

Analisis Sentimen Ulasan Aplikasi *MyRepublic* di *Google Play Store* Menggunakan Algoritma *Support Vector Machine (SVM)*

Leo Fernando¹, Maisyaroh^{2,*}, Martua Hami Siregar²

¹ Teknologi Informasi; Universitas Bina Sarana Informatika; Jl. SMA Kapin No.292A, Jakarta Timur 13450; e-mail: leo.fernando@gmail.com

² Teknologi Informasi, Universitas Bina Sarana Informatika; Jl. Kramat Raya no.98 Jakarta Pusat 10450; e-mail: maysaroh.msy@bsi.ac.id, martua.mhe@bsi.ac.id

* Korespondensi: e-mail: maysaroh.msy@bsi.ac.id

Diterima: 10 Juli 2026; Review: 26 Juli 2026; Disetujui: 03 Agustus 2026

Cara sitasi: Fernando L, Maisyaroh, Siregar MH. 2026. Analisis Sentimen Ulasan Aplikasi *MyRepublic* di *Google Play Store* Menggunakan Algoritma *Support Vector Machine (SVM)*. *Bina Insani ICT Journal*. Vol. 13 (1), 107 – 116.

Abstrak: Perkembangan teknologi mendorong penyedia layanan internet seperti *MyRepublic* untuk terus mengevaluasi layanannya melalui ulasan pengguna di *Google Play Store*. Namun, ribuan ulasan yang tidak terstruktur dan didominasi oleh keluhan pelanggan menyulitkan pihak pengembang untuk melakukan analisis secara manual dan efisien. Penelitian ini bertujuan untuk menerapkan teknik *Text Mining* menggunakan algoritma *Support Vector Machine (SVM)* guna mengklasifikasikan sentimen pengguna ke dalam kategori positif dan negatif secara otomatis sebagai bahan referensi evaluasi layanan. Metodologi yang digunakan mengacu pada kerangka kerja *Knowledge Discovery in Database (KDD)* yang meliputi ekstraksi 1.000 data ulasan, *preprocessing*, pelabelan *hybrid* menggunakan *InSet Lexicon* dan validasi manual, transformasi data dengan TF-IDF, serta pemodelan SVM *Linear Kernel* dengan rasio pembagian data latih dan data uji sebesar 80:20. Hasil pengujian *Confusion Matrix* menunjukkan bahwa model algoritma SVM mencapai tingkat akurasi sebesar 94,50%. Namun, tingginya angka akurasi ini dipengaruhi oleh ketidakseimbangan data (*data imbalance*) yang ekstrem dengan proporsi 94% kelas negatif. Hal ini dibuktikan dari evaluasi metrik yang menunjukkan tingkat *Recall* kelas negatif mencapai 100%, sementara kelas positif hanya sebesar 8,33%. Dengan demikian, model ini memiliki sensitivitas yang sangat tinggi dalam mendeteksi mayoritas keluhan pelanggan secara presisi, namun belum dapat diklaim handal secara keseluruhan karena masih kesulitan dalam mengklasifikasikan sentimen positif akibat ketimpangan data pelatihannya.

Kata kunci: Analisis Sentimen, *Machine Learning*, *MyRepublic*, *Support Vector Machine*, *Text Mining*

Abstract: The advancement of technology encourages internet service providers like *MyRepublic* to continuously evaluate their services through user reviews on the *Google Play Store*. However, thousands of unstructured reviews dominated by customer complaints make it difficult for developers to conduct manual and efficient analysis. This study aims to apply *Text Mining* techniques using the *Support Vector Machine (SVM)* algorithm to automatically classify user sentiment into positive and negative categories as a reference for service evaluation. The methodology applied refers to the *Knowledge Discovery in Database (KDD)* framework, which includes the extraction of 1,000 review data, *preprocessing*, hybrid labeling using the *InSet lexicon* and manual validation, data transformation with TF-IDF, and SVM *Linear Kernel* modeling with an 80:20 ratio of training and testing data split. The *Confusion Matrix* evaluation results show that the SVM algorithm model achieved an accuracy rate of 94.50%. However, this high accuracy is heavily influenced by extreme data imbalance, with 94% of the data belonging

