 1 dan 2 akan dikategorikan sebagai sentimen negatif, sedangkan ulasan dengan rating 4 dan 5 akan dikategorikan sebagai sentimen positif. Terkait rating 3 akan dianalisis secara manual untuk menentukan apakah ulasan tersebut lebih condong ke sentimen positif atau negatif, berdasarkan penggunaan kata-kata positif dan negatif di dalamnya [2]. Pada tabel 2 menunjukkan jumlah data setelah dilakukan pelabelan berdasarkan rating.

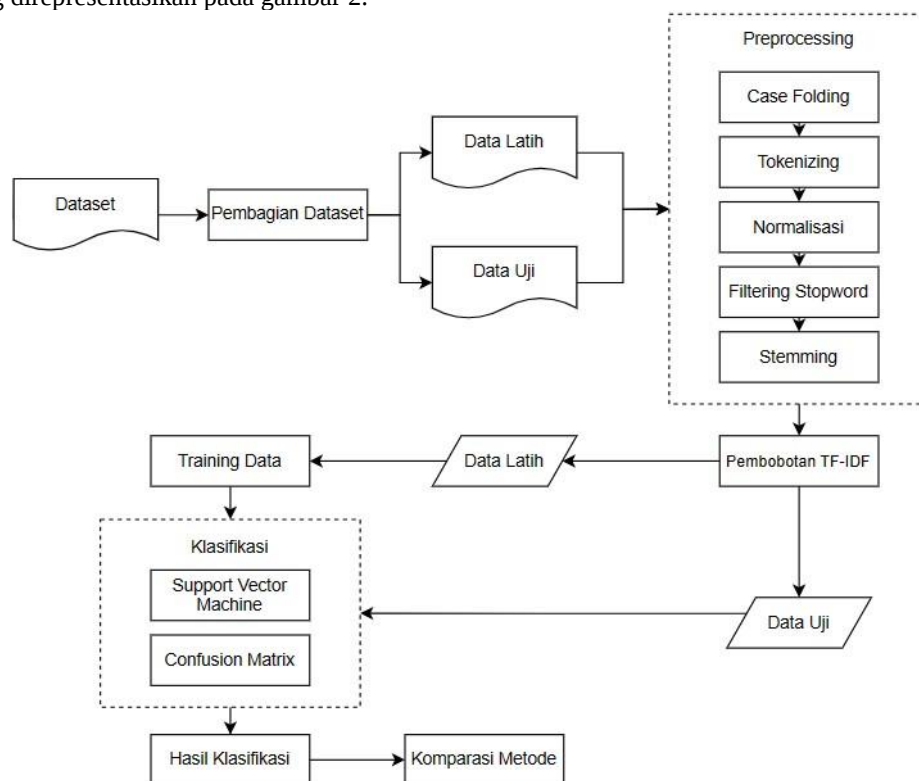
Tabel 2. Jumlah Pelabelan Data

Sentimen	Jumlah Data
Positif	585
Negatif	915
Total	1500

Penulis membagi dataset sebanyak 1500 data menjadi dua bagian, yaitu data latih dan data uji dengan beberapa rasio perbandingan. Hal ini dipilih karena perbandingan antara data latih dan data uji memiliki dampak besar terhadap tingkat akurasi yang akan dicapai.

Flowchart

Dalam penelitian ini, berikut adalah sejumlah kerangka alur yang digunakan untuk menjalankan proses analisis sentimen yang direpresentasikan pada gambar 2.



Gambar 2. Kerangka Alur Penelitian

III. HASIL DAN PEMBAHASAN

Preprocessing Data

Proses ini melibatkan beberapa tahap, mengingat data ulasan tak selalu menggunakan kata baku. Terdapat beberapa tahapan yang digunakan pada text preprocessing.

1. Case Folding

Tahapan case folding bertujuan untuk mengonversi seluruh huruf dalam teks menjadi huruf kecil. Pada tabel 3 menampilkan hasil dari tahap case folding.

