

# Analisis Sentimen Terhadap Ulasan Aplikasi Mobile JKN Menggunakan Metode Machine Learning Logistic Regression, SVM, dan CSVM

Moch.Firman Fernando<sup>1</sup>, Davin Anezta Ahmad<sup>2\*</sup>, Nugroho Fajar Rachmanto<sup>3</sup>,  
Shindi Shella May Wara<sup>4</sup>, Kartika Maulida Hindrayani<sup>5</sup>

<sup>12345</sup>Program Studi Sains Data, Fakultas Ilmu Komputer, Universitas  
Pembangunan Nasional Veteran Jawa Timur, Surabaya, 60294, Indonesia

\* Corresponding author, email: 23083010097@student.upnjatim.ac.id

## Abstract

*One of the digital-based public service innovations in the health sector is the Mobile JKN application developed by BPJS Kesehatan. This application allows people to get health services more easily, effectively, and integrated. The purpose of this study is to evaluate user perceptions of the Mobile JKN application through collecting reviews from the Google Play Store. The collected data was analyzed using TF-IDF text mining technique and Chi-Square feature selection. Furthermore, logistic regression, support vector machine (SVM), and clustered SVM (CSVM) algorithms were used to perform sentiment classification. Comments were categorized into three categories: positive, neutral, and negative. The evaluation results show that CSVM has an accuracy value of 93%, precision of 94%, recall of 84%, and F1 value of 89%. Although features such as online registration and digital cards received positive feedback, sentiment analysis showed that most reviews were negative, especially regarding technical issues. The results show that ML-based algorithms can be effectively used to assess how people perceive digital services. These results can be used as a basis for BPJS Kesehatan to improve and develop new services.*

**Keywords:** Mobile JKN, Health facilities, Web scraping, Sentiment analysis, User reviews, Service evaluat.

## Abstrak

Salah satu inovasi layanan publik berbasis digital di bidang kesehatan adalah aplikasi Mobile JKN yang dikembangkan oleh BPJS Kesehatan. Aplikasi ini memungkinkan masyarakat untuk mendapatkan layanan kesehatan secara lebih mudah, efektif, dan terintegrasi. Tujuan penelitian ini adalah untuk mengevaluasi persepsi pengguna terhadap aplikasi Mobile JKN melalui pengumpulan ulasan dari Google Play Store. Data yang dikumpulkan dianalisis menggunakan teknik text mining TF-IDF dan seleksi fitur Chi-Square. Selanjutnya, algoritma regresi logistik, support vector machine (SVM), dan clustered SVM (CSVM) digunakan untuk melakukan klasifikasi sentimen. Komentar dikategorikan ke dalam tiga kategori: positif, netral, dan negatif. Hasil evaluasi menunjukkan bahwa CSVM memiliki skor akurasi 93%, precision 94%, recall 84%, dan skor F1 89%. Sementara fitur seperti pendaftaran online dan kartu digital mendapat tanggapan positif, analisis sentimen menunjukkan bahwa mayoritas ulasan bersifat negatif, terutama terkait masalah teknis. Hasilnya menunjukkan bahwa algoritma berbasis ML dapat digunakan secara efektif untuk menilai bagaimana masyarakat melihat layanan digital. Hasil ini dapat digunakan sebagai dasar untuk BPJS Kesehatan memperbaiki dan mengembangkan layanan baru.

