

Perbandingan Performa *Rating-Based Labeling* dan *Lexicon-Based* pada Analisis Sentimen Ulasan Mobile JKN dengan SVM

Alya Nurjannah^{a,1*}, Wahyuni^{a,2}, Ahmad Fahrijal Pukeng^{a,3}

^a STMIK Widya Cipta Dharma, Jl. M. Yamin No.25, Gn. Kelua, Kec. Samarinda Ulu, Kota Samarinda, Kalimantan Timur, Indonesia

e-mail: ¹ 2241033@wicida.ac.id, ² wahyuni@wicida.ac.id,
³ pukeng@wicida.ac.id
*Corresponding author

Submitted Date: April 08th, 2026
Revised Date: May 14th, 2026

Reviewed Date: April 27th, 2026
Accepted Date: June 15th, 2026

Abstract

This study aims to identify user sentiment toward the Mobile JKN application by comparing lexicon-based labeling using the InSet sentiment lexicon with rating-based labeling derived from users' star ratings. The dataset consisted of 50,000 Google Play Store reviews, which were reduced to 34,317 valid reviews after preprocessing. Sentiment classification was performed using the Support Vector Machine (SVM) algorithm with CountVectorizer for feature extraction, while class imbalance was handled using the Synthetic Minority Over-sampling Technique (SMOTE). The lexicon-based approach achieved 99% accuracy, precision, recall, and F1-score. In contrast, the rating-based approach achieved 76% accuracy, 62% precision, 61% recall, and a macro-average F1-score of 60%, mainly due to lower performance in classifying neutral reviews. Labeling validity was evaluated using 300 manually annotated reviews created by two annotators as the ground truth. The results indicate that lexicon-based labeling is more effective than rating-based labeling for SVM-based sentiment classification. WordCloud visualization shows that negative sentiment is dominated by technical issues, whereas positive sentiment mainly reflects easier access to healthcare services.

Keywords: Mobile JKN; Sentiment Analysis; Support Vector Machine; Rating-Based; Lexicon-Based

Abstrak

Studi ini bertujuan untuk mengidentifikasi sentimen pengguna terhadap aplikasi Mobile JKN dengan membandingkan pelabelan berbasis *lexicon* menggunakan *lexicon-based InSet* dengan pelabelan berbasis *Rating-based* yang berasal dari peringkat bintang pengguna. Dataset terdiri dari 50.000 ulasan Google Play Store, yang dikurangi menjadi 34.317 ulasan valid setelah pra-pemrosesan. Klasifikasi sentimen dilakukan menggunakan *algoritma Support Vector Machine* (SVM) dengan *CountVectorizer* untuk ekstraksi fitur, sementara ketidakseimbangan kelas ditangani menggunakan *Synthetic Minority Over-sampling Technique* (SMOTE). Pendekatan berbasis *lexicon* mencapai *accuracy*, *precision*, *recall*, dan *F1-score* sebesar 99%. Sebaliknya, pendekatan berbasis peringkat mencapai *accuracy* 76%, *precision* 62%, *recall* 61%, dan *F1-score* rata-rata makro sebesar 60%, terutama karena kinerja yang lebih rendah dalam mengklasifikasikan ulasan netral. Validitas pelabelan dievaluasi menggunakan 300 ulasan yang dianotasi secara manual yang dibuat oleh dua *annotator* sebagai kebenaran dasar. Hasil penelitian menunjukkan bahwa pelabelan berbasis *lexicon* lebih efektif daripada pelabelan berbasis peringkat untuk klasifikasi sentimen berbasis SVM. Visualisasi *WordCloud* menunjukkan bahwa sentimen negatif didominasi oleh masalah teknis, sedangkan sentimen positif terutama mencerminkan kemudahan akses ke layanan kesehatan.

Kata kunci: Mobile JKN; Analisis Sentimen; *Support Vector Machine*; *Rating-Based*; *Lexicon-Based*

1. Pendahuluan

Salah satu prioritas utama pemerintah dalam upayanya meningkatkan kualitas layanan yang diberikan pada masyarakat umum adalah transformasi digital di sektor kesehatan. Diharapkan pemanfaatan teknologi informasi akan meningkatkan efisiensi prosedur administratif, mendorong transparansi, serta meningkatkan kepuasan pelanggan pada layanan yang diberikan. Dalam konteks tersebut, BPJS Kesehatan sebagai penyelenggara program JKN mengadakan aplikasi Mobile JKN yang diluncurkan secara resmi pada 15 November 2017 sebagai sarana layanan digital untuk para peserta program JKN. Aplikasi Mobile JKN dirancang sebagai sarana layanan mandiri (*self-service platform*) yang memungkinkan peserta untuk mengakses berbagai layanan kesehatan secara daring, seperti pengecekan status kepesertaan, perubahan data, pendaftaran antrean fasilitas kesehatan, hingga penyampaian pengaduan [1], Mobile JKN dikembangkan untuk mempermudah akses layanan tanpa harus datang ke kantor BPJS Kesehatan. Hingga tahun 2026, Mobile JKN telah diunduh lebih dari 50 juta kali dengan lebih dari 960 ribu ulasan dan mendapatkan rating rata-rata yakni 4,4 di Google Play Store [2], Tingginya jumlah unduhan dan ulasan menunjukkan luasnya penggunaan aplikasi serta tingginya interaksi pengguna yang mencerminkan beragam pengalaman dan tingkat kepuasan terhadap performa aplikasi.

Penelitian sebelumnya telah menggunakan berbagai metode, baik pada tahap pelabelan maupun klasifikasi. Metode pelabelan yang umum digunakan meliputi manual *labeling* dan *lexicon-based labeling* [3], sedangkan algoritma klasifikasi yang banyak diterapkan meliputi *Naïve Bayes*, *Random Forest*, algoritma *Support Vector Machine* (SVM), serta metode *deep learning* untuk mengidentifikasi polaritas sentimen [4][5][6]. Penelitian ini berfokus pada perbandingan metode *rating-based* dan *lexicon-based labeling* dengan algoritma *Support Vector Machine* (SVM) sebagai model klasifikasi. Dalam analisis sentimen berbasis *machine learning*, tahap pelabelan data menjadi proses penting sebelum klasifikasi dilakukan, dan algoritma *Support Vector Machine* dipilih karena efektif untuk menangani dataset berdimensi besar dengan stabilitas kinerja yang terjaga. Penelitian yang dilakukan oleh [7] menunjukkan bahwa algoritma *Support Vector Machine* (SVM) memperoleh akurasi sebesar 83,3% pada analisis sentimen ulasan aplikasi Kredivo, sehingga menunjukkan bahwa algoritma *Support Vector Machine* (SVM) memiliki kemampuan yang baik dalam mengklasifikasikan sentimen pada dataset tersebut.

Rating-Based memanfaatkan skor bintang sebagai representasi sentimen sehingga proses pelabelan menjadi cepat, namun skor pengguna tidak selalu sesuai dengan isi ulasan. Sebaliknya, *lexicon-based* menentukan sentimen berdasarkan polaritas kata dalam teks, tetapi sangat bergantung pada kualitas kamus sentimen yang digunakan [8].