Tabel 3. Hasil Tahap Case Folding

Teks	Case Folding
Susah di akses	susah di akses
Kadang suka susah mau bayar	kadang suka susah mau bayar
Aplikasinya bagus dan bermanfaat	aplikasinya bagus dan bermanfaat

2. Tokenizing

Tahapan tokenizing digunakan untuk memisahkan kata, simbol, frase, dan entitas penting lainnya atau token dari suatu dokumen. Pada tabel 4 menampilkan hasil dari tahap tokenizing.

Tabel 4. Hasil Tahap Tokenizing

Teks	Tokenizing
Susah di akses	['susah', 'di', 'akses']
Kadang suka susah mau bayar	['kadang', 'suka', 'susah', 'mau', 'bayar']
Aplikasinya bagus dan bermanfaat	['aplikasinya', 'bagus', 'dan', 'bermanfaat']

3. Normalisasi Kata

Tahapan normalisasi kata akan mengubah frasa yang tidak baku menjadi baku dan juga menggantikan singkatan dengan kata-kata aslinya. Pada tabel 5 menampilkan hasil dari tahap normalisasi.

Tabel 5. Hasil Tahap Normalisasi

Teks	Normalisasi Kata
Susah di akses	['susah', 'di', 'akses']
Kadang suka susah mau bayar	['kadang', 'suka', 'susah', 'mau', 'bayar']
Aplikasinya bagus dan bermanfaat	['aplikasinya', 'bagus', 'dan', 'bermanfaat']

4. Filtering Stopword

Tahapan filtering stopwords bertujuan untuk mengekstraksi kata-kata penting dari hasil token. Pada tabel 6 menampilkan hasil dari tahap filtering stopwords.

Tabel 6. Hasil Tahap Filtering Stopword

Teks	Filtering Stopword
Susah di akses	['susah', 'akses']
Kadang suka susah mau bayar	['kadang', 'suka', 'susah', 'bayar']
Aplikasinya bagus dan bermanfaat	['aplikasi', 'bagus', 'bermanfaat']

5. Stemming

Tahapan stemming digunakan untuk mengubah sebuah kata menjadi bentuk kata dasarnya dengan menghilangkan imbuhan yang ada di awalan dan akhiran kata. Pada tabel 7 menampilkan hasil dari tahap stemming.

Tabel 7. Hasil Tahap Stemming

Teks	Stemming
Susah di akses	['susah', 'akses']
Kadang suka susah mau bayar	['kadang', 'suka', 'susah', 'bayar']
Aplikasinya bagus dan bermanfaat	['aplikasi', 'bagus', 'manfaat']

Setelah tahap preprocessing data selesai, langkah selanjutnya adalah menyimpan hasil preprocessing ini ke dalam sebuah file baru. File tersebut akan digunakan sebagai dataset dalam proses klasifikasi.

Penerapan Metode Support Vector Machine

Dalam metode SVM, terdapat beberapa jenis kernel yang dapat digunakan, antara lain kernel Linear, kernel Polynomial, kernel Radial Basis Function (RBF), dan kernel Sigmoid. Setiap kernel memiliki tingkat akurasi yang berbeda dalam implementasinya [16]. Pada penelitian ini, dataset akan dibagi menjadi tiga perbandingan rasio yang berbeda, dan setiap kernel dari metode SVM akan diujikan pada masing-masing perbandingan tersebut. Berikut adalah hasil perbandingan dari pengujian yang telah dilakukan pada setiap kernel.

1. Hasil pengujian dengan rasio perbandingan data latih dan data uji sebesar 70% : 30%. Pada tabel 8 terdapat beberapa jenis kernel antara lain kernel linear, polynomial, rbf, dan sigmoid untuk mengetahui nilai akurasi. Berdasarkan tabel 8 kernel rbf memiliki nilai akurasi paling tinggi sebesar 87,98%. Sedangkan kernel polynomial memiliki nilai akurasi terendah sebesar 80,15%.