**Kata Kunci:** Mobile JKN, Fasilitas Kesehatan, Web Scraping, Analisis sentimen, Ulasan Pengguna, Evaluasi Layanan

## 1. Pendahuluan

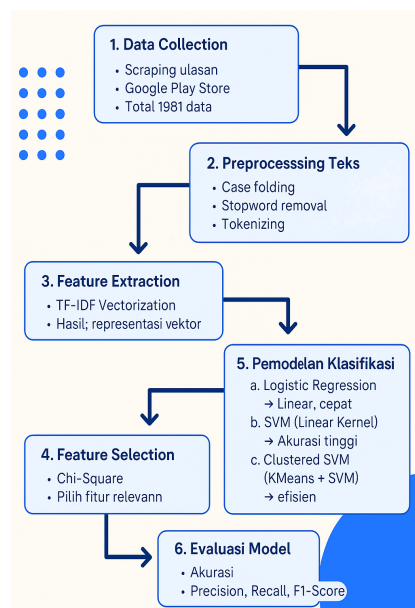
Teknologi informasi telah menjadi elemen penting dalam upaya meningkatkan kualitas layanan publik, dengan memberikan manfaat signifikan dalam efisiensi dan kemudahan akses [1]. Teknologi informasi memiliki peran yang penting dalam bidang kesehatan di era digital, termasuk dalam operasional rumah sakit. Kemajuan teknologi informasi dan komunikasi memberikan pengaruh besar terhadap peradaban modern, sehingga memungkinkan berbagai pekerjaan dalam organisasi dapat diselesaikan secara cepat, tepat, dan efisien [2]. Dalam bidang kesehatan, teknologi informasi berkontribusi besar dalam mempercepat akses layanan dan meningkatkan efisiensi. Aplikasi *mobile* seperti *Mobile JKN* merupakan inovasi penting yang memungkinkan masyarakat memperoleh layanan kesehatan secara lebih mudah, cepat, dan dapat diakses kapan saja dan di mana saja [3]. Dengan *rating* aplikasi yang berada pada angka 3.5, ulasan yang diterima oleh aplikasi *Mobile JKN Faskes* mencerminkan pandangan yang beragam, baik positif maupun negatif. Ulasan ini memberikan wawasan yang berharga bagi pengembang untuk mengevaluasi dan meningkatkan kualitas aplikasi. Penelitian ini bertujuan untuk menganalisis sentimen dalam ulasan pengguna aplikasi *Mobile JKN Faskes*, dengan fokus pada identifikasi polaritas ulasan-positif, netral, atau negative-berdasarkan *rating* yang diberikan oleh pengguna. Analisis sentimen merupakan proses otomatis untuk mengekstrak, memahami, dan mengolah data teks yang tidak terstruktur, dengan tujuan untuk memperoleh informasi mengenai sentimen yang terkandung dalam sebuah kalimat pendapat atau opini [4]. Analisis sentimen memiliki tujuan untuk menentukan apakah suatu teks memiliki sentimen positif, negatif, atau netral [5]. Dalam penelitian ini, analisis sentimen dilakukan terhadap ulasan pengguna aplikasi *Mobile JKN Faskes* di *Google Play Store*. Teknik *web scraping* digunakan untuk mengumpulkan data ulasan, sementara *machine learning* diterapkan untuk mengklasifikasikan sentimen. Algoritma seperti *Naive Bayes*, *Support Vector Machine (SVM)*, dan *Logistic Regression* digunakan untuk mengklasifikasikan ulasan pengguna ke dalam tiga kategori sentimen: positif, netral, dan negatif.

Penelitian ini bertujuan memberikan gambaran yang lebih jelas mengenai pandangan pengguna terhadap aplikasi *Mobile JKN Faskes*, serta memberikan rekomendasi perbaikan berdasarkan hasil analisis sentimen. Memahami pendapat dan perasaan pengguna dapat membantu pengembang untuk meningkatkan aspek-aspek yang perlu diperbaiki atau diperkuat dalam aplikasi, yang pada gilirannya akan meningkatkan pengalaman pengguna dan memperbaiki *rating* aplikasi. Penelitian ini diharapkan memberikan kontribusi signifikan terhadap pemahaman penggunaan aplikasi kesehatan digital di Indonesia dan memberikan panduan untuk mengoptimalkan layanan kesehatan berbasis teknologi.

## 2. Material dan Metode

### 2.1. Proses Tahapan Penelitian

Penelitian dimulai dengan pengumpulan data melalui *web scraping* dari *Google Play Store* untuk memperoleh ulasan aplikasi *Mobile JKN Faskes*. Data ulasan yang terkumpul kemudian diproses melalui tahap *preprocessing* teks untuk membersihkan dan menyiapkan teks dengan teknik seperti *case folding*, *stopword removal*, dan *tokenization*. Tahap ekstraksi fitur dilakukan dengan menggunakan metode *TF-IDF* untuk mengubah teks menjadi representasi numerik yang dapat diproses oleh model. Seleksi fitur dilakukan menggunakan metode *Chi-Square* untuk meningkatkan efisiensi model. Pemodelan klasifikasi diterapkan dengan menggunakan tiga model, yaitu *Logistic Regression*, *SVM*, dan *Clustered SVM*, untuk mengklasifikasikan *sentimen* ulasan menjadi tiga kategori: positif, netral, dan negatif. Evaluasi model dilakukan untuk mengukur kinerja masing-masing model menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-Score*. *Flowchart* berikut memberikan gambaran yang jelas dan menyederhanakan pemahaman tentang urutan tahapan dalam penelitian.