to the negative class. This is evidenced by the metric evaluation showing a negative class Recall rate reaching 100%, while the positive class is only 8.33%. Therefore, while the model is highly sensitive in precisely detecting the majority of customer complaints, it cannot be claimed as highly reliable overall due to its difficulty in accurately classifying positive sentiments caused by the imbalanced training data.

Keywords: Sentiment Analysis, Machine Learning, MyRepublic, Support Vector Machine, Text Mining

1. Pendahuluan

Perkembangan teknologi mendorong penyedia layanan internet untuk terus meningkatkan kualitas layanan yang mereka tawarkan, termasuk pengembangan aplikasi seluler. Aplikasi MyRepublic dirancang untuk mempermudah pelanggan dalam memantau layanan mereka, melakukan pembayaran, dan melaporkan gangguan layanan kepada perusahaan. Sebagai sarana komunikasi utama dengan pengguna, umpan balik pelanggan di *Google Play Store* menjadi sumber informasi penting bagi perusahaan untuk mengevaluasi seberapa baik sistem dan layanannya berfungsi.

Meskipun MyRepublic menyediakan layanan internet cepat, banyak pengguna telah mengungkapkan ketidakpuasan mereka melalui berbagai platform daring. Sebuah penelitian yang menganalisis pendapat pelanggan mengenai berbagai penyedia layanan internet di Indonesia menemukan bahwa sekitar 80,9% umpan balik bersifat negatif, dengan faktor-faktor utama berupa koneksi yang tidak stabil, cakupan layanan yang terbatas, dan dukungan pelanggan yang kurang memadai [1]. Tren keluhan serupa sering kali terlihat terutama di bagian ulasan aplikasi MyRepublic di *Google Play Store*, yang merupakan saluran utama bagi pelanggan untuk melaporkan masalah teknis maupun administratif.

Ulasan yang terkumpul di *Google Play Store* bisa mencapai ribuan seiring dengan meningkatnya jumlah unduhan; umumnya ditampilkan sebagai komentar yang tidak teratur. Hal ini menimbulkan masalah bagi pengelola aplikasi karena sulit dan memakan waktu untuk menyaring atau mengidentifikasi pola apa pun di antara ulasan tersebut, sehingga proses ini rentan terhadap bias pribadi [2]. Dengan demikian, diperlukan metode analisis sentimen berbasis penambangan teks (*Text Mining*) untuk mengkategorikan data tidak terstruktur ini secara efisien dan efektif [3].

Penelitian ini mengusulkan penggunaan algoritma Support Vector Machine (SVM) untuk mengelompokkan ulasan pengguna aplikasi MyRepublic berdasarkan sentimen. Berdasarkan beberapa penelitian yang menunjukkan keunggulannya sebagai alat kategorisasi teks yang lebih akurat dibandingkan metode lain seperti Naive Bayes dan K-Nearest Neighbor, SVM menjadi pilihan [4], [2]. Beberapa studi kasus ulasan aplikasi berikut ini semakin memperkuat keunggulan SVM dalam berbagai penelitian sebelumnya: Dengan pembagian data pelatihan dan pengujian sebesar 90:10, penerapan SVM pada ulasan aplikasi Mola menghasilkan akurasi sebesar 84,61% [5]; ulasan aplikasi LinkedIn diklasifikasikan dengan akurasi 82% melalui optimasi *grid search* [6]; ulasan aplikasi Gojek pada 8.000 data mencapai akurasi 86% [7]; serta studi komparatif pada ulasan Jamsostek Mobile membuktikan SVM (82,25%) lebih unggul dibandingkan *Naive Bayes* (78,50%) [8].