Tabel 8. Rasio Perbandingan Data 70% : 30%

Kernel	Akurasi
Linear	87,07%
Polynomial	80,15%
Radial Basis Function (RBF)	87,98%
Sigmoid	86,52%

2. Hasil pengujian dengan rasio perbandingan data latih dan data uji sebesar 80% : 20%. Pada tabel 9 terdapat beberapa jenis kernel antara lain kernel linear, polynomial, rbf, dan sigmoid untuk mengetahui nilai akurasi. Berdasarkan tabel 9 kernel rbf memiliki nilai akurasi paling tinggi sebesar 90,16%. Sedangkan kernel polynomial memiliki nilai akurasi terendah sebesar 84,15%.

Tabel 9. Rasio Perbandingan Data 80% : 20%

Kernel	Akurasi
Linear	87,16%
Polynomial	84,15%
Radial Basis Function (RBF)	90,16%
Sigmoid	87,71%

3. Hasil pengujian dengan rasio perbandingan data latih dan data uji sebesar 90% : 10%. Pada tabel 10 terdapat beberapa jenis kernel antara lain kernel linear, polynomial, rbf, dan sigmoid untuk mengetahui nilai akurasi. Berdasarkan tabel 10 kernel rbf memiliki nilai akurasi paling tinggi sebesar 90,71%. Sedangkan kernel polynomial memiliki nilai akurasi terendah sebesar 85,25%.

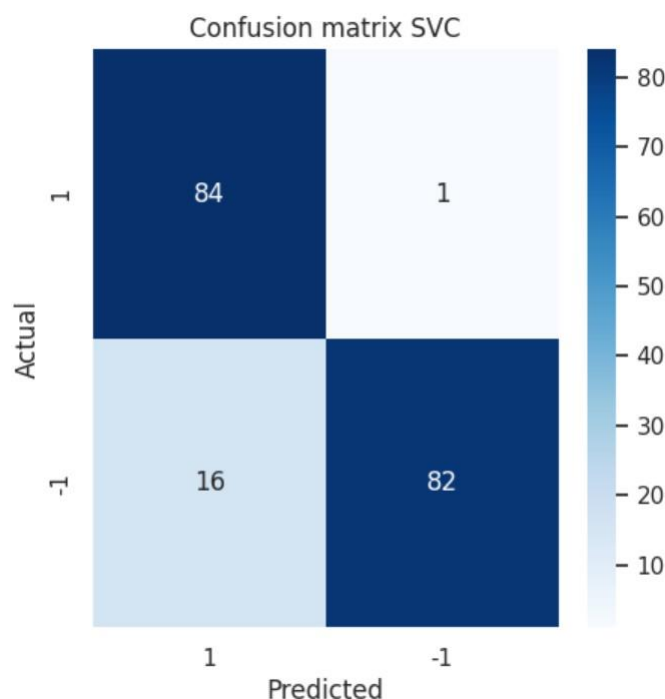
Tabel 10. Rasio Perbandingan Data 90% : 10%

Kernel	Akurasi
Linear	87,43%
Polynomial	85,25%
Radial Basis Function (RBF)	90,71%
Sigmoid	87,43%

Dalam tabel yang telah ditampilkan, terlihat bahwa kernel rbf memiliki tingkat akurasi yang paling tinggi dibandingkan dengan kernel-kernel lainnya. Oleh karena itu, kernel rbf dipilih untuk digunakan dalam proses klasifikasi berikutnya dalam metode SVM. Proses klasifikasi dilakukan dengan membuat confusion matrix untuk mengetahui tingkat akurasi, recall, dan precision. Matriks ini digunakan untuk mengevaluasi performa model yang dibentuk oleh setiap algoritma klasifikasi. Berdasarkan tabel 8 pengujian metode SVM menunjukkan akurasi yang paling tinggi di perbandingan 90% : 10% yaitu dengan nilai akurasi 90,71 %.