**Gambar 1.** *Flowchart* tahapan penelitian

## 2.2. Dataset

Penelitian ini menggunakan data ulasan pengguna dari aplikasi layanan publik pemerintah, yaitu *Mobile JKN Faskes* yang diperoleh dari platform *Google Play Store*. Data dikumpulkan melalui teknik *web scraping*. *Web scraping* merupakan metode untuk mengekstraksi berbagai jenis informasi, seperti teks, tautan, media, dan dokumen, dari halaman web tertentu [6]. Proses ini dilakukan menggunakan paket *Python google-play-scraper*.

Total 1.981 komentar yang mencakup informasi berupa teks ulasan, tanggal, skor *rating* (1–5), dan nama pengguna. Untuk keperluan analisis sentimen, ulasan dikategorikan sebagai positif apabila memiliki skor  $\geq 4$ , netral jika skor = 3, dan negatif

apabila skor  $\leq 2$ .

### 2.3. Text Preprocessing

Pra-pemrosesan teks dilakukan untuk membersihkan dan menyederhanakan data teks sebelum dianalisis lebih lanjut. Tahapan ini mencakup *case folding*, tujuan dari *case folding* adalah mengubah semua kata menjadi huruf kecil bertujuan untuk memastikan bahwa seluruh data teks yang diproses berada dalam format yang seragam [7], penghapusan angka dan tanda baca, penggantian kata tidak baku dengan kata baku melalui kamus sinonim, penghapusan *stopword* yaitu penghapusan dilakukan terhadap kata-kata yang memiliki frekuensi kemunculan tinggi namun tidak memberikan kontribusi berarti terhadap makna keseluruhan kalimat atau teks [8], dan proses *stemming* untuk menemukan bentuk dasar dari sebuah kata turunan [9]. Tujuan utama dari tahap ini adalah untuk menghilangkan *noise* dan meningkatkan konsistensi data teks.

### 2.4. Word Weighting

Pada ekstraksi fitur, digunakan metode *Term Frequency–Inverse Document Frequency (TF-IDF)* untuk menilai sejauh mana pentingnya suatu kata (*term*) dalam sebuah dokumen, dibandingkan dengan koleksi dokumen yang lebih besar [10]. *TF* menghitung frekuensi kemunculan kata dalam sebuah dokumen, sedangkan *IDF* menilai seberapa unik kata tersebut dengan menghitung kebalikannya terhadap jumlah dokumen yang memuat kata tersebut. Metode *TF-IDF* menghasilkan representasi *vektor* dari dokumen teks yang mencerminkan pentingnya kata dalam konteks dokumen dan *korpus* secara keseluruhan.

$$TF - IDF(t, d) = TF(t, d) \times \log\left(\frac{N}{DF(t)}\right) \quad (1)$$

### 2.5 Feature Selection

Agar model klasifikasi bekerja lebih efektif, dilakukan seleksi fitur menggunakan metode *Chi-Square*. Metode ini mengukur tingkat ketergantungan antara fitur (kata) dengan kelas target (sentimen).

#### Hipotesis:

$H_0$  : Tidak ada hubungan antara kata dengan label sentimen.

$H_1$  : Ada hubungan antara kata dengan label sentimen.

