

ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI MOBILE JKN BERBASIS GOOGLE PLAY STORE MENGGUNAKAN NAÏVE BAYES CLASSIFIER DAN SUPPORT VECTOR MACHINE

Talitha Syahda Aguslin^{1*}, Pismia Sylvi²

^{1,2}Program Studi Statistika, Fakultas Sains dan Teknologi, Universitas Terbuka, Indonesia

**Penulis korespondensi: talithasyahda12@gmail.com*

ABSTRAK

Pada era digital sekarang perkembangan teknologi dapat membantu beberapa aktivitas masyarakat dimana salah satunya untuk mengakses informasi dan layanan kesehatan. Pengembangan inovasi di bidang kesehatan adalah aplikasi Mobile JKN yang dikembangkan oleh BPJS Kesehatan. Aplikasi ini menyediakan beberapa fitur terbaru seperti Telehealth, Bugar, New Rehab, Minum Obat, Tren Penyakit Daerah, BPJS Keliling, serta layanan *e-commerce* yang mempermudah peserta mengurus berbagai kebutuhan tanpa harus mendatangi kantor cabang. Selain itu ulasan pengguna berbasis Google Play Store masih memiliki kendala seperti tidak dapat login ke aplikasi, proses loading yang lama dan verifikasi akun gagal. Kendala tersebut menjadi bahan evaluasi bagi BPJS Kesehatan kedepannya. Penelitian ini dengan tujuan menentukan model klasifikasi paling efektif serta tingkat akurasi terbaik yang selanjutnya akan mengkategorikan menjadi sentimen positif dan negatif. Penelitian ini dengan menggunakan analisis sentimen untuk menggambarkan seperti apa pengalaman pengguna. Dua algoritma yang digunakan adalah Naïve Bayes Classifier (NBC) dan Support Vector Machine (SVM) untuk memberikan perbandingan performa dari masing-masing model dalam memproses ulasan aplikasi tersebut. Perbandingan yang akan digunakan yaitu metrik Akurasi, Presisi, dan Recall. Penelitian ini menggunakan 2.500 ulasan terbaru dengan langkah scrapping data pada halaman website aplikasi “Mobile JKN” berbasis Google Play Store. Hasil penelitian menunjukkan bahwa Algoritma Naïve Bayes Classifier memperoleh tingkat akurasi tertinggi yaitu 91%, sedangkan SVM memperoleh akurasi sebesar 84%. Hasil ini mengindikasikan bahwa NBC memiliki performa yang lebih baik dalam mengklasifikasikan sentimen serta ulasan pengguna terhadap aplikasi Mobile JKN cenderung lebih banyak bersifat negatif.

Kata kunci: Analisis Sentimen, Ulasan, Mobile JKN, Naïve Bayes Classifier, Support Vector Machine.

1 PENDAHULUAN

Kesehatan merupakan hak dasar setiap individu, selain kebutuhan lain seperti sandang, pangan, dan papan yang layak. Kesadaran masyarakat mengenai pentingnya perlindungan sosial terus meningkat seiring amanat perubahan UUD 1945 Pasal 34 ayat (2) yang menegaskan bahwa negara bertanggung jawab mengembangkan Sistem Jaminan Sosial. Lahirnya Undang-Undang Nomor 40 Tahun 2004 tentang Sistem Jaminan Sosial Nasional (SJSN) menjadi bukti nyata komitmen pemerintah dan pemangku kepentingan dalam mewujudkan kesejahteraan sosial bagi seluruh rakyat. Melalui SJSN, pemerintah berupaya memberikan jaminan agar masyarakat dapat memenuhi kebutuhan hidup yang layak (Dharma, 2022). Program Jaminan Kesehatan Nasional

yang dijalankan oleh BPJS Kesehatan sendiri telah diimplementasikan selama lebih dari sembilan tahun (Kedeputan Bidang Manajemen Klaim dan Utilisasi, 2023).