Ringkasan beberapa penelitian terdahulu yang relevan dengan penerapan algoritma SVM pada analisis sentimen ulasan aplikasi disajikan pada Tabel 1.

Tabel 1. Ringkasan Penelitian Terkait

Peneliti (Tahun)	Objek Penelitian	Jumlah Data	Akurasi
Hendriyanto dkk. (2022) [5]	Aplikasi Mola	520 ulasan	84,61% (rasio 90:10)
Maarif & Setiyawati (2024) [6]	Aplikasi LinkedIn	3.000 ulasan	82,00%
Nugroho & Sudarno (2025) [7]	Aplikasi Gojek	8.000 ulasan	86,00%
Setyani dkk. (2026) [8]	Aplikasi Jamsostek Mobile	6000 ulasan	82,25% (SVM) vs 78,50% (Naive Bayes)
Andiansyah & Wahyudin. (2023) [2]	Aplikasi Home Credit	2845 ulasan	84% (rasio 80:20)

Sumber : Hasil Penelitian (2026)

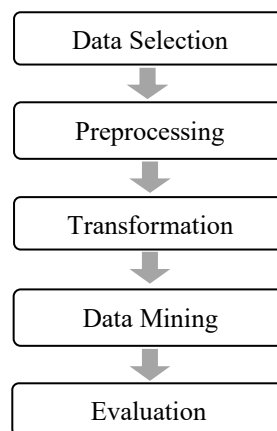
Penelitian dilakukan dengan mengikuti kerangka kerja Knowledge Discovery in Databases (KDD) [9], yang mencakup lima fase *Data Selection*, *Preprocessing*, *Transformation*, *Data Mining*, dan *Evaluation* untuk mengkategorikan sentimen pengguna terkait aplikasi MyRepublic yang terdapat di *Google Play Store* ke dalam dua kelompok: positif dan negatif, dengan memanfaatkan algoritma SVM. Penelitian ini bertujuan untuk mengungkap pola dalam umpan balik pengguna dan menilai keefektifan model SVM melalui analisis *Confusion Matrix*, dengan asumsi bahwa model tersebut dapat mencapai akurasi klasifikasi minimal 75% atau lebih.

Nilai dari penelitian ini terletak pada penerapan kerangka kerja KDD secara menyeluruh, yang mencakup pengumpulan data awal melalui *web scraping*, penerapan pendekatan campuran dalam pelabelan sentimen yang menggabungkan metode leksikon dan pemeriksaan manual, serta pelaksanaan pemodelan dan evaluasi SVM. Hal ini dieksplorasi dalam studi kasus yang berfokus pada aplikasi MyRepublic, sebuah topik yang belum mendapat perhatian signifikan dalam penelitian yang ada. Temuan dari penelitian ini diharapkan dapat memberikan landasan analitis berbasis data bagi para pengembang yang bertujuan untuk mengevaluasi dan meningkatkan pengalaman pengguna (UX) serta kualitas layanan secara keseluruhan dari aplikasi MyRepublic, sekaligus berkontribusi pada pengetahuan akademis mengenai penggunaan algoritma SVM pada teks dalam bahasa Indonesia.

2. Metode Penelitian

2.1 Tahapan Penelitian (Kerangka KDD)

Penelitian dilaksanakan dengan memanfaatkan kerangka kerja *Knowledge Discovery in Database* (KDD) [9] untuk memastikan bahwa setiap tahap pengolahan data selama penilaian evaluasi aplikasi MyRepublic dapat dilaksanakan secara terstruktur dan dapat diukur. Kerangka kerja ini dipilih karena kemampuannya yang telah teruji dalam menggali wawasan berharga dari data tidak terstruktur. Tahapan penelitian secara umum mencakup lima proses utama sebagaimana ditunjukkan pada Gambar 1.