#### Statistik uji:

$$\chi^2 = \sum_{i=1}^n \frac{(x_i - \mu)^2}{\mu} \quad (2)$$

## 2.6 Chi-Square

Uji statistik yang disebut *Chi-Square* bertujuan untuk menentukan apakah dua peristiwa, seperti munculnya kata dan kelas sentimen, berkorelasi atau tidak independen. *Chi-square* adalah alat untuk mengukur tingkat hubungan statistik antar variabel. Dalam penelitian ini Dalam penelitian ini, metode *chi-square* digunakan untuk menguji sejauh mana hubungan antara fitur-fitur dalam dataset [11] Adapun persamaannya dapat dituliskan pada persamaan berikut:

$$\chi^2 = \sum_{i=1}^n \frac{(O - E)^2}{E} \quad (3)$$

## 2.7. Classification with Logistic Regression

Penelitian ini menggunakan algoritma *regresi logistik* untuk klasifikasi sentimen sebagai metode pendamping. *Regresi logistik* merupakan teknik yang digunakan untuk menggambarkan hubungan antara satu atau lebih variabel independen dengan satu variabel dependen yang bersifat biner [12]. Metode klasifikasi yang didasarkan pada fungsi *logistik (sigmoid)* dikenal sebagai regresi logistik digunakan untuk memodelkan kemungkinan bahwa suatu data akan termasuk dalam kelas tertentu. Algoritma ini menggunakan hubungan linier antara fitur dan *logit*, yang dikenal sebagai *log odds*, untuk menghitung kemungkinan bahwa suatu entri teks termasuk ke dalam kategori sentimen tertentu—positif, negatif, atau netral. Dalam penelitian ini, *regresi logistik* adalah salah satu metode analisis statistik yang digunakan untuk mensimulasikan hubungan antara variabel independen dan variabel dependen yang berskala dalam data *nominal* dan *ordinal* [13]. Model *regresi logistik* dilatih untuk mengidentifikasi pola sentimen dari data ulasan pengguna aplikasi *Mobile JKN Faskes* setelah proses *preprocessing* dan representasi fitur dilakukan menggunakan metode *TF-IDF*.

$$P(y = 1|x) = \frac{1}{1 + e^{(\beta_0 + \beta_1 x_1 + \dots + \beta_p x_p)}} \quad (4)$$

## 2.8. Classification with Support Vector Machine

Klasifikasi *sentimen* dalam penelitian ini dilakukan menggunakan algoritma *Support Vector Machine (SVM)* dengan *kernel linear*. *SVM* merupakan algoritma klasifikasi yang memanfaatkan konsep *margin* untuk memisahkan data yang bersifat positif, netral dan negatif [14] . Metode klasifikasi *SVM* mencari *hyperplane* optimal untuk memisahkan data antar kelas dengan *margin* maksimum, sehingga mampu melakukan *generalisasi* dengan baik. Data ulasan yang telah diproses diubah menjadi *vektor fitur* menggunakan metode *TF-IDF*, menghasilkan representasi berdimensi tinggi yang sesuai untuk *SVM*. Algoritma ini dipilih karena efektif menangani data teks yang bersifat *sparse* dan

berdimensi tinggi, serta dikenal memiliki performa yang baik dalam tugas klasifikasi.

$$f_{svm}(x) = \sum_{i=1}^{sv} \alpha_i y_i K(x_i, x_i) + b \quad (5)$$

dengan  $\alpha_i$  adalah vektor bobot,  $y_i$  adalah label kelas, dan  $b$  adalah bias. *SVM* digunakan karena kemampuannya dalam menangani data berdimensi tinggi dan menghasilkan performa tinggi pada masalah klasifikasi teks.

## 2.9. Clustered Support Vector Machine

Untuk meningkatkan efisiensi komputasi pada data berskala besar, digunakan pendekatan *Clustered Support Vector Machine (CSVM)*. Data pelatihan dibagi menjadi beberapa *klaster* menggunakan metode *K-Means*. *SVM* kemudian diterapkan secara terpisah pada tiap klaster untuk menangkap relasi *non-linier* dalam ruang lokal. Uji data baru akan diasosiasikan ke *klaster* terdekat berdasarkan jarak rata-rata terhadap sampel dalam setiap klaster. Jumlah *klaster* terbaik ditentukan menggunakan *indeks* Salah satu keunggulan utama dari *CSVM* adalah efisiensinya, dengan waktu pemrosesan yang dua kali lebih cepat dibandingkan *SVM*, menjadikannya pilihan yang tepat untuk memproses *Calinski-Harabasz. dataset* besar dengan lebih efisien [15].