Pada era digital dan teknologi yang canggih bisa membantu aktivitas masyarakat contohnya penyebaran informasi dan penyediaan layanan kesehatan. Salah satunya inovasi dari teknologi yaitu aplikasi Mobile JKN oleh BPJS Kesehatan. Aplikasi ini bertujuan untuk memfasilitasi peserta dalam mengakses layanan mulai dari pengecekan status kepesertaan, pencarian fasilitas kesehatan tingkat pertama dan lanjut serta pendaftaran peserta secara daring. Google Play Store adalah platform tempat pengguna dapat mengunduh berbagai aplikasi termasuk Mobile JKN. Ulasan yang ada pada platform tersebut menggambarkan pengalaman langsung yang dirasakan oleh pengguna. Saat ini Mobile JKN memiliki lebih dari 905 ribu ulasan. Rating yang diberikan oleh pengguna tidak hanya rating tinggi saja tetapi juga ada rating rendah. Perbedaan ini menunjukkan adanya kesenjangan antara harapan dan pengalaman pengguna yang perlu dianalisis lebih dalam. Salah satu metode yang dapat digunakan untuk memahami tingkat kepuasan pengguna adalah analisis sentimen. Teknik ini memanfaatkan *Natural Language Processing* (NLP) dan *Machine Learning* (ML) untuk mengidentifikasi kecenderungan emosi dalam sebuah teks, apakah bersifat positif, negatif, atau netral (Purnamasari, 2023).

Sejumlah penelitian terdahulu telah memanfaatkan algoritma Naïve Bayes Classifier (NBC) maupun metode lainnya. Penelitian yang dilakukan Yunizar (2023) terkait sentimen pengguna Twitter terhadap Mobile JKN menggunakan 1001 tweet dan menghasilkan akurasi sebesar 69,65%. Syafrizal (2024), yang meneliti ulasan aplikasi PLN Mobile dengan metode NBC dan KNN, menemukan bahwa NBC memiliki performa lebih baik dengan akurasi 77,69%. Sementara itu, Sugihartono (2024) menerapkan metode optimasi Particle Swarm Optimization (PSO) untuk meningkatkan parameter pada model SVM dalam analisis ulasan Mobile JKN, sehingga akurasi meningkat dari 81% menjadi 85%. Penelitian Sriani (2024) yang membandingkan NBC dan C4.5 menggunakan 790 data ulasan menyimpulkan bahwa C4.5 memiliki kinerja klasifikasi yang lebih unggul dibanding NBC. Penelitian lain oleh Hakim (2025), yang menggunakan 100.000 ulasan Mobile JKN dengan metode Stochastic Gradient Descent (SGD), menunjukkan bahwa metode tersebut efektif dan efisien dalam menangani data berskala besar dengan akurasi mencapai 91%. Selain itu, penelitian Rohman (2021) menggunakan metode Maximum Entropy yang dikombinasikan dengan seleksi fitur Gini Index Text, dan menghasilkan akurasi tertinggi sebesar 85,36% pada threshold 80%. Sementara penelitian Hendra (2021) mengenai analisis sentimen aplikasi Halodoc menggunakan NBC menunjukkan akurasi sebesar 81,68%. Sementara itu, penelitian yang dilakukan oleh Nuris (2021) mengenai penerapan Particle Swarm Optimization (PSO) pada analisis sentimen ulasan aplikasi Halodoc dengan algoritma Naïve Bayes Classifier menunjukkan adanya peningkatan performa model. Tingkat Akurasi sebelum menggunakan AUC sebesar 88,50% dengan nilai AUC 0,535 dan kemudian meningkat menjadi 90,50% dengan nilai AUC 0,525 setelah penerapan PSO. Hal ini menghasilkan penggunaan teknik seleksi fitur berbasis PSO efektif dalam meningkatkan kualitas prediksi model. Penelitian lain oleh Rizaldi (2023) mengenai sentimen terhadap aplikasi JMO menggunakan metode NBC memperlihatkan bahwa mayoritas pengguna memberikan komentar bernada negatif, meskipun tetap terdapat sejumlah ulasan positif. Model NBC pada penelitian tersebut menghasilkan akurasi sebesar 95%. Selanjutnya, studi yang dilakukan Darmawan (2023) terhadap aplikasi MyPertamina menggunakan 3.948 data ulasan di Google Play Store menunjukkan bahwa mayoritas ulasan pengguna bersentimen negatif, dengan akurasi klasifikasi sebesar 91% menggunakan algoritma NBC.