Perbedaan karakteristik kedua metode tersebut dapat menghasilkan ketidakkonsistenan pelabelan yang memengaruhi proses pembelajaran model klasifikasi. Pada *rating-based*, ketidaksesuaian antara rating dan isi ulasan berpotensi menimbulkan *noise*, sedangkan pada *lexicon-based* keterbatasan kamus dapat menyebabkan kesalahan interpretasi konteks.

Berdasarkan permasalahan tersebut, penelitian ini mengimplementasikan analisis sentimen dengan membandingkan dua pendekatan pelabelan, yaitu *rating-based labeling* dan *lexicon-based labeling*, menggunakan algoritma *Support Vector Machine* (SVM). Selain itu, teknik *Synthetic Minority Over-sampling Technique* (SMOTE) diterapkan pada data pelatihan guna mengatasi ketidakseimbangan distribusi kelas (*class imbalance*), sehingga distribusi data menjadi lebih proporsional dan kinerja model meningkat. Penelitian sebelumnya yang dilakukan oleh [8], menunjukkan bahwa pendekatan *lexicon-based* menggunakan kamus InSet memperoleh akurasi sebesar 89,7%, lebih tinggi dibandingkan pendekatan *rating-based* yang memperoleh akurasi sebesar 87% pada analisis sentimen ulasan aplikasi GoBiz. Namun, penelitian tersebut dilakukan pada domain aplikasi bisnis dengan skala dataset yang lebih kecil dan tanpa pembahasan mendalam mengenai penanganan ketidakseimbangan kelas (*class imbalance*). Sejauh ini, belum terdapat studi yang menilai secara menyeluruh kedua pendekatan pelabelan tersebut pada domain layanan kesehatan publik dengan dataset skala besar yang memperhitungkan teknik *balancing* data.

Berdasarkan hasil identifikasi terhadap kesenjangan pada penelitian terdahulu, penelitian ini menggunakan 50.000 ulasan yang diperoleh dari Google Play Store. Proses *preprocessing* menghasilkan 34.317 data bersih yang kemudian digunakan pada tahap pemodelan. Untuk mengurangi dampak ketidakseimbangan kelas, metode *Synthetic Minority Over-sampling Technique* (SMOTE) diterapkan pada data latih. Selanjutnya, performa model dievaluasi menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score* agar diperoleh penilaian yang lebih menyeluruh terhadap kemampuan klasifikasi model.



Berbeda dengan penelitian sebelumnya yang umumnya berfokus pada pencapaian performa klasifikasi berdasarkan satu metode pelabelan[9], penelitian ini memberikan penguatan melalui evaluasi terhadap dua pendekatan pelabelan, yaitu *rating-based* dan *lexicon-based*, serta melakukan pengujian kesesuaian label menggunakan *ground truth* hasil anotasi manual. Pendekatan ini memungkinkan penelitian tidak hanya menilai kemampuan model dalam melakukan klasifikasi sentimen, tetapi juga mengevaluasi kualitas label yang digunakan sebagai dasar pembelajaran model. Dengan demikian, penelitian ini memberikan kontribusi dalam memperjelas hubungan antara metode pelabelan, kualitas data, dan performa algoritma *Support Vector Machine* (SVM) pada analisis sentimen ulasan aplikasi layanan kesehatan digital.

Dengan demikian, penelitian ini bertujuan membandingkan performa *rating-based labeling* dan *lexicon-based labeling* untuk mengevaluasi bagaimana karakteristik masing-masing metode pelabelan memengaruhi keterpelajaran (*learnability*) label oleh algoritma *Support Vector Machine* (SVM) pada ulasan aplikasi Mobile JKN.

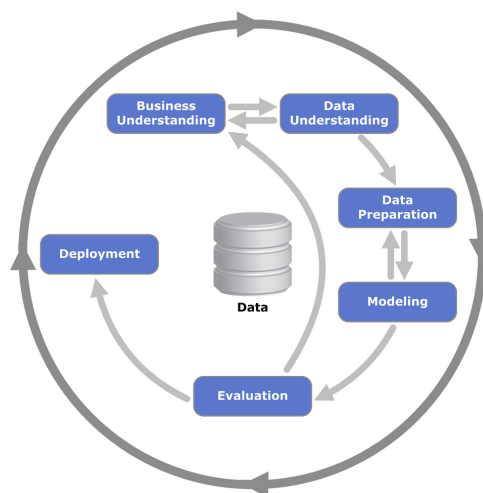
2. Metode Penelitian

Support Vector Machine (SVM) merupakan algoritma *machine learning* untuk klasifikasi yang memaksimalkan margin antarkelas sehingga menghasilkan *hyperplane* pemisah yang optimal [10], SVM bekerja dengan menentukan *hyperplane* optimal yang memaksimalkan margin antar kelas untuk meningkatkan akurasi klasifikasi. Melalui kernel function, algoritma ini juga mampu menangani data nonlinier serta memiliki kemampuan generalisasi yang baik sehingga banyak digunakan dalam analisis sentimen berbasis teks [11].

a. CRISP-DM

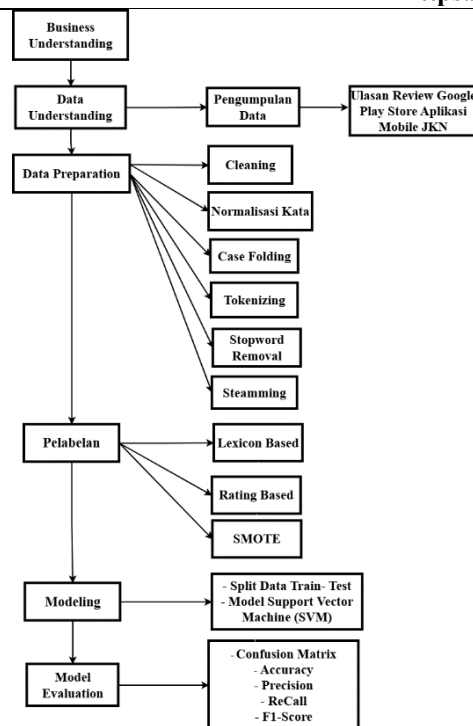
Metode CRISP-DM (*Cross Industry Standard Process for Data Mining*) memiliki sejumlah tahapan yaitu: *Business Understanding* (tahapan memahami proses bisnis), *Data Understanding* (tahap memahami data), *Data Preparation* (tahap mempersiapkan data), *Modeling* (pemodelan), *Evaluation* (evaluasi) dan *Deployment* (tahap implementasi) [12].

Penelitian ini tidak mencakup tahap *Deployment* karena berfokus pada perbandingan performa *rating-based labeling* dan *lexicon-based labeling* dalam analisis sentimen menggunakan algoritma *Support Vector Machine* (SVM). Oleh karena itu, penelitian difokuskan pada proses pengujian dan evaluasi model untuk menilai kinerja masing-masing metode pelabelan [13].