$$d(i, k) = \sqrt{\sum_{i=1}^p |x_{ij} - x_{kj}|^2} \quad (6)$$

## 2.10 Evaluation

Evaluasi performa model dilakukan untuk memberikan gambaran kinerja model dalam mengklasifikasi

$$Akurasi = \frac{Prediksi\ benar}{Total\ data} \quad (7)$$

untuk mengukur ketepatan dan kelengkapan klasifikasi:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (8)$$

## 3. Hasil dan Diskusi

### 3.1. Sumber Data

Data yang digunakan dalam penelitian ini berasal dari ulasan pengguna terhadap aplikasi *mobile JKN Faskes* yang diperoleh dari *Google Play Store*. Pengumpulan data dilakukan dengan menggunakan teknik *web scraping* dengan *skrip Python* bernama *google-play-scraper*. Proses ini berhasil mengumpulkan 1.981 ulasan secara total, yang meliputi informasi identitas pengguna, tanggal, *rating* (1-5), dan teks ulasan. Terkait







**Tabel 1.** Simulasi dari *Feature Selection*

| no   | feature | chi2_score | p-value |
|------|---------|------------|---------|
| 1    | bagus   | 120.5330   | 6.7079  |
| 2    | mantap  | 56.5483    | 5.2566  |
| ⋮    | ⋮       | ⋮          | ⋮       |
| 1980 | ancur   | 0.3621     | 0.34383 |
| 1981 | bukan   | 0.3602     | 0.35189 |

Dari proses seleksi fitur, jumlah variabel yang dianggap memiliki pengaruh terhadap klasifikasi diperoleh sebagai berikut. Sejumlah kata (*fitur*) dengan skor *chi2\_tinggi* dan p-nilai signifikan ditemukan selama proses seleksi fitur menggunakan metode *Chi-Square*. Nilai-nilai ini menunjukkan adanya hubungan kuat antara kata-kata tersebut dengan *label sentimen* tertentu. Kata-kata seperti "bagus" dan "mantap" memiliki nilai *chi2\_score* yang tinggi, yang menunjukkan bahwa kata-kata tersebut memiliki korelasi yang sangat kuat dengan kelas sentimen positif. Sebaliknya, kata-kata seperti "ancur" dan "bukan" memiliki nilai *chi2\_score* yang rendah dan *p-value* yang tinggi, sehingga mereka dianggap tidak relevan untuk membedakan kelas sentimen. Hasilnya menunjukkan bahwa *fitur* yang memiliki kontribusi signifikan terhadap pemisahan kelas akan dipertahankan selama proses pelatihan model *klasifikasi*, sedangkan *fitur* dengan nilai *chi2\_score* rendah akan dieliminasi untuk meningkatkan efisiensi dan akurasi *model*. Agar model dapat memahami konteks sentimen pengguna terhadap aplikasi *Mobile JKN*, proses ini membantu mengekstrak informasi yang paling penting dari teks ulasan.

### 3.4. Hasil Analisis Parameter Terbaik

Hasil pengujian parameter terbaik dari ketiga metode dapat dilihat pada Tabel 2 berikut:

**Tabel 2.** Hasil Pengujian Parameter Terbaik

| Model               | Accuracy | Precision | Recall | F1-Score | Macro Avg | Weight Avg |
|---------------------|----------|-----------|--------|----------|-----------|------------|
| Clustered SVM       | 0.93     | 0.94      | 0.84   | 0.89     | 0.63      | 0.92       |
| Logistic Regression | 0.89     | 0.90      | 0.98   | 0.94     | 0.56      | 0.90       |

|        |      |      |      |      |      |      |
|--------|------|------|------|------|------|------|
| SVM    | 0.91 | 0.92 | 0.98 | 0.95 | 0.57 | 0.91 |
| Linear |      |      |      |      |      |      |