Dalam analisis sentimen terdapat beberapa algoritma seperti Long Short Term Memory, K Nearest Neighbor, Lexicon Based ataupun lainnya. Dua metode yang paling banyak diterapkan untuk penelitian adalah NBC dan SVM. Naïve Bayes merupakan algoritma *machine learning* yang kuat dan efisien dalam klasifikasi serta tidak sulit diimplementasikan untuk menghasilkan prediksi dalam waktu relatif singkat. Sedangkan, Support Vector Machine merupakan metode *supervised learning* yang digunakan untuk klasifikasi maupun regresi dan memiliki keunggulan dalam pemrosesan data berukuran besar. Perbedaan penelitian ini dengan sejumlah studi sebelumnya terletak pada fokus perbandingan antara dua algoritma yaitu NBC dan SVM dalam menganalisis sentimen ulasan pengguna Mobile JKN. Penelitian ini bertujuan untuk menentukan model klasifikasi serta tingkat akurasi terbaik dalam mengolah ulasan pengguna aplikasi Mobile JKN berbasis Google Play Store dengan menggunakan algoritma Naïve Bayes Classifier dan Support Vector Machine. Selain itu, hasil penelitian ini diharapkan dapat menjadi bahan evaluasi BPJS Kesehatan untuk mengetahui ulasan pengguna secara objektif sehingga dapat digunakan dalam memajukan kualitas layanan informasi digital kesehatan di masa depan.

2 METODE

2.1 Jenis Penelitian

Jenis Penelitian ini merupakan penelitian terapan, yaitu penelitian yang fokus pada penerapan, pengujian, dan evaluasi suatu teori untuk menyelesaikan permasalahan praktis. Dalam hal ini, peneliti menerapkan algoritma NBC dan SVM.

2.2 Sumber Data

rating pengguna aplikasi Mobile JKN yang terdapat pada Google Playstore, dengan total sebanyak 2.500 ulasan. Data dikumpulkan dalam rentang waktu September 2020 hingga November 2025 menggunakan Google Colab.

2.3 Metode Penelitian

Adapun langkah-langkah analisis sentimen yang dilakukan adalah sebagai berikut:

a) Kajian literatur

Kajian literatur yang dilakukan penulis mencakup buku-buku, jurnal, serta bahan bacaan lainnya yang memiliki kaitan langsung dengan penelitian ini.

b) Import Data

pengambilan data dilakukan dengan cara mengambil ulasan data bertipe teks.

c) Pelabelan Manual

memberi label pada semua ulasan menjadi dua kelas yaitu positif dan negatif dalam pilihan angka 1 (satu) sampai dengan 5 (lima), dan pada rating ini dengan ketentuan rating 5 (lima) dan 4 (empat) diberikan label positif sedangkan 2 (dua) dan 1 (satu) diberikan label negatif.

d) Text Pre-processing

Text pre-processing adalah langkah awal yang digunakan untuk memformat data teks sehingga dapat diproses dalam proses data mining. Adapun langkah-langkah text pre-processing yang dilakukan sebagai berikut:

1. Cleaning

Proses membersihkan data teks dari karakter yang tidak relevan atau inkonsisten, seperti menghapus karakter khusus yaitu tanda baca, simbol, angka serta mengubah format penulisan menjadi huruf kecil.

2. Case Folding
Proses menyeragamkan semua huruf dalam dokumen seperti huruf kapital akan diubah menjadi huruf kecil.
 3. Stopword Removal
Proses menghilangkan kata tidak penting (stopword) dari teks karena tidak mewakili makna utama dokumen. Proses stopwords removal akan memproses data dari daftar stopwords. Jika terdapat kata tidak penting dalam daftar stopwords, maka kata tersebut akan dihapus sehingga kata yang tersisa hanya kata yang menandakan isi dari suatu dokumen.
 4. Tokenizing
Proses memisahkan keseluruhan teks menjadi unit-unit kecil yang disebut token yakni berupa kalimat atau kata.
 5. Stemming
Proses mengubah kata menjadi kata dasar dengan tujuan mengurangi jumlah indeks yang berbeda dari suatu dokumen.
- e) Data Training
Data training bertujuan untuk mengembangkan model.
- f) Data Testing
Data testing bertujuan untuk mengukur performansi model.
- g) Vektorisasi Data
Proses vektorisasi merupakan proses pemvektoran kata berdasarkan jumlah munculnya kata pada attribute text. Penelitian ini penulis menggunakan Pembobotan Kata/ Ekstraksi TF IDF dengan tujuan untuk mengetahui seberapa penting suatu kata terhadap Kumpulan ulasan berdasarkan frekuensi kemunculannya. Semakin sering suatu istilah muncul, maka semakin besar nilai TF-nya. Sementara Inverse Document Frequency (IDF) adalah ukuran kelangkaan suatu istilah dalam sebuah kumpulan dokumen. Semakin jarang suatu istilah muncul, maka akan semakin besar nilai IDF-nya Rumus dar TF-IDF adalah:
- $$TF = \frac{\text{Frekuensi term dalam satu dokumen}}{\text{total kata dalam satu dokumen}}$$
- $$IDF = \log \frac{\text{Total dokumen} + 1}{\text{Frekuensi dokumen mengandung term}}$$
- h) K-Fold Cross Validation
Merupakan suatu metode untuk memprediksi tingkat kesalahan dalam teknik klasifikasi.
- i) Naïve Bayes Classifier
NBC merupakan algoritma *machine learning* yang efektif dan sering digunakan dalam tugas klasifikasi. Keunggulan utamanya terletak pada kemudahan untuk mengimplementasikan smodel prediktif dalam waktu singkat. Sifat yang sederhana memungkinkan Naïve Bayes untuk memprediksi kelas sampel data dengan cepat dan menunjukkan performa yang baik pada skenario prediksi dengan banyak kategori (*multi-kelas*). Berikut rumus NBC:

$$P(A|B) = \frac{P(B|A) \times P(A)}{P(B)}$$

- j) Support Vector Machine
Algoritma *machine learning* ini adalah jenis *Supervised Learning* yang dapat melakukan prediksi kelas setelah melalui tahap pelatihan.
- k) Evaluasi Model
Kinerja model bertujuan untuk mengukur data menggunakan set data pengujian dan metrik evaluasi umum (akurasi, presisi, recall, dan F1-score). Tujuannya adalah untuk menilai kemampuan model secara optimal dalam mengklasifikasikan sentimen dari ulasan pengguna.

3 HASIL DAN PEMBAHASAN

Data ulasan dikumpulkan dengan mengimpor data menggunakan bahasa pemrograman Python di Google Colab. Proses ini mencakup ulasan yang berasal dari lima tahun terakhir (September 2020 hingga November 2025) dan menghasilkan total 2500 data ulasan. Berikut data dapat dilihat pada Tabel 1.

Tabel 1. Data Scrapping

Content	Score	Label
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif
makin kesini aplikasi jkn payah makin gila, no nik sdh benar koq bisanya katanya salah.	1	Negatif

Data yang telah terkumpul kemudian dilakukan pelabelan manual dan preprocessing langkah pertamanya adalah cleaning data. Berikut hasil cleaning pada Tabel 2.

Tabel 2. Data Cleaning

Content	Score	Label
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif
makin kesini aplikasi jkn payah makin gila, no nik sdh benar koq bisanya katanya salah.	1	Negatif

Proses berikutnya adalah Case Folding dalam python menggunakan modul lower(). Berikut hasilnya pada Tabel 3.

Tabel 3. Case Folding

Content	Score	Label	Text clean
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif	oksudah bisahanya saja harus reset password
makin kesini aplikasi jkn payah makin	1	Negatif	makin kesini aplikasi jkn payah

gila, no nik sdh benar koq bisanya katanya salah.

makin gila no nik sdh benar koq bisanya katanya salah

Selanjutnya, adalah Stopword Removal. Proses penghapusan kata seperti kata hubung ataupun kata depan. Berikut hasil pada Tabel 4.

Tabel 4. Stop Removal

content	Score	Label	Text Clean	Text StopWord
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif	oksudah bisahanya saja harus reset password	oksudah bisahanya reset password
makin kesini aplikasi jkn payah makin gila, no nik sdh benar koq bisanya katanya salah.	1	Negatif	makin kesini aplikasi jkn payah makin gila no nik sdh benar koq bisanya katanya salah	kesini aplikasi jkn payah gila no nik sdh koq bisanya salah

Proses berikutnya adalah tokenizing, merubah kalimat yang terdapat pada attribute text menjadi potongan kata-kata. Pada NTLK, tokenisasi dapat di eksekusi dengan menggunakan fungsi word_tokenize. Hasil dari tahapan tokenizing dapat diamati pada Tabel 5.