Gambar 1. Tahapan CRISP-DM Sumber: [14]

Penelitian diawali dengan pengumpulan dan *preprocessing* data ulasan pengguna untuk menghasilkan dataset yang siap dianalisis. Data diperoleh melalui *web scraping* ulasan aplikasi Mobile JKN di Google Play Store, kemudian dikelompokkan ke dalam tiga kelas sentimen, yaitu positif, netral, dan negatif. Tahapan penelitian disajikan pada Gambar 2.



Gambar 2. Tahapan Penelitian

1. Business Understanding

Tahap *Business Understanding* difokuskan pada identifikasi permasalahan dalam ulasan pengguna Mobile JKN di Google Play Store. Analisis sentimen digunakan untuk mengelompokkan opini pengguna sebagai dasar evaluasi layanan digital. Selanjutnya, penelitian ini membandingkan *rating-based Labeling* dan *Lexicon-Based Labeling* menggunakan algoritma *Support Vector Machine* (SVM) untuk menentukan metode pelabelan dengan kinerja terbaik.

Hasil analisis diharapkan dapat membantu BPJS Kesehatan dalam mengevaluasi layanan digital berbasis data. Penelitian ini juga membandingkan metode *rating-based* dan *lexicon-based* untuk menentukan pendekatan dengan performa klasifikasi terbaik menggunakan algoritma *Support Vector Machine* (SVM).

2. Data Understanding

Data penelitian dikumpulkan melalui *web scraping* ulasan aplikasi Mobile JKN di Google Play Store menggunakan *library* google-play-scraper. Sebanyak 50.000 ulasan berhasil diperoleh, kemudian melalui proses *data cleaning* sehingga tersisa 34.317 data yang digunakan dalam penelitian. Dataset terdiri atas lima atribut, yaitu Review ID, Username, Rating, Review Text, dan Date, yang digunakan pada proses pelabelan, pembagian data latih dan data uji, serta pembangunan model *Support Vector Machine* (SVM).

Tabel 1. Hasil Pengumpulan Data

No	Review Text	Rating
12777	aplikasi, sangat membantu dan mempermudah dlm pelayanan Kesehatan bagi pesertanya.	5
49980	sampah bgt habis lebaran bukannya bagus makin gak jelas apk ini, login sampe coba di 2 perangkat android dan ios gak bisa login semua tombol login gak fungsi.	2
49992	aplikasi ga jelas Cuma nyusahin ga bermutu, udah coba daftar pake hp saya dan nik saya tapi anak saya ga masuk aplikasi padahal masih 1 KK nyoba daftar lagi pake nik anak hp anak no baru, pas buka aplikasi ada gambar tanda gembok semua ga bisa dibuka tulisannya belum verifikasi padahal jelas2 udah verifikasi pulsanya juga baru di isi, nyoba lagi verifikasi tapi tulisannya nomer hp kosong astagfirullah tobaaaat dah coba aplikasi perbaiki bikin versi lebih canggih kaya aplikasi orang2 luar gitulohh	1

a. Wordcloud

Pada penelitian ini, *WordCloud* dimanfaatkan sebagai metode visualisasi untuk menampilkan pola kemunculan kata dalam data teks. Semakin tinggi frekuensi suatu kata, semakin besar ukuran tampilannya, sehingga kata-kata yang paling dominan dalam ulasan dapat diidentifikasi dengan lebih mudah [15].

Dalam visualisasi tersebut, *WordCloud* sentimen positif, kata-kata seperti bagus, baik, mudah, akses, dan daftar muncul dengan frekuensi tinggi, yang menunjukkan bahwa sebagian pengguna merasakan kemudahan dalam menggunakan aplikasi Mobile JKN. Di sisi lain, *WordCloud* untuk sentimen negatif memperlihatkan dominasi kata-kata seperti *error*, susah, daftar, login, verifikasi, dan wajah, yang mengindikasikan bahwa keluhan pengguna lebih banyak berkaitan dengan kendala teknis pada proses registrasi, autentikasi, dan verifikasi identitas. Temuan ini menunjukkan bahwa peningkatan kualitas layanan aplikasi sebaiknya difokuskan pada stabilitas sistem, penyederhanaan proses pendaftaran, serta peningkatan keandalan fitur verifikasi wajah agar pengalaman pengguna menjadi lebih baik.



Gambar 3. Wordcloud Positif



Gambar 4. WordCloud Negatif

3. Data Preparation

Tahap *data preparation* bertujuan menyiapkan dataset agar siap digunakan pada proses pelabelan dan klasifikasi. Proses ini mencakup *cleansing*, *case folding*, *tokenizing*, *stopword removal*, dan *stemming* untuk meningkatkan kualitas data sebelum dianalisis menggunakan algoritma *machine learning*.

a. Cleaning

Cleaning dilakukan dengan menghapus emoji, angka, dan karakter khusus untuk mengurangi noise sehingga data lebih siap digunakan pada proses ekstraksi fitur.

Tabel 2. Hasil Cleaning Data

Review Text	Cleaning
Akses layanan cepat, sangat memuaskan ❤️❤️❤️	Akses layanan cepat sangat memuaskan
verif nomor satu jam gak ada muncul kode nya, gak jelas!!	verif nomor satu jam tidak ada muncul kode ya tidak jelas

b. Case Folding

Tahap *Case folding* mengubah seluruh teks menjadi huruf kecil agar representasi kata menjadi konsisten selama proses analisis [16].

Tabel 3. Hasil Case Folding

Cleaning	Case Folding
Akses layanan cepat sangat memuaskan	akses layanan cepat sangat memuaskan
verif nomor satu jam tidak ada muncul kode ya tidak jelas	verif nomor satu jam tidak ada muncul kode ya tidak jelas

c. Normalisasi Kata

Normalisasi dilakukan untuk membuat kata yang tidak baku, singkatan, dan slang menjadi bentuk baku sesuai kaidah bahasa Indonesia [17].

Tabel 4. Normalisasi Kata

Case Folding	Normalisasi Kata
percuma punya bpjs mau daftar berobat aplikasinya susah ga jelas tolong di perbaiki dong percuma kita bayar iuran tiap bulan mau berobat susah pak knp hp saya tdk bisa dwnlod mobile jkn ya dan saya mau berobat pun susah jd solusinya gmna ya	percuma punya bpjs mau daftar berobat aplikasinya susah tidak jelas tolong di perbaiki dong percuma kita bayar iuran tiap bulan mau berobat susah pak kenapa hp saya tidak bisa dwnlod mobile jkn ya dan saya mau berobat pun susah jadi solusinya bagaimana ya

d. Tokenizing

Tokenizing membagi teks menjadi token berupa kata-kata yang seterusnya digunakan sebagai fitur pada proses klasifikasi.

Tabel 5. Hasil Tokenizing

Case Folding	Tokenizing
akses layanan cepat sangat memuaskan verif nomor satu jam tidak ada muncul kode ya tidak jelas	[akses, layanan, cepat, sangat, memuaskan] [verif, nomor, satu, jam, gak, ada, muncul, kode, nya, gak, jelas]

e. Stopword Removal

berfokus pada eliminasi kata kata yang frekuensinya tinggi tetapi minim kontribusi terhadap pemahaman isi teks.