Pada Tabel 2 yang disajikan menggambarkan hasil pengujian parameter terbaik dari tiga model *klasifikasi*, yaitu *SVM Linear*, *Logistic Regression*, dan *CSVM*, dalam upaya mengklasifikasikan ulasan *Mobile JKN* ke dalam tiga kategori *label sentimen*: Negatif, Netral, dan Positif. Secara keseluruhan *Clustered SVM* menunjukkan hasil yang paling stabil dengan *accuracy* dan *F1-Score* tertinggi di antara model lainnya, namun dengan *recall* yang sedikit lebih rendah dibandingkan dengan model *SVM Linear*. *Logistic Regression* sangat baik dalam mendeteksi data positif dengan *recall* yang sangat tinggi, tetapi performanya sedikit menurun pada kelas yang tidak seimbang (*macro avg* rendah). *SVM Linear* memiliki performa yang sangat baik dalam hal *precision* dan *recall*, tetapi kesulitan dalam menangani kelas yang tidak seimbang.

Hasil pengujian parameter terbaik perlabel kelas label sentimen ulasan dapat dilihat pada Tabel 3 berikut:

**Tabel 3.** Hasil Pengujian Parameter Terbaik Perkelas

| Model                  | Label   | Precision | Recall | F1-Score |
|------------------------|---------|-----------|--------|----------|
| SVM<br>Linear          | Negatif | 0.92      | 0.98   | 0.95     |
|                        | Netral  | 0.00      | 0.00   | 0.00     |
|                        | Positif | 0.85      | 0.74   | 0.80     |
| Logistic<br>Regression | Negatif | 0.90      | 0.98   | 0.94     |
|                        | Netral  | 0.00      | 0.00   | 0.00     |
|                        | Positif | 0.98      | 0.69   | 0.78     |
| CSVM                   | Negatif | 0.93      | 0.99   | 0.96     |
|                        | Netral  | 1.00      | 0.06   | 0.11     |
|                        | Positif | 0.94      | 0.84   | 0.89     |

Pada Tabel 3 menyajikan hasil pengujian parameter terbaik perkelas dari tiga *model klasifikasi*, yaitu *SVM Linear*, *Logistic Regression*, dan *CSVM*, dalam upaya mengklasifikasikan ulasan *Mobile JKN* ke dalam tiga kategori *label sentimen*: Negatif, Netral, dan Positif. Metrik yang digunakan untuk mengevaluasi kinerja masing-masing model adalah *Precision*, *Recall*, dan *F1-Score*, yang masing-masing mencerminkan ketepatan, keterbacaan, dan keseimbangan antara keduanya dalam mengidentifikasi setiap label. Model *SVM Linear* menunjukkan kinerja terbaik pada *label* Negatif, dengan nilai *precision* 0.92 dan *recall* 0.98, menghasilkan *F1-Score* 0.95. Hal ini mengindikasikan bahwa model ini cukup efektif dalam mengklasifikasikan ulasan negatif, dengan tingkat kesalahan yang relatif rendah. Namun, model ini tidak berhasil dalam mengklasifikasikan *label* Netral, dengan nilai *precision* dan *recall* 0.00, yang

menunjukkan bahwa model ini tidak mampu membedakan ulasan netral dari ulasan lainnya. Pada label Positif, *SVM Linear* menunjukkan *precision* 0.85 dan *recall* 0.74, dengan *F1-Score* 0.80. Meskipun performanya lebih rendah dibandingkan dengan label Negatif, model ini tetap menunjukkan hasil yang cukup baik dalam mengidentifikasi ulasan positif. *Logistic Regression* menunjukkan hasil yang sebanding pada label Negatif, dengan *precision* 0.90 dan *recall* 0.98, menghasilkan *F1-Score* 0.94. Model ini juga berhasil mengklasifikasikan ulasan negatif dengan akurat, meskipun kurang optimal pada label Netral, dengan *precision* dan *recall* yang keduanya 0.00. Sementara itu, untuk label Positif, *Logistic Regression* menunjukkan *precision* yang sangat tinggi (0.98), meskipun *recall* sedikit lebih rendah pada 0.69, yang menghasilkan *F1-Score* 0.78. Meskipun *recall* lebih rendah, *precision* yang tinggi menunjukkan bahwa model ini sangat selektif dalam mengklasifikasikan ulasan positif.