Tabel 5. Tokenizing

Content	Score	Label	Text Clean	Text StopWord	Text Tokens
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif	oksudah bisahanya saja harus reset password	oksudah bisahanya reset password	['oksudah', 'bisahanya', 'reset', 'password']
makin kesini aplikasi jkn payah makin gila, no nik sdh benar koq bisanya katanya salah.	1	Negatif	makin kesini aplikasi jkn payah makin gila no nik sdh benar koq bisanya katanya salah	kesini aplikasi jkn payah gila no nik sdh koq bisanya salah	['kesini', 'aplikasi', 'jkn', 'payah', 'gila', 'no', 'nik', 'sdh', 'koq', 'bisanya', 'salah']

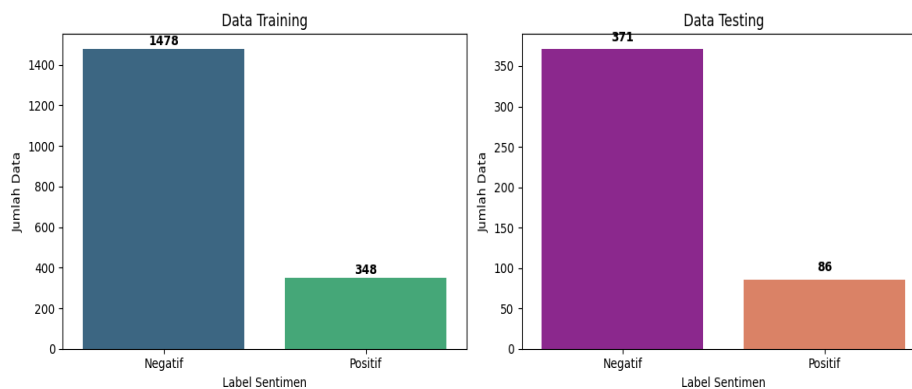
Stemming merupakan proses terakhir dari preprocessing. Pada algoritma proses stemming di eksekusi menggunakan library Sastrawi. Hasil dari tahapan pre-processing pada Tabel 6.

Tabel 6. Stemming

Content	Score	Label	Text Clean	Text StopWord	Text Tokens	Text Steamindo
Ok,,,sudah bisa,,,hanya saja harus reset password,,,	4	Positif	oksudah bisahanya saja harus reset password	oksudah bisahanya reset password	['oksudah', 'bisahanya', 'reset', 'password']	oksudah bisahanya reset password
makin kesini aplikasi jkn payah makin gila, no nik sdh benar koq bisanya katanya salah.	1	Negatif	makin kesini aplikasi jkn payah makin gila no nik sdh benar koq bisanya katanya salah	kesini aplikasi jkn payah gila no nik sdh koq bisanya salah	['kesini', 'aplikasi', 'jkn', 'payah', 'gila', 'no', 'nik', 'sdh', 'koq', 'bisanya', 'salah']	kesini aplikasi jkn payah gila no nik sdh koq bisa salah

Tahapan selanjutnya adalah Pra-Processing yaitu pemisahan data training dan data testing ulasan pengguna aplikasi Mobile JKN. Tahap ini membagi data training sebesar 80% dan data testing 20%. Berikut hasilnya pada Gambar 1.

Data Training dan Data Testing



Gambar 1. Data Training dan Data Testing

Dalam analisis sentimen, data dapat divisualisasikan dalam bentuk kumpulan potongan kata yang merepresentasikan frekuensi kemunculan kata pada ulasan, yang biasanya ditampilkan melalui sebuah gambar bernama *word cloud*. Tampilan visualisasi data melalui *word cloud* bertujuan untuk memudahkan pemahaman sentimen dimana menunjukkan kata-kata yang paling sering terlihat pada sentimen positif maupun negatif. Pada sentimen positif terdapat tiga kata dengan frekuensi tertinggi yaitu “aplikasi”, “mudah”, dan “bantu”. Kemunculan kata-kata tersebut menginformasikan bahwa pengguna merasa aplikasi Mobile JKN memberikan

Vol. 3 No. 1 (2026)

Word Cloud Sentimen POSITIF



Gambar 2. Wordcloud Positif

Word Cloud Sentimen NEGATIF



Gambar 3. Wordcloud Negatif

Sementara itu, pada sentimen negatif terdapat tiga kata dengan frekuensi kemunculan tertinggi, yaitu “verifikasi”, “sulit”, dan “error”. Kata-kata yang muncul dari hasil keluhan pengguna yang merasa terganggu karena tidak dapat mengakses layanan digital secara optimal sehingga pengguna berharap agar BPJS Kesehatan melaksanakan evaluasi terhadap fitur-fitur yang masih belum membantu pengguna serta dapat meningkatkan kualitas layanan demi menghasilkan kepuasan pengguna. Visualisasi *word cloud* untuk sentimen negatif pada aplikasi Mobile JKN ditampilkan pada gambar tersebut.