Sumber : Hasil Penelitian (2026)

Gambar 1. Tahapan Penelitian Bagan alur KDD

1. *Data Selection*: mengumpulkan informasi umpan balik pengguna langsung dari halaman aplikasi MyRepublic Indonesia di *Google Play Store* [10] melalui metode *web scraping* dengan memanfaatkan pustaka Python bernama *google-play-scraper*.
2. *Preprocessing*: proses persiapan dan standarisasi teks mentah melibatkan upaya untuk menjadikannya konsisten dan sesuai untuk analisis lebih lanjut. Hal ini mencakup mengubah huruf besar-kecil semua karakter menjadi huruf kecil, serta menghilangkan URL, sebutan, tagar, angka, tanda baca, dan emoji dengan menggunakan ekspresi reguler. Teks tersebut kemudian dibagi menjadi token kata-kata individual. Selanjutnya, dilakukan normalisasi untuk memperbaiki istilah yang salah eja atau kata-kata non-standar dengan memanfaatkan kamus

normalisasi khusus. Setelah itu, kata-kata pengisi seperti kata penghubung yang tidak menyampaikan sentiment dihilangkan menggunakan kamus NLTK bahasa Indonesia. Terakhir, dilakukan *stemming* untuk mengubah kata-kata menjadi bentuk akarnya melalui pustaka Sastrawi, yang mengarah ke tahap awal pelabelan sentimen.

3. *Transformation*: transformasi teks yang telah disterilkan ke dalam format numerik dengan menggunakan teknik TF-IDF (*Term Frequency-Inverse Document Frequency*) [3].
4. *Data Mining*: penerapan metode *Support Vector Machine* (SVM) untuk menganalisis kumpulan data pelatihan dan memprediksi kategori sentimen pada kumpulan data uji.
5. *Evaluation*: Menganalisis model yang dikembangkan dengan bantuan *Confusion Matrix* untuk mengevaluasi tingkat ketepatannya.

2.2 Landasan Algoritma Support Vector Machine dan TF-IDF

Support Vector Machine, yang umumnya dikenal sebagai SVM, adalah metode yang digunakan dalam (*Supervised-Learning*) untuk mengidentifikasi *Hyperplane* optimal yang membedakan dua kategori data, sekaligus memperbesar jarak antara kelas-kelas yang berbeda [2], [4]. *Hyperplane* yang memisahkan data tersebut dapat direpresentasikan sebagai:

$$w \cdot x + b = 0 \quad \text{..... (1)}$$

sedangkan fungsi keputusan yang digunakan untuk mengklasifikasikan data baru dirumuskan sebagai:

$$f(x) = \text{sign}(w \cdot x + b) \quad \text{..... (2)}$$

di mana x mewakili vektor data uji, w melambangkan vektor bobot yang menentukan orientasi *Hyperplane*, b menunjukkan istilah bias, dan sign adalah fungsi yang menghasilkan nilai +1 untuk kelas positif dan -1 untuk kelas negatif [11]. Jarak pemisahan antara kedua kategori tersebut didefinisikan sebagai 2 dibagi dengan besarnya w oleh karena itu, tujuan utama pelatihan SVM adalah mengurangi nilai $\|w\|$ untuk memperoleh margin terbesar. Dalam kasus di mana data tidak dapat dipisahkan secara linier, SVM menggunakan fungsi kernel untuk mentransformasi data ke dalam ruang berdimensi lebih tinggi untuk penelitian ini, digunakan *Linear Kernel*, yang telah terbukti efektif dalam mengkategorikan teks berdimensi tinggi seperti hasil yang diberi bobot TF-IDF. [2].