Hasil pengujian antar parameter terbaik terhadap model terbaik dapat dilihat pada Tabel 4 berikut:

**Tabel 4.** Hasil Pengujian Antar Tiga Model *CSVM*, *SVM*, dan *Logistic Regression*

| Model               | Akurasi | Precision | F1 Score | Waktu (s) |
|---------------------|---------|-----------|----------|-----------|
| SVM                 | 0.9068  | 0.8865    | 0.8951   | 0.16      |
| Logistic Regression | 0.8992  | 0.8808    | 0.8856   | 0.07      |
| CSVM                | 0.9339  | 0.9358    | 0.9221   | 0.17      |

Hasil pengujian tiga *model klasifikasi*, yaitu *CSVM* (*C-Support Vector Machine*), *SVM* (*Support Vector Machine*), dan *Logistic Regression*, menunjukkan perbedaan yang signifikan dalam hal kinerja dan efisiensi. Berdasarkan hasil pengujian, model *CSVM* menunjukkan performa terbaik dengan *accuracy* 0.9339, *precision* 0.9358, dan *F1-score* 0.9221, meskipun memerlukan waktu proses yang sedikit lebih lama yaitu 0.17 detik. Model ini terbukti sangat efektif dalam mengklasifikasikan data ulasan *Mobile JKN Faskes* dengan sedikit kesalahan pada kategori positif, yang menjadikannya pilihan terbaik ketika prioritas utama adalah *accuracy* dan *precision* tinggi. Sementara itu model *SVM* memberikan *accuracy* sangat baik 0.9068, *precision* 0.8865, dan *F1-score* 0.8951, meskipun sedikit lebih rendah dari *CSVM* dalam hal *accuracy* dan *precision*. Model ini membutuhkan waktu yang sedikit lebih lama daripada *Logistic Regression*, yaitu 0.16 detik, tetapi masih cukup efisien dalam menghasilkan prediksi yang akurat dan cepat. Di sisi lain, *Logistic Regression* menunjukkan hasil yang lebih efisien dalam waktu proses, yaitu hanya 0.07 detik, tetapi memiliki *accuracy* 0.8992, *precision* 0.8808, dan *F1-score* 0.8856 yang sedikit lebih rendah dibandingkan dengan kedua model sebelumnya.

#### 4. Kesimpulan

Hasil penelitian menunjukkan bahwa teknik *text mining* dan algoritma klasifikasi berbasis *machine learning* dapat digunakan untuk menganalisis *sentimen* pengguna aplikasi *Mobile JKN* tentang layanan yang diberikan oleh fasilitas kesehatan. Teknik *web scraping* dari *Google Play Store* digunakan untuk mengumpulkan data ulasan, yang kemudian diproses melalui tahap *preprocessing* teks untuk mengubah data mentah menjadi representasi numerik menggunakan metode *TF-IDF*. Dalam proses ini, dua algoritma klasifikasi, *Support Vector Machine (SVM)* dan *Clustered SVM*, digunakan untuk membagi ulasan ke dalam tiga kategori *sentimen* positif, netral, dan negatif. Hasil evaluasi menunjukkan bahwa algoritma *Clustered SVM* lebih baik daripada *SVM* konvensional, seperti yang ditunjukkan oleh perbandingan metrik evaluasi seperti *accuracy*, ketepatan, *recall*, dan skor *F1* yang lebih tinggi. Metode pengelompokan data sebelum klasifikasi digunakan oleh model *Clustered SVM*. Ini membantu mengurangi *noise* dan meningkatkan ketepatan klasifikasi pada *dataset* ulasan pengguna yang memiliki karakteristik yang berbeda. Secara keseluruhan, penelitian ini menunjukkan bahwa teknik *analisis sentimen* yang didasarkan pada pembelajaran mesin dapat digunakan untuk menilai persepsi masyarakat terhadap layanan publik yang disediakan oleh aplikasi digital. Hasil dapat digunakan oleh *BPJS Kesehatan* untuk meningkatkan layanan, terutama yang berkaitan dengan layanan yang diberikan di fasilitas kesehatan melalui aplikasi *Mobile JKN*.