Tahap selanjutnya adalah perhitungan vektorisasi TF-IDF pada kedua algoritma NBC dan SVM dengan tujuan memperoleh nilai atau bobot pada setiap kata. Proses perhitungan TF-IDF dilakukan menggunakan Python dengan bantuan pustaka Scikit-Learn. Berikut hasilnya pada Gambar 4.


```

MultinomialNB Accuracy: 0.912472647702407
MultinomialNB Precision: 0.912718204488778
MultinomialNB Recall: 0.9865229110512129
MultinomialNB f1_score: 0.9481865284974094
confusion_matrix:
[[366  5]
 [ 35 51]]
=====

```

	precision	recall	f1-score	support
Negatif	0.91	0.99	0.95	371
Positif	0.91	0.59	0.72	86
accuracy			0.91	457
macro avg	0.91	0.79	0.83	457
weighted avg	0.91	0.91	0.90	457

Gambar 4. Hasil Perhitungan Vektorisasi TF-IDF dengan NBC

```

... SVM Accuracy: 0.8358862144420132
SVM Precision: 0.9134078212290503
SVM Recall: 0.8814016172506739
SVM f1-score: 0.897119341563786

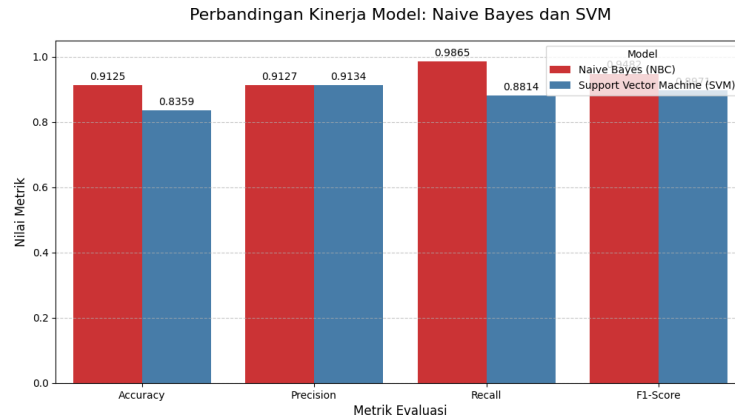
Confusion matrix:
[[327 44]
 [ 31 55]]
=====

```

	precision	recall	f1-score	support
Negatif	0.91	0.88	0.90	371
Positif	0.56	0.64	0.59	86
accuracy			0.84	457
macro avg	0.73	0.76	0.75	457
weighted avg	0.85	0.84	0.84	457

Gambar 5. Hasil Perhitungan Vektorisasi TF-IDF dengan SVM

Hasil pemodelan dari kedua algoritma tersebut kemudian akan dibandingkan untuk mendapatkan model terbaik berdasarkan nilai akurasi, precision, recall, dan F1-Score. Hasil perbandingan model NBC dan SVM dapat dilihat pada Gambar 6.



Gambar 6. Hasil Perbandingan NBC dan SVM

Berdasarkan hasil perbandingan kedua model, diperoleh bahwa algoritma NBC menunjukkan performa yang lebih unggul dibandingkan SVM. NBC berhasil mencapai akurasi sebesar 91,25%, recall 98,65%, precision 91,27%, serta F1-Score 94,82%. Sementara itu, algoritma SVM hanya memperoleh akurasi 83,59%, recall 88,14%, precision 91,34%, dan F1-Score 89,71%. Penelitian ini diperoleh hasil algoritma NBC dapat menghasilkan klasifikasi lebih baik dibandingkan algoritma SVM. NBC mampu memberikan kinerja yang efektif walaupun algoritma sederhana. Sedangkan, algoritma SVM dapat mengolah data teks tapi hasilnya tidak lebih baik dari NBC. Selain itu, proses data yang dilakukan SVM membutuhkan serangkaian pengujian tambahan seperti transformasi data ke format yang lebih besar dan penggunaan kernel tertentu sehingga perlu waktu komputasi yang lebih lama dibandingkan NBC agar dapat menghasilkan performa yang lebih optimal.