Evaluasi Confusion Matrix

Proses evaluasi bertujuan untuk mengukur performa algoritma Support Vector Machine (SVM) dalam menganalisis sentimen. Evaluasi ini menggunakan confusion matrix sebagai alat untuk mengevaluasi precision, f-measure, recall untuk setiap kelas sentimen, serta menghitung akurasi keseluruhan. Hasil confusion matrix dengan rasio perbandingan data sebesar 90% : 10% dapat dilihat pada gambar 3.



Gambar 3. Hasil Confusion Matrix SVM

Seperti yang terlihat pada gambar 3, confusion matrix berupa matriks dengan ukuran 2x2, dimana setiap kolomnya mewakili setiap sentimen yaitu positif (1) dan negatif (-1). Apabila ingin mendapatkan nilai true positive, true negative, false positive dan false negative dalam confusion matrix dengan ukuran matriks 2x2 di setiap kelas, maka akan menjadi seperti pada tabel 11.

Tabel 11. Keterangan Confusion Matrix

Confusion Matrix	Jumlah Data
True Positive (TP)	84
False Positive (FP)	1
True Negative (TN)	82
False Negative (FN)	16

Berdasarkan hasil analisis sentimen terhadap data ulasan pada tabel 8, 9, dan 10, evaluasi dilakukan dengan menghitung classification report sebagai berikut :

1. Perhitungan Nilai Precision :

$$\text{Precision Sentimen Positif} = \frac{TP}{TP + FP} = \frac{84}{84 + 1} = 0,9882$$

$$\text{Precision Sentimen Negatif} = \frac{TN}{TN + FN} = \frac{82}{82 + 16} = 0,8367$$

2. Perhitungan Nilai Recall :

$$\text{Recall Sentimen Positif} = \frac{TP}{TP + FN} = \frac{84}{84 + 16} = 0,84$$

$$\text{Recall Sentimen Negatif} = \frac{TN}{TN + FP} = \frac{82}{82 + 1} = 0,988$$

3. Perhitungan Nilai F₁-Score :

$$F_1\text{-Score Sentimen Positif} = 2 \times \frac{(\text{recall} \times \text{precision})}{(\text{recall} + \text{precision})} = 2 \times \frac{0,84 \times 0,9882}{0,84 + 0,9882} = 0,908$$

$$F_1\text{-Score Sentimen Negatif} = 2 \times \frac{(\text{recall} \times \text{precision})}{(\text{recall} + \text{precision})} = 2 \times \frac{0,988 \times 0,8367}{0,988 + 0,8367} = 0,9061$$

4. Perhitungan Nilai Akurasi :

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} = \frac{84 + 82}{84 + 82 + 1 + 16} = 0,9071 \times 100\% = 90,71\%$$

	precision	recall	f1-score	support
-1	0.84	0.99	0.91	85
1	0.99	0.84	0.91	98
accuracy			0.91	183
macro avg	0.91	0.91	0.91	183
weighted avg	0.92	0.91	0.91	183

Gambar 4. Evaluasi Confusion Matrix SVM

Berdasarkan hasil evaluasi dengan menggunakan metode Confusion Matrix, diperoleh nilai precision untuk sentimen negatif (-1) dan positif (1) antara lain sebesar 84% dan 99%. Hasil recall sentimen negatif (-1) dan positif (1) sebesar 99%, dan 84%. Sedangkan untuk nilai dari f1-score untuk sentimen negatif (-1) dan positif (1) berturut-turut sebesar 91%, dan 91%. Melalui evaluasi dengan menggunakan metode Confusion Matrix pada gambar 4, metode SVM mampu memperoleh accuracy sebesar 91%.