#### Daftar Pustaka

- [1] Sakir, A. R. Tinjauan Literatur: Pemanfaatan Teknologi Informasi untuk Meningkatkan Mutu Pelayanan Publik. *Jurnal Administrasi Publik dan Bisnis*, 6(2), 165–171, 2024, doi: 10.36917/japabis.v6i2.170
- [2] Amelia, A., & Solikhah, M. Web-based employee attendance information system on CV. Syntax Corporation Indonesia. *JIST*, 4(12), 2436–2442, 2023, doi: 10.59141/jist.v4i12.824
- [3] Baskila, N. A., Farisni, T. N., Fitriani, F., & Jihad, F. F. Pemanfaatan Inovasi Pelayanan Kesehatan Mobile JKN pada Masyarakat di Kota Meulaboh. *Jurnal Kesehatan Tambusai*, 4(3), p. 2859, 023, doi: 10.31004/jkt.v4i3.17914
- [4] Brahimi, B., Touahria, M., & Tari, A. Improving sentiment analysis in Arabic: A combined approach. *J. King Saud Univ. – Comput. Inf. Sci.*, 33(10), 1242–1250, 2021, doi: 10.1016/j.jksuci.2019.07.011
- [5] Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. Analisis Sentimen Ulasan Aplikasi Dana dengan Metode Random Forest. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 6(9), 4305–4313, 2022.

- [6] Kusumo, S., & Somya, R. Penerapan Web Scraping Deskripsi Produk Menggunakan Selenium Python Dan Framework Laravel. *Jurnal Teknik Informatika dan Sistem Informasi*, 9(4), 3426-3435, 2022.
- [7] Permana, A. Y., & Effendi, M. M. Optimasi Stemming Porter KBBI dan Cross Validation Naïve Bayes untuk Klasifikasi Topik Soal UN Bahasa Indonesia. *J. Ilm. Komputasi*, 17(4), 2018.
- [8] Angelina, S. J., Negara, A. B. P., & Muhardi, H. Analisis Pengaruh Penerapan Stopword Removal Pada Performa Klasifikasi Sentimen Tweet Bahasa Indonesia. *JUARA (Jurnal Aplikasi dan Riset Informatika)*, 2(1), pp. 165–173, 2023, doi: 10.26418/juara.v2i1.69680.
- [9] Guterres, A., Gunawan, G., & Santoso, J. Stemming Bahasa Tetun Menggunakan Pendekatan Rule Based. *Teknika*, 8(2), pp. 142-150, 2019, doi: 10.34148/teknika.v8i2.224
- [10] Septiani, D., & Isabela, I. Analisis Term Frequency Inverse Document Frequency (TF-IDF) dalam Temu Kembali Informasi pada Dokumen Teks. *SINTESIA: Jurnal Sistem dan Teknologi Informasi Indonesia*, 1(2), pp. 81-88, 2022.
- [11] Sami'un, C., Sugiharto, A., & Jie, F. Chi Square Feature Selection for Improving Sentiment Analysis of News Data Privacy Treats. *Journal of Theoretical and Applied Information Technology*, 102(18), 6601-6609, 2024.
- [12] Situngkir, R. H., & Sembiring, P. Analisis Regresi Logistik Untuk Menentukan Faktor-Faktor Yang Mempengaruhi Kesejahteraan Masyarakat Kabupaten/Kota Di Pulau Nias. *FARABI: Jurnal Matematika dan Pendidikan Matematika*, 6(1), pp. 25-31, 2023.
- [13] Muflihah, I. Z. Analisis Financial Distress Perusahaan Manufaktur Di Indonesia dengan Regresi Logistik. *Maj. Ekon.*, vol. 22(2), 254–269, 2017.
- [14] Aisah, I. S., Irawan, B., & Suprpti, T. Algoritma Support Vector Machine (SVM) untuk Analisis Sentimen Ulasan Aplikasi Al-Qur'an Digital. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(6), pp. 3759-3765, 2023.
- [15] Wara, S. S. M., Adziima, A. F., Pratama, A. R., & Nasrudin, M. Predictive analysis of government application comment on Playstore with clustered support vector machine. in *Proc. 2024 IEEE 10th Information Technology International Seminar (ITIS)*, Surabaya, Indonesia, Nov. 6–8, 2024, doi: 10.1109/ITIS64716.2024.10845453