4 KESIMPULAN

Proses penelitian ini terdiri dari sebelas langkah yaitu pengumpulan data ulasan pengguna, pemberian label pada data ulasan, preprocessing, pembagian dataset, vektorisasi untuk model NBC dan SVM, serta evaluasi performa masing-masing model tersebut. Penelitian ini menghasilkan model NBC dengan kinerja paling baik dalam mengklasifikasikan sentimen ulasan. Hasil ini diperoleh informasi ulasan pengguna terhadap aplikasi Mobile JKN cenderung bersentimen negatif. Hasil penelitian ini diharapkan dapat menjadi bahan evaluasi bagi pengembang Mobile JKN untuk meningkatkan kualitas serta pengalaman pengguna dalam layanan digital kesehatan dengan memperhatikan ulasan pengguna. Pengembang aplikasi dapat membuat fitur lebih relevan dengan kebutuhan pengguna. Oleh karena itu, penelitian ini memberikan arahan bagi pengembangan aplikasi kesehatan yang lebih optimal ke depannya. Berdasarkan hasil penelitian, maka saran yang dapat diberikan yaitu penelitian selanjutnya diharapkan dapat memperluas cakupan dengan menganalisis ulasan berbahasa asing seperti Bahasa Inggris atau bahasa lainnya. Selain itu, penelitian selanjutnya juga dapat membandingkan beberapa aplikasi layanan kesehatan digital lainnya untuk mendapatkan gambaran dalam sisi lain.

UCAPAN TERIMA KASIH

Ucapan terimakasih kepada pihak-pihak yang telah membantu dalam penelitian. Semoga hasil penelitian bermanfaat dan menjadi acuan penelitian selanjutnya.

DAFTAR PUSTAKA

- Darmawan, G., Alam, S., & Sulisty, I. M. (2023). Analisis sentimen berdasarkan ulasan pengguna aplikasi mypertamina pada google play store menggunakan metode naïve bayes. *Jurnal Ilmiah Teknik dan Ilmu Komputer*, 100-108.
- Hendra, A., & Fitriyani. (2021). Analisis Sentimen Review Halodoc Menggunakan Nai'Ve Bayes Classifier. *JISKa*, 78-89.
- Kedeputan Bidang Manajemen Klaim dan Utilasi. (2023). *Profil Pelayanan Kesehatan BPJS Kesehatan Tahun 2014-2022*. Jakarta: BPJS Kesehatan.
- Nuris, N., Yulia, E. R., & Solecha, K. (2021). Implementasi Particle Swarm Optimization (PSO) Pada Analysis Sentiment Review Aplikasi Halodoc Menggunakan Algoritma Naïve Bayes. *Jurnal Teknologi Informasi*, 17-23.
- Purnamasari, D., Aji, B. A., & A, W. D. (2023). *Pengantar Metode Analisis Sentimen*. Depok: Gunadarma.
- Rizaldi, S. A., Alam, S., & Kurniawan, I. (2023). Analisis sentimen pengguna aplikasi jmo (jamsostek mobile) pada google play store menggunakan metode naïve bayes. *Jurnal Ilmiah Teknik dan Ilmu Komputer*, 109-117.
- Rohman, M. M., Indriati, & Adinugroho, S. (2021). Analisis Sentimen pada Ulasan Aplikasi Mobile JKN Menggunakan Metode Maximum Entropy dan Seleksi Fitur Gini Index Text . *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 2646-2654.
- Sriani, Suhardi, & Yardha, L. S. (2024). Analisis sentimen pengguna aplikasi mobile jkn menggunakan algoritma naïve bayes. *Journal of Science and Social Research*, 555-563.
- Sugihartono, T., & Chrisna, R. R. (2024). Penerapan Metode Support Vector Machine dalam Classifikasi Ulasan . *SKANIKA: Sistem Komputer dan Teknik Informatika*, 144-153.
- Syafrizal, Afdal, M., & Novita, R. (2024). Analisis Sentimen Ulasan Aplikasi PLN Mobile Menggunakan Algoritma Naïve Bayes Classifier dan K-Nearest Neighbor. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 10-19.
- Yunizar, Z., Rusnani, & Ardian, Z. (2023). Analisis sentimen pada twitter terhadap aplikasi mobile jkn. *Journal of Informatics and Computer Science*, 103-111