Klasifikasi Metode Naive Bayes

Hasil pengujian metode naive bayes menggunakan beberapa rasio pembagian dataset. Pada tabel 12 terdapat tiga rasio pembagian dataset antara lain 70% : 30%, 80% : 20%, dan 90% : 10% untuk mengetahui nilai akurasi. Berdasarkan tabel 12 rasio pembagian 80% : 20% memiliki nilai akurasi paling tinggi sebesar 83,88%. Sedangkan rasio pembagian 90% : 10% memiliki nilai akurasi terendah sebesar 83,06%.

Tabel 12. Hasil Pengujian Naive Bayes

Rasio Pembagian Dataset	Akurasi
Data Latih 70% : Data Uji 30%	83,42%
Data Latih 80% : Data Uji 20%	83,88%
Data Latih 90% : Data Uji 10%	83,06%

Selanjutnya untuk mengevaluasi hasil akurasi Naive bayes, akan digunakan metode confusion matrix yang ditampilkan pada gambar 5.

	precision	recall	f1-score	support
-1	0.78	0.91	0.84	169
1	0.91	0.78	0.84	197
accuracy			0.84	366
macro avg	0.84	0.84	0.84	366
weighted avg	0.85	0.84	0.84	366

Gambar 5. Evaluasi Confusion Matrix Naive Bayes

Berdasarkan hasil evaluasi dengan menggunakan metode Confusion Matrix, diperoleh nilai precision untuk sentimen negatif (-1) dan positif (1) antara lain sebesar 78% dan 91%. Hasil recall sentimen negatif (-1) dan positif (1) sebesar 91%, dan 78%. Sedangkan untuk nilai dari f1-score untuk sentimen negatif (-1) dan positif (1) berturut-turut sebesar 84%, dan 84%. Melalui evaluasi dengan menggunakan metode Confusion Matrix pada gambar 6, metode Naive Bayes mampu memperoleh accuracy sebesar 84%.

Metode LSTM

Hasil pengujian metode LSTM menggunakan beberapa rasio pembagian dataset. Pada tabel 13 terdapat tiga rasio pembagian data antara lain 70% : 30%, 80% : 20%, dan 90% : 10% untuk mengetahui nilai akurasi. Berdasarkan tabel 13 rasio pembagian 90% : 10% memiliki nilai akurasi paling tinggi sebesar 85,33%. Sedangkan rasio pembagian 70% : 30% memiliki nilai akurasi terendah sebesar 84,22%.

Tabel 13. Hasil Pengujian LSTM

Rasio Pembagian Dataset	Akurasi
Data Latih 70% : Data Uji 30%	84,22%
Data Latih 80% : Data Uji 20%	85%
Data Latih 90% : Data Uji 10%	85,33%

Selanjutnya untuk mengevaluasi hasil akurasi metode LSTM, akan digunakan metode confusion matrix yang ditampilkan pada gambar 6.

	precision	recall	f1-score	support
0	0.85	0.92	0.89	93
1	0.86	0.74	0.79	57
accuracy			0.85	150
macro avg	0.85	0.83	0.84	150
weighted avg	0.85	0.85	0.85	150

Gambar 6. Evaluasi Confusion Matrix LSTM

Berdasarkan hasil evaluasi dengan menggunakan metode Confusion Matrix, diperoleh nilai precision untuk sentimen negatif (0) dan positif (1) antara lain sebesar 85% dan 86%. Hasil recall sentimen negatif (0) dan positif (1) sebesar 92%, dan 74%. Sedangkan untuk nilai dari f1-score untuk sentimen negatif (0) dan positif (1) berturut-turut sebesar 89%, dan 79%. Melalui evaluasi dengan menggunakan metode Confusion Matrix pada gambar 6, metode LSTM mampu memperoleh accuracy sebesar 85%.