Tabel 6. Hasil Stopword Removal

Tokenizing	Stopword Removal
[akses, layanan, cepat, sangat, memuaskan] [verif, nomor, satu, jam, gak, ada, muncul, kode, nya, gak, jelas]	[akses, layanan, cepat, memuaskan] ['verif', 'nomor', 'jam', 'gak', 'muncul', 'kode', 'nya', 'gak']

f. Stemming

setiap kata diubah ke bentuk dasar melalui penghilangan imbuhan agar kata kata dengan makna serupa dapat dikenali sebagai representasi yang sama oleh model. Proses ini memanfaatkan Pustaka sastrawi khusus untuk Bahasa Indonesia.

Tabel 7. Hasil Stemming

Stopword Removal	Stemming
[akses, layanan, cepat, memuaskan] ['verif', 'nomor', 'jam', 'gak', 'muncul', 'kode', 'nya', 'gak']	akses layan cepat muas verif nomor jam gak muncul kode nya gak

4. Pelabelan Data

Penentuan label sentimen menjadi salah satu tahapan utama dalam analisis sentimen karena hasil pelabelan digunakan sebagai acuan bagi model *machine learning* dalam proses pembelajaran. Dalam penelitian ini, proses tersebut dilakukan dengan membandingkan dua pendekatan, yaitu *rating-based labeling* dan *lexicon-based labeling*, untuk mengetahui pengaruhnya terhadap performa algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan ulasan Mobile JKN. Seluruh data yang telah melewati proses preprocessing kemudian dikelompokkan ke dalam tiga kelas sentimen, yaitu positif, netral, dan negatif.

Tabel 8. Hasil Klasifikasi Sentimen

	Positif	Negatif	Netral
Jumlah	9367	6977	17973
Persentase	27,4%	20,3%	52,3%

Sentimen terbagi menjadi tiga kategori dari hasil pelabelan terhadap 34.317 data, yaitu positif, negatif, dan netral. Sentimen netral mendominasi dengan persentase 52,3%, diikuti sentimen positif sebesar 27,4% dan

sentimen negatif sebesar 20,3%. Distribusi tersebut menunjukkan bahwa sebagian besar ulasan bersifat netral, sedangkan sisanya terbagi antara sentimen positif dan negatif. dan negatif.

a. *Lexicon InSet*

Metode *Lexicon InSet* digunakan untuk menentukan polaritas sentiment dengan memanfaatkan kamus leksikon yang berisi daftar kata beserta nilai bobot sentimennya [18], Dalam penelitian ini, *Lexicon InSet* diterapkan setelah tahap *preprocessing*. Setiap kata pada ulasan dicocokkan dengan kamus sentimen yang memiliki nilai polaritas positif maupun negatif. Seluruh skor kata kemudian dijumlahkan untuk memperoleh skor sentimen akhir. Jika skor lebih besar dari nol (>0), ulasan dikategorikan sebagai sentimen positif, jika lebih kecil dari nol (<0) sebagai sentimen negatif, dan jika sama dengan nol ($=0$) sebagai sentimen netral.

Berdasarkan proses pelabelan menggunakan pendekatan *lexicon-based* dalam penelitian ini, diperoleh distribusi sentimen yang terdiri dari 9367 pada ulasan sentimen positif, 6977 pada ulasan dengan sentimen negatif, serta 17973 pada ulasan dengan sentimen netral.

Tabel 9. Hasil Lexicon InSet

Review Text	InSetPos	InSetNeg	Sentimen
selalu saja logout heran gue supaya apa sih begitu terus disaat urgent mau digunakan logout sendiri dikira kita selalu ingat apa kata sandi nya bikin kebijakan itu yg mudah jangan mempersulit rakyat rasa ingin berkata kasar diperbaiki klo ada yg komplain jangan kyk org dongo itu bagian ITE nya nyebelin	11	-60	Negatif
mantap sangat memudahkan pendaftaran online nya	12	-6	Positif
gmana bih ngatasin offline	0	0	Netral

b. *Rating-Based*

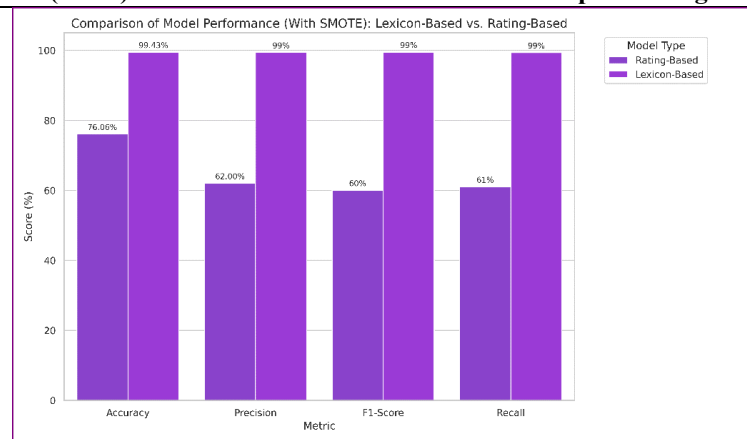
Pendekatan *Rating-Based* merupakan metode pelabelan sentimen yang memanfaatkan nilai rating pengguna sebagai indikator polaritas sentimen suatu ulasan. Pada penelitian ini, label sentimen ditentukan berdasarkan skala rating Google Play Store, di mana rating 5-4 merepresentasikan sentimen positif, rating 1-2 merepresentasikan sentimen negatif, sedangkan rating 3 dikategorikan sebagai sentimen netral [19].

Tabel 10. Hasil Rating-Based

Review Text	Rating	Sentimen
selalu saja logout heran gue supaya apa sih begitu terus disaat urgent mau digunakan logout sendiri dikira kita selalu ingat apa kata sandi nya bikin kebijakan itu yg mudah jangan mempersulit rakyat rasa ingin berkata kasar diperbaiki klo ada yg komplain jangan kyk org dongo itu bagian ITE nya nyebelin	1	Negatif
mantap sangat memudahkan pendaftaran online nya	5	Positif
gmana bih ngatasin offline	3	Netral

c. *Perbandingan Metode Lexicon InSet dan Rating-Based*

Perbandingan metode *lexicon-based* dan *rating-based* dilakukan untuk menganalisis pengaruh pendekatan pelabelan terhadap performa klasifikasi sentimen menggunakan SVM. *Lexicon-Based* menentukan label berdasarkan polaritas kata menggunakan kamus InSet, sedangkan *rating-based* memanfaatkan skor bintang yang diberikan pengguna. Perbedaan mekanisme pelabelan tersebut memengaruhi konsistensi label yang dihasilkan, sehingga berdampak pada proses pembelajaran model dan performa klasifikasi. Evaluasi terhadap kedua pendekatan disajikan pada bagian hasil dan pembahasan.