IV. KESIMPULAN

Berdasarkan hasil pengujian, ditemukan bahwa penggunaan kernel RBF pada penelitian ini menghasilkan akurasi tertinggi dibandingkan dengan kernel lainnya. Oleh karena itu, penggunaan metode Support Vector Machine dengan total data sebanyak 1.500 data berhasil mencapai tingkat akurasi yang signifikan, yaitu mencapai 90,71% dan evaluasi hasil pengujian dengan menggunakan Confusion Matrix juga memvalidasi nilai akurasi sebesar 91% pada pembagian dataset dengan nilai rasio perbandingan 90% : 10% serta didapatkan akurasi sebesar berturut-turut 87,98% dan 90,16% untuk rasio perbandingan 70% : 30% dan 80% : 20%. Sedangkan hasil pengujian metode Naive Bayes dan LSTM berturut-turut mendapatkan tingkat akurasi sebesar 83,88% dan 85,33%. Berdasarkan penelitian kali ini dapat disimpulkan bahwa metode Support Vector Machine mampu menghasilkan akurasi lebih tinggi dibanding Naive Bayes dan LSTM pada rasio dan jumlah dataset yang sama.

UCAPAN TERIMA KASIH

Penulis dengan tulus berterima kasih kepada semua yang telah memberikan kontribusi berharga dalam penelitian ini. Ucapan terima kasih yang mendalam kami sampaikan kepada :

1. Kedua orang tua penulis yang sudah mendukung dan memberikan semangat selama penyusunan karya tulis ini.
2. Rekan mahasiswa Program Studi Informatika Universitas Muhammadiyah Sidoarjo.

REFERENSI

- [1] C. Malisi, "ANALISIS KINERJA SISTEM RESERVASI DAN PEMBAYARAN TIKET FERI (FERIZY) DENGAN METODE IMPORTANCE PERFORMANCE ANALYSIS (IPA) PADA PT. ASDP INDONESIA FERRY," 2021.
- [2] S. Fransiska dan A. Irham Gufroni, "Sentiment Analysis Provider by.U on Google Play Store Reviews with TF-IDF and Support Vector Machine (SVM) Method," *Scientific Journal of Informatics*, vol. 7, no. 2, hlm. 2407–7658, 2020, [Daring]. Tersedia pada: <http://journal.unnes.ac.id/nju/index.php/sji>