Gambar 5. Perbandingan Rating-based dan Lexicon-based

d. Validasi Ground Truth

Untuk mengevaluasi kualitas pelabelan otomatis, dilakukan validasi menggunakan 300 ulasan yang dipilih secara acak dari dataset hasil *preprocessing*. Seluruh ulasan dianotasi secara independen oleh dua anotator untuk menghasilkan label acuan (*ground truth*). Selanjutnya, hasil pelabelan menggunakan metode *lexicon-based* dan *rating-based* dibandingkan dengan label hasil anotasi manual menggunakan confusion matrix serta metrik *accuracy*, *precision*, *recall*, dan *F1-score*.

Tabel 11. Label Ground Truth

Label Ground Truth	Jumlah
Negatif	100
Netral	30
Positif	170
Total	300

5. SMOTE

Synthetic Minority Over-sampling Technique (SMOTE) merupakan metode *oversampling* yang digunakan untuk mengatasi ketidakseimbangan kelas dengan membangkitkan sampel sintetis pada kelas minoritas melalui interpolasi data [20]. SMOTE diterapkan hanya pada data pelatihan, sedangkan data pengujian tetap menggunakan distribusi asli agar evaluasi model tetap objektif. Pendekatan ini membantu model mempelajari karakteristik kelas minoritas tanpa memengaruhi hasil pengujian.

6. Modeling

Model klasifikasi sentimen dibangun menggunakan algoritma *Support Vector Machine* (SVM) berdasarkan data hasil *preprocessing* [21]. Dataset dibagi menjadi 80% data pelatihan dan 20% data pengujian melalui *train_test_split* dengan *random_state*=42 dan *stratify* untuk menjaga keseimbangan proporsi setiap kelas. Representasi fitur dilakukan menggunakan *CountVectorizer* berbasis *bag-of-words* (*ngram_range*=(1,1), *min_df*=1, *max_features*=None), sedangkan model dikembangkan menggunakan SVC dengan *kernel linear* dan parameter *C*=1.0. Metode SMOTE diterapkan hanya pada data pelatihan untuk mengurangi ketidakseimbangan kelas, sementara data pengujian dibiarkan dalam distribusi aslinya agar evaluasi tetap objektif. Kinerja model selanjutnya diukur menggunakan *confusion matrix*, *accuracy*, *precision*, *recall*, dan *F1-score*. Penelitian ini tidak menerapkan *hyperparameter tuning* maupun *cross-validation*, sehingga seluruh pengujian menggunakan parameter bawaan yang telah ditetapkan.

7. Evaluation Model

Untuk mengevaluasi performa dari model dalam klasifikasi sentimen menggunakan Data Uji masing-masing, kemudian membandingkan kinerjanya menggunakan instrumen *Confusion Matrix* [22], Karena penelitian ini menggunakan tiga kelas sentimen (negatif, netral, dan positif), maka perhitungan

akurasi dilakukan berdasarkan *confusion matrix multiclass*, yaitu dengan membandingkan jumlah seluruh prediksi yang benar pada setiap kelas terhadap jumlah seluruh data uji. Kriteria evaluasi yang dibandingkan meliputi:

- a. *Accuracy*, Rasio ketepatan prediksi keseluruhan model terhadap data aktual.

$$Accuracy = \frac{\sum_{i=1}^k TP_i}{N} \quad (1)$$

di mana:

- TP negatif= Jumlah prediksi benar pada kelas negatif.
- TP netral= Jumlah prediksi benar pada kelas netral.
- TP positif= Jumlah prediksi benar pada kelas positif.
- N= Jumlah seluruh data uji.

- b. *Precision*, Proporsi prediksi yang akurat dibandingkan dengan total ulasan yang dianalisis oleh model.

$$Precision_i = \frac{TP_i}{TP_i + FP_i} \quad (2)$$

- c. *Recall*, Kemampuan model dalam menemukan kembali seluruh ulasan yang sebenarnya.

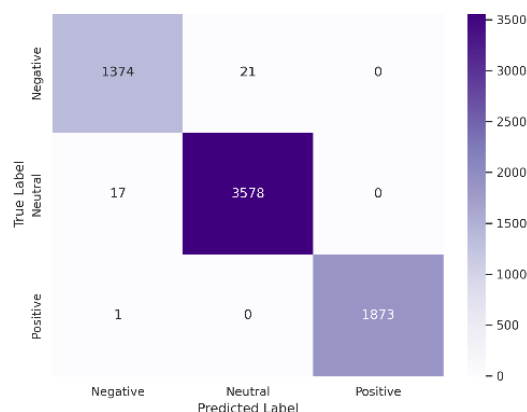
$$Recall_i = \frac{TP_i}{TP_i + FN_i} \quad (3)$$

- d. *F1-Score*, dihitung sebagai rata-rata harmonik antara nilai presisi dan recall.

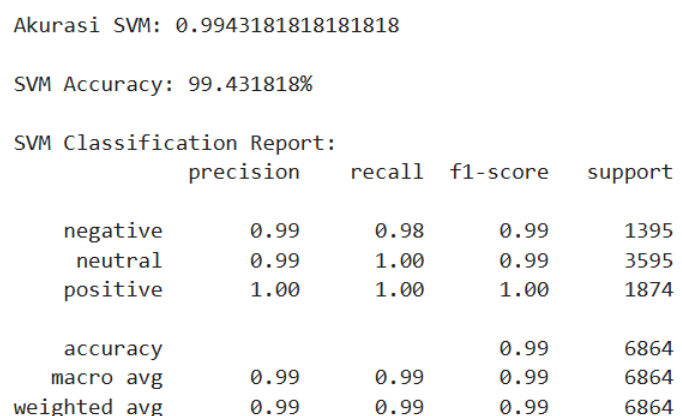
$$F1_i = \frac{2 \times Precision_i \times Recall_i}{Precision_i + Recall_i} \quad (4)$$

Evaluasi pada penelitian ini menggunakan pendekatan *macro average* karena melibatkan klasifikasi multikelas. Prosesnya diawali dengan menghitung *precision*, *recall*, dan *F1-score* untuk setiap kelas sentimen, kemudian nilai tersebut dirata-ratakan. Pendekatan ini memastikan bahwa setiap kelas memiliki bobot evaluasi yang sama, sehingga hasil pengukuran tidak bias terhadap kelas yang memiliki jumlah sampel lebih banyak [23].

Hasil klasifikasi divisualisasikan menggunakan Matplotlib dan *WordCloud* untuk menampilkan distribusi sentimen, perbandingan jumlah kelas prediksi, Termasuk Kata-kata yang paling banyak muncul dalam ulasan pengguna Mobile JKN juga termasuk dalam analisis ini. Selanjutnya, hasil prediksi digunakan untuk membentuk *confusion matrix* sebagai alat evaluasi dalam mengukur performa model melalui nilai akurasi, presisi, dan recall. Terlihat pada gambar 6 dan gambar 7 merupakan hasil evaluasi dalam bentuk *confusion matrix*.



Gambar 6. Confusion Matrix Lexicon



Gambar 7. Evaluasi Model SVM Lexicon



Gambar 8. Confusion Matrix Rating-Based

SVM Accuracy: 76.06%

SVM Classification Report (Rating Based):

	precision	recall	f1-score	support
negative	0.86	0.76	0.81	3325
neutral	0.08	0.26	0.12	325
positive	0.92	0.81	0.86	3214
accuracy			0.76	6864
macro avg	0.62	0.61	0.60	6864
weighted avg	0.85	0.76	0.80	6864

Gambar 9. Evaluasi Model SVM Rating-Based

Dari *confusion matrix lexicon* tersebut program berhasil menampilkan hasil jumlah *accuracy* 0.99.43% dan 77.06%, contoh untuk hasil perhitungan manual *accuracy* dapat dihitung manual dengan rumus $Accuracy = TP / \text{Jumlah Data}$

Tabel 12. Hasil Perhitungan Manual Lexicon

Accuracy			
	Negatif	Netral	Positif
TP	1374	3578	1873
Total	$(1374 + 3578 + 1873) / 6864 = 0,9943$		

Tabel 13. Hasil Perhitungan Manual Rating-Based

Accuracy			
	Negatif	Netral	Positif
TP	2532	86	2603
Total	$(2532 + 86 + 2603) / 6864 = 0,7606$		

Perhitungan manual terhadap confusion matrix menunjukkan bahwa metode *lexicon* pada tabel 12 mencapai *accuracy*, *precision*, dan *recall* sebesar 99.43%. Sementara itu, metode *rating-based* pada tabel 13 menghasilkan *accuracy* 77%, *precision* 62%, dan *recall* 61%. Hasil ini mengindikasikan perbedaan performa kedua metode dalam melakukan klasifikasi sentimen.

3. Hasil dan Pembahasan

Berdasarkan hasil penelitian, tampak perbedaan performa yang menonjol antara *lexicon-based labeling* dan *rating-based labeling*. Berdasarkan hasil evaluasi pada *confusion matrix*, pendekatan *lexicon-based* menghasilkan nilai *accuracy*, *precision*, dan *recall* sebesar 99%, sedangkan pendekatan *rating-based* hanya mencapai *accuracy* sebesar 76.06%, dengan *precision* sebesar 62% dan *recall* sebesar 61%. Jika dikaitkan dengan penelitian sebelumnya [8] menunjukkan bahwa kombinasi metode pelabelan *Inset Lexicon* dan algoritma SVM lebih efektif dalam mengklasifikasikan sentimen ulasan pengguna serta memberikan pemahaman yang lebih akurat mengenai persepsi pengguna terhadap aplikasi GoBiz, dibandingkan *rating-based*, hasil penelitian ini bukan hanya menguatkan temuan tersebut, tetapi juga menunjukkan peningkatan performa yang lebih tinggi pada skala dataset yang lebih besar dan menggunakan teknik penyeimbangan data (SMOTE).

Temuan ini menunjukkan bahwa pada subset validasi, pelabelan *rating-based* memiliki tingkat kesesuaian yang lebih tinggi terhadap anotasi manual dibandingkan pelabelan *lexicon-based*. *Rating-Based* memperoleh *accuracy* sebesar 85,67%, sedangkan *lexicon-based* memperoleh *accuracy* sebesar 46,67%. Temuan ini menunjukkan bahwa rating pengguna pada subset validasi lebih mendekati hasil anotasi manual. Namun demikian, model SVM yang menggunakan label *lexicon-based* tetap menghasilkan performa klasifikasi yang lebih tinggi karena pola pelabelannya lebih konsisten dan lebih mudah dipelajari oleh model berbasis *CountVectorizer*.

Perbedaan performa tersebut diduga dipengaruhi oleh karakteristik masing-masing metode pelabelan. *Lexicon-based* membentuk label berdasarkan polaritas kata pada teks, sedangkan *rating-based* menggunakan nilai bintang yang tidak selalu sejalan dengan isi ulasan. Ketidaksesuaian tersebut dapat menghasilkan label yang kurang konsisten sehingga memengaruhi proses pembelajaran model.

a. Penerapan SMOTE Terhadap Performa Model

Untuk mengevaluasi dampak penerapan *Synthetic Minority Over-sampling Technique* (SMOTE) terhadap kinerja klasifikasi, penelitian ini membandingkan dua skenario pengujian, yaitu model yang dilatih tanpa SMOTE dan model yang menggunakan SMOTE pada data pelatihan. Berdasarkan hasil evaluasi, penerapan SMOTE tidak memberikan peningkatan performa pada dataset yang digunakan. Bahkan, pada pendekatan *lexicon-based labeling*, nilai akurasi mengalami sedikit penurunan, dari 100% menjadi 99,43%, sedangkan pada pendekatan *rating-based* akurasi menurun dari 86,45% menjadi 76,06%. Temuan ini menunjukkan bahwa efektivitas SMOTE bergantung pada karakteristik dataset dan distribusi label yang digunakan, sehingga penerapannya tidak selalu menghasilkan peningkatan performa klasifikasi.

Tabel 14. Perbandingan Performat Model

Metode Labeling	SMOTE	Accuracy	Precision (Macro Avg)	Recall (Macro Avg)	F1-Score (Macro Avg)
Rating-Based	Dengan SMOTE	76.06%	62%	61%	60%
Rating-Based	Tanpa SMOTE	86.45%	62%	61%	60%
Lexicon-Based	Dengan SMOTE	100%	100%	100%	100%
Lexicon-Based	Tanpa SMOTE	99.43%	99%	99%	99%

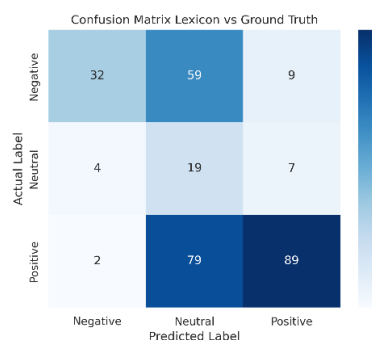
b. Validasi Ground Truth

Untuk mengevaluasi kualitas pelabelan otomatis, dilakukan validasi menggunakan 300 ulasan yang dipilih secara acak dari dataset hasil *preprocessing*. Seluruh ulasan dianotasi secara independen oleh dua anotator untuk menghasilkan label acuan (*ground truth*). Konsistensi penilaian antara kedua anotator dievaluasi menggunakan koefisien Cohen's Kappa, yang menghasilkan nilai sebesar 0,80, yang menunjukkan tingkat kesepakatan kuat (*substantial agreement*) sehingga hasil anotasi layak digunakan sebagai ground truth. Selanjutnya, hasil pelabelan menggunakan metode *lexicon-based* dan *rating-based* dibandingkan dengan label hasil anotasi manual tersebut.