- [3] T. M. Permata Aulia, N. Arifin, dan R. Mayasari, "Perbandingan Kernel Support Vector Machine (Svm) Dalam Penerapan Analisis Sentimen Vaksinasi Covid-19," *SINTECH (Science and Information Technology) Journal*, vol. 4, no. 2, hlm. 139–145, 2021, doi: 10.31598/sintechjournal.v4i2.762.
- [4] N. Herlinawati, Y. Yuliani, S. Faizah, W. Gata, dan S. Samudi, "Analisis Sentimen Zoom Cloud Meetings di Play Store Menggunakan Naïve Bayes dan Support Vector Machine," *CESS (Journal of Computer Engineering, System and Science)*, vol. 5, no. 2, hlm. 293, 2020, doi: 10.24114/cess.v5i2.18186.
- [5] Z. Alhaq, A. Mustopa, S. Mulyatun, dan J. D. Santoso, "PENERAPAN METODE SUPPORT VECTOR MACHINE UNTUK ANALISIS Informatika Universitas Amikom Yogyakarta Abstraksi Keywords : Pendahuluan Tinjauan Pustaka," *Jurnal of Information System Management*, vol. Vol. 3, no. 2, hlm. 44–49, 2021, [Daring]. Tersedia pada: <https://jurnal.amikom.ac.id/index.php/joism/article/view/558>
- [6] M. I. H. A. D. Akbari, A. Novianty, dan C. Setianingsih, "Analisis Sentimen Menggunakan Metode Learning Vector Quantization Sentiment Analysis Using Learning Vector Quantization Method," 2017.
- [7] L. Ardiani, H. Sujaini, dan T. Tursina, "Implementasi Sentiment Analysis Tanggapan Masyarakat Terhadap Pembangunan di Kota Pontianak," *Jurnal Sistem dan Teknologi Informasi (Justin)*, vol. 8, no. 2, hlm. 183, Apr 2020, doi: 10.26418/justin.v8i2.36776.
- [8] A. C. Pradikdo dan A. Ristyawan, "Model Klasifikasi Abstrak Skripsi Menggunakan Text Mining Untuk Pengkategorian Skripsi Sesuai Bidang Kajian," *Jurnal SIMETRIS*, vol. 9, no. 2, hlm. 1091–1098, 2018.
- [9] B. B. Baskoro, I. Susanto, dan S. Khomsah, "Analisis Sentimen Pelanggan Hotel di Purwokerto Menggunakan Metode Random Forest dan TF-IDF (Studi Kasus: Ulasan Pelanggan Pada Situs TRIPADVISOR)," *INISTA (Journal of Informatics Information System Software Engineering and Applications)*, vol. 3, no. 2, hlm. 21–29, 2021, doi: 10.20895/INISTA.V3I2.
- [10] S. Khairunnisa, A. Adiwijaya, dan S. Al Faraby, "Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi COVID-19)," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 5, no. 2, hlm. 406, Apr 2021, doi: 10.30865/mib.v5i2.2835.
- [11] D. Alita, Y. Fernando, dan H. Sulistiani, "IMPLEMENTASI ALGORITMA MULTICLASS SVM PADA OPINI PUBLIK BERBAHASA INDONESIA DI TWITTER," *Jurnal TEKNOKOMPAK*, vol. 14, no. 2, hlm. 86, 2020.
- [12] U. Khaira, R. Johanda, P. E. P. Utomo, dan T. Suratno, "Sentiment Analysis Of Cyberbullying On Twitter Using SentiStrength," *Indonesian Journal of Artificial Intelligence and Data Mining*, vol. 3, no. 1, hlm. 21, Mei 2020, doi: 10.24014/ijaidm.v3i1.9145.
- [13] M. Nurjannah dan I. Fitri Astuti, "PENERAPAN ALGORITMA TERM FREQUENCY-INVERSE DOCUMENT FREQUENCY (TF-IDF) UNTUK TEXT MINING," *Jurnal Informatika Mulawarman*, vol. 8, no. 3, hlm. 110–113, 2013.
- [14] N. L. Ratniasih, M. Sudarma, dan N. Gunantara, "Penerapan Text Mining Dalam Spam Filtering Untuk Aplikasi Chat," *Majalah Ilmiah Teknologi Elektro*, vol. 16, no. 3, hlm. 13, 2017, doi: 10.24843/mite.2017.v16i03p03.
- [15] V. I. Santoso, G. Virginia, dan Y. Lukito, "PENERAPAN SENTIMENT ANALYSIS PADA HASIL EVALUASI DOSEN DENGAN METODE SUPPORT VECTOR MACHINE," 2017.
- [16] F. F. Irfani, M. Triyanto, A. D. Hartanto, dan Kusnawi, "Analisis Sentimen Review Aplikasi Ruangguru Menggunakan Algoritma Support Vector Machine," *JBMI (Jurnal Bisnis, Manajemen, dan Informatika)*, vol. 16, no. 3, hlm. 258–266, 2020, doi: 10.26487/jbmi.v16i3.8607.
- [17] D. P. Daryfayi Edyt dan I. Asror, "Sentimen Analisis pada Ulasan Google Play Store Menggunakan Metode Naïve Bayes," *Sentimen Analisis pada Ulasan Google Play Store Menggunakan Metode Naïve Bayes*, vol. 7, no. Ulasan Pada Google Play Store, hlm. 11, 2020.
- [18] F. D. Ananda dan Y. Pristyanto, "Analisis Sentimen Pengguna Twitter Terhadap Layanan Internet Provider Menggunakan Algoritma Support Vector Machine," *MATRIK : Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, vol. 20, no. 2, hlm. 407–416, Mei 2021, doi: 10.30812/matrik.v20i2.1130.

Conflict of Interest Statement:

The author declares that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.