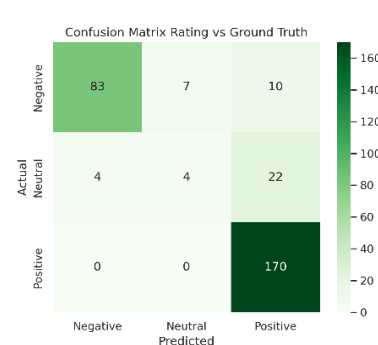
Tabel 15. Perbandingan Label Ground Truth

Metode	Accuracy	Precision	Recall	F1-Score
Lexicon-Based	46,67%	60,36%	49,23%	43,81%
Rating-Based	85,67%	71,97%	65,44%	66,56%

Validasi yang dilakukan menggunakan 300 data ground truth menunjukkan bahwa pendekatan *rating-based labeling* memiliki tingkat keselarasan yang lebih baik terhadap hasil anotasi manual. dibandingkan *lexicon-based*. Meskipun demikian, pada eksperimen klasifikasi menggunakan SVM, pendekatan *lexicon-based* menghasilkan performa model yang lebih tinggi. Temuan ini menunjukkan bahwa kualitas pelabelan dan performa model klasifikasi merupakan dua aspek yang perlu dievaluasi secara terpisah.



Gambar 10. Lexicon vs Ground Truth



Gambar 11. Rating vs Ground Truth

Confusion matrix menunjukkan bahwa metode *lexicon-based* masih mengalami kesalahan klasifikasi yang cukup besar, terutama pada kecenderungan memprediksi data ke kelas netral. Sebaliknya, *rating-based* menghasilkan jumlah kesalahan yang lebih sedikit dan memiliki tingkat kesesuaian yang lebih tinggi terhadap *ground truth*. Temuan ini sejalan dengan hasil evaluasi yang menunjukkan *accuracy* sebesar 85,67% pada *rating-based* dan 46,67% pada *lexicon-based*.

Sebaliknya, pendekatan *rating-based labeling* menunjukkan kesesuaian yang lebih baik dengan data *ground truth*. Model mampu mengklasifikasikan 83% data sentimen negatif dan seluruh data sentimen positif secara benar. Walaupun masih terdapat kesalahan pada kelas netral, jumlahnya lebih sedikit dibandingkan *Lexicon-Based Labeling*, sehingga menghasilkan *accuracy* sebesar 85,67%, lebih tinggi daripada 46,67% yang diperoleh *Lexicon-Based Labeling*.

c. Keterbatasan Penelitian

Meskipun menunjukkan hasil yang baik, penelitian ini memiliki beberapa keterbatasan. Pendekatan *Lexicon-Based* masih bergantung pada kualitas kamus sentimen sehingga belum sepenuhnya mampu menangani bahasa tidak formal, singkatan, maupun konteks tertentu dalam ulasan pengguna. Selain itu, meskipun validasi menggunakan *ground truth* telah dilakukan pada 300 ulasan, jumlah tersebut masih merupakan sebagian kecil dari keseluruhan dataset sehingga belum sepenuhnya merepresentasikan seluruh variasi sentimen pengguna. Penelitian ini juga hanya menggunakan algoritma SVM dan dataset dari satu domain aplikasi sehingga hasil penelitian masih perlu diuji pada domain dan algoritma.

4. Kesimpulan

Penelitian ini mengevaluasi dan membandingkan kinerja dua pendekatan pelabelan, yaitu *rating-based labeling* dan *lexicon-based labeling*, dalam analisis sentimen ulasan aplikasi Mobile JKN menggunakan algoritma *Support Vector Machine* (SVM). Seluruh tahapan penelitian dilaksanakan dengan mengacu pada kerangka kerja *Cross-Industry Standard Process for Data Mining* (CRISP-DM).

Analisis yang dilakukan menunjukkan bahwa penerapan algoritma *Support Vector Machine* (SVM) dengan *lexicon-based labeling* memperoleh *accuracy* sebesar 99%, sedangkan *rating-based labeling* memperoleh *accuracy* sebesar 76%. Namun, hasil validasi menggunakan 300 anotasi manual menunjukkan bahwa *rating-based* memiliki tingkat kesesuaian yang lebih tinggi terhadap *ground truth* dibandingkan *lexicon-based*. Temuan ini menunjukkan bahwa performa model klasifikasi dan kualitas pelabelan merupakan dua aspek yang berbeda.

Kontribusi penelitian ini terletak pada perbandingan dua metode pelabelan pada domain layanan kesehatan publik menggunakan dataset berskala besar serta validasi kualitas pelabelan menggunakan *ground truth*. Penelitian ini juga mengkaji dampak penerapan *Synthetic Minority Over-sampling Technique* (SMOTE) terhadap kinerja model klasifikasi. Hasil evaluasi menunjukkan bahwa keberhasilan penerapan SMOTE dipengaruhi oleh karakteristik dataset yang digunakan.

Oleh sebab itu, penerapan temuan penelitian pada domain yang berbeda perlu dikaji lebih lanjut melalui penelitian lanjutan yang melibatkan dataset yang lebih beragam dan membandingkan berbagai metode klasifikasi.

Hasil visualisasi *WordCloud* memperlihatkan bahwa kelompok sentimen positif didominasi oleh kata-kata yang mencerminkan kemudahan dalam penggunaan layanan. Sebaliknya, pada kelompok sentimen negatif, kata-kata yang paling sering muncul berkaitan dengan permasalahan *error* sistem, kendala pada proses pendaftaran, serta verifikasi wajah, yang menunjukkan aspek-aspek layanan yang masih banyak dikeluhkan oleh pengguna. Temuan ini menunjukkan bahwa BPJS Kesehatan dapat memprioritaskan peningkatan stabilitas aplikasi, penyederhanaan proses registrasi, serta optimalisasi fitur verifikasi identitas untuk meningkatkan kualitas layanan digital.

Penelitian selanjutnya disarankan menggunakan jumlah data *ground truth* yang lebih besar dengan melibatkan lebih banyak anotator, serta membandingkan metode pelabelan menggunakan algoritma lain seperti LSTM, BERT, atau model berbasis transformer agar diperoleh gambaran yang lebih komprehensif mengenai kualitas pelabelan dan performa klasifikasi.

Referensi

- [1] A. N. Arofah, V. Maiga, M. Noor, F. Endra, B. Setyawan, dan D. Azmi, "Dampak Implementasi Program JKN Terhadap Biaya Kesehatan di Fasilitas Kesehatan Tingkat Lanjutan," vol. 3, no. 2, hal. 64–72, 2022, doi: <https://doi.org/10.37148/comphijournal.v3i2.104>.
- [2] B. Kesehatan, "Mobile JKN," 2026. <https://play.google.com/store/apps/details?id=app.bpjs.mobile&hl=id> (diakses 8 Maret 2026).
- [3] M. S. Assyifa dan H. Pratiwi, "Sentiment Analysis of STMIK Widya Cipta Dharma Using a Lexicon-Based Approach Analisis Sentimen Terhadap STMIK Widya Cipta Dharma Menggunakan Pendekatan Lexicon," hal. 1–7, 2025, [Daring]. Tersedia pada: <http://repository.wicida.ac.id/id/eprint/6082>
- [4] R. R. Suryono, "Perbandingan kinerja model support vector machine dan naïve bayes untuk analisis sentimen super app polri 1) 1,2,3)," vol. 11, no. 1, hal. 1962–1976, 2026, doi: <https://doi.org/10.36341/rabit.v11i1.7582>.
- [5] W. T. Kusuma dan L. C. Kurniawan, "Sentiment Analysis User Reviews of Acupuncture Applications on Play Store Using Naïve Bayes Analisis Sentimen Ulasan Pengguna Aplikasi Akupunktur di Play Store Menggunakan Naïve Bayes," vol. 6, no. April, hal. 632–639, 2026, doi: <https://doi.org/10.57152/malcom.v6i2.2620>.
- [6] D. Analisis, S. Terhadap, dan P. Publik, "Komparasi model svm, naïve bayes, dan random forest dalam analisis sentimen terhadap percakapan publik tentang Pertamina," vol. 11, no. 1, hal. 92–105, 2026, doi: <https://doi.org/10.29100/jipi.v11i1.7699>.
- [7] T. Setyani, K. Sari, H. N. Heidy, dan R. R. Suryono, "Sentiment Analysis of Jamsostek Mobile Application Reviews on Google Play Store Using Support Vector Machine and Naive Bayes Algorithms Analisis Sentimen Pengguna Aplikasi Jamsostek Mobile Berdasarkan Ulasan Google Play Store Menggunakan Algoritma Support," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 6, no. January, hal. 373–384, 2026, doi: <https://doi.org/10.57152/malcom.v6i1.2526>.
- [8] H. Firda *et al.*, "Perbandingan Pelabelan Rating - based dan Inset Lexicon - based dalam Analisis Sentimen Menggunakan SVM (Studi Kasus : Ulasan Aplikasi GoBiz di Google Play Store) Comparison of Rating - based and Inset Lexicon - based Labeling in Sentiment Analysis usin," *Sist. J. Sist. Inf.*, vol. 14, hal. 516–528, 2025, doi: <https://doi.org/10.32520/stmsi.v14i2.4795>.
- [9] L. E. Naufal, N. Surojudin, dan I. Afriantoro, "Analisis Sentimen Aplikasi Mobile JKN di Google Play Store Menggunakan Algoritma Naive Bayes," vol. 6, no. 4, hal. 1999–2009, 2025, doi: <https://doi.org/10.47065/josh.v6i4.7846>.
- [10] Y. Wang, W. Liao, H. Shen, Z. Jiang, dan J. Zhou, "Some notes on the basic concepts of support vector machines," *J. Comput. Sci.*, vol. 82, no. July, hal. 102390, 2024, doi: <https://doi.org/10.1016/j.jocs.2024.102390>.
- [11] M. Ali dan Y. Fissaha, "Propose Adjustable Support Vector Machine Approach for Classifying Imbalanced Work Travel Mode Choice Data," *Transp. Res. Interdiscip. Perspect.*, vol. 35, no. September 2024, 2026, doi: <https://doi.org/10.1016/j.trip.2025.101786>.
- [12] Y. A. Singgalen, "Analisis Sentimen Wisatawan terhadap Kualitas Layanan Hotel dan Resort di Lombok Menggunakan SERVQUAL dan CRISP-DM," vol. 4, no. 4, hal. 1870–1882, 2023, doi: <https://doi.org/10.47065/bits.v4i4.3199>.
- [13] H. P. Jelita dan M. I. Saad, "Penerapan Algoritma Naïve Bayes Dalam Analisis Sentimen Masyarakat Terhadap STMIK Widya Cipta Dharma," vol. 6, no. 2, hal. 148–160, 2025, doi: <https://doi.org/10.47065/bit.v5i2.2029>.
- [14] I. Kurniawan *et al.*, "Implementasi Algoritma Random Forest Untuk Menentukan Implementation of Random Forest Algorithm for Determining Recipients of Raskin," vol. 10, no. 2, hal. 421–428, 2023, doi: <https://doi.org/10.25126/jtiik.202396225>.
- [15] J. J. A. Limbong, I. Sembiring, K. D. Hartomo, U. Kristen, S. Wacana, dan P. Korespondensi, "Analisis Klasifikasi Sentimen Ulasan pada E-Commerce Shopee Berbasis Word Cloud dengan Metode Naive Bayes dan K-Nearest Analysis of Review Sentiment Classification on E-Commerce Shopee Word Cloud Based with Naïve Bayes And K-Nearest Neighbor Methods," vol. 9, no. 2, hal. 347–356, 2022, doi: <https://doi.org/10.25126/jtiik.202294960>.
- [16] S. A. H. Bahtiar, "Perbandingan Naïve Bayes dan Logistic Regression dalam Sentiment Analysis pada Review Marketplace menggunakan Rating-Based Labeling," Universitas Islam Indonesia, 2023. [Daring]. Tersedia pada: <https://dspace.uui.ac.id/handle/123456789/46497>
- [17] R. C. A. S, Y. Suhari, S. Informasi, dan U. S. Unisbank, "Analisis Sentimen Ulasan Pelanggan Restoran Menggunakan Support Vector Machine," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 10, no. 1, hal. 572–576, 2026, doi: <https://doi.org/10.36040/jati.v10i1.16813>.
- [18] R. H. Muhammadi, T. G. Laksana, dan A. B. Arifa, "Combination of Support Vector Machine and Lexicon-Based Algorithm in Twitter Sentiment Analysis," *Khazanah Inform. J. Ilmu Komput. dan Inform.*, vol. 8, no.



- 1, hal. 59–71, 2022, doi: <https://doi.org/10.23917/khif.v8i1.15213>.
- [19] R. I. Alhaqq, I. M. K. Putra, dan Y. Ruldeviyani, “Analisis Sentimen terhadap Penggunaan Aplikasi MySAPK BKN di Google Play Store,” *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 11, no. 2, hal. 105–113, 2022, doi: <https://doi.org/10.22146/jnteti.v11i2.3528>.
- [20] A. Nugroho dan E. Rilvani, “Penerapan Metode Oversampling SMOTE Pada Algoritma Random Forest Untuk Prediksi Kebangkrutan Perusahaan,” *Techno.COM*, vol. 22, no. 1, hal. 207–214, 2023, doi: <https://doi.org/10.33633/tc.v22i1.7527>.
- [21] M. K. Anam, T. Arita, dan M. Bambang, “Sentiment Analysis for Online Learning using The Lexicon-Based Method and The Support Vector Machine Algorithm,” *Ilk. J. Ilm.*, vol. 15, no. 2, hal. 290–302, 2023, doi: <https://doi.org/10.33096/ilkom.v15i2.1590.290-302>.
- [22] D. Ismafillah, T. Rohana, dan Y. Cahyana, “Analisis Algoritma Pohon Keputusan untuk Memprediksi Penyakit Diabetes Menggunakan Oversampling SMOTE,” *INFOTECH J. Inform. Teknol.*, vol. 4, no. 1, hal. 27–36, 2023, doi: <https://doi.org/10.37373/infotech.v4i1.452>.
- [23] F. Syahputra, D. Yani, H. Tanjung, R. Dewi, dan H. E. Br, “Analisis Sentimen Masyarakat terhadap Isu Transformasi Digital Pemerintah di Twitter Menggunakan Naive Bayes dan Support Vector Machine,” *INFOSYS (Journal Inf. Syst.*, vol. 10, no. 2, hal. 73–85, 2026, doi: <https://doi.org/10.22303/infosys.10.2.2026.73-85>.