Sebelum dilanjutkan ke tahap pemodelan, teks ulasan yang telah disempurnakan diubah ke dalam format numerik dengan menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) [3]. Bagian *Term Frequency* (TF) menghitung frekuensi kemunculan istilah t dalam suatu dokumen d :

$$TF(t,d) = \frac{f_{t,d}}{n_d} \quad \text{..... (3)}$$

Sementara itu, *Inverse Document Frequency* (IDF) mengukur seberapa jarang suatu kata muncul di seluruh kumpulan dokumen (*corpus*):

$$IDF(t) = \log \left(\frac{N}{df(t)} \right) \quad \text{..... (4)}$$

Nilai TF-IDF akhir diperoleh dengan mengalikan kedua elemen tersebut:

$$TF - IDF(t,d) = TF(t,d) \times IDF(t) \quad \text{..... (5)}$$

di mana N mewakili jumlah total dokumen dalam kumpulan tersebut dan $df(t)$ menunjukkan berapa banyak dokumen yang memuat istilah t . Matriks bobot TF-IDF yang dihasilkan kemudian digunakan sebagai masukan untuk algoritma SVM pada tahap penambangan data.

2.3 Sumber dan Teknik Pengumpulan Data

Data primer dikumpulkan melalui pengamatan langsung dan pengumpulan umpan balik pengguna mengenai aplikasi MyRepublic Indonesia dari Google Play Store [10] dengan menggunakan pendekatan *web scraping* yang dijalankan dengan bahasa pemrograman Python. Data sekunder diperoleh dari tinjauan pustaka yang mencakup jurnal ilmiah, buku, dan makalah akademis terkait penambangan teks, analisis sentimen, serta algoritma SVM yang telah diterbitkan dalam lima tahun terakhir.

2.4 Populasi dan Sampel

Kelompok penelitian ini terdiri dari seluruh ulasan yang ditulis dalam bahasa Indonesia yang dipublikasikan oleh pengguna di halaman aplikasi MyRepublic Indonesia di *Google Play Store*, sebuah angka yang terus meningkat seiring berjalannya waktu. Proses seleksi menggunakan metode *purposive sampling* dengan persyaratan khusus: ulasan harus ditulis dalam bahasa Indonesia, memiliki isi yang bermakna, dan berasal dari pengguna yang secara rutin menggunakan layanan MyRepublic. Berdasarkan kriteria tersebut, terkumpul total 1.000 ulasan, yang dianggap memadai dan representatif untuk tujuan melatih dan mengevaluasi model klasifikasi dalam *machine learning* [2].

2.5 Teknik Analisis Data

Dengan menggunakan bahasa pemrograman Python, analisis data dilakukan dengan memanfaatkan pustaka *pandas*, *scikit-learn*, *NLTK*, *Sastrawi*, *matplotlib/seaborn*, dan *WordCloud*. Tahapan penelitian tersebut meliputi:

- a. Pelabelan sentimen dilakukan secara *hybrid* dalam dua tahap. Tahap pertama adalah pelabelan otomatis yang menggunakan pendekatan *Lexicon-Based* dengan kamus *InSet Lexicon* (Sentimen Indonesia) [12], di mana setiap ulasan dihitung skor sentimennya berdasarkan bobot kata (skor > 0 dikategorikan positif, dan skor < 0 dikategorikan negatif). Tahap kedua adalah validasi manual oleh peneliti, yang secara khusus menargetkan ulasan dengan skor 0 (netral/ambigu), ulasan yang mengandung sarkasme, serta meninjau ulang ketepatan label otomatis guna memastikan kualitas data sebelum masuk ke tahap pemodelan.
- b. Menggunakan TF-IDF untuk ekstraksi fitur guna membuat matriks bobot kata sebagai masukan bagi model klasifikasi.
- c. Dengan menggunakan teknik *train-test split* [5], untuk membagi kumpulan data dengan perbandingan 80:20, diperoleh data uji dan data pelatihan.
- d. Pelatihan model SVM dengan *Linear Kernel* dan parameter regularisasi $C=1.0$ (nilai *default*), karena kernel ini terbukti efektif untuk klasifikasi teks berdimensi tinggi [2].
- e. Model tersebut dievaluasi menggunakan alat uji matriks kebingungan (*Confusion Matrix*) untuk mengukur *Accuracy*, *Precision*, *Recall*, dan *F1-Score* [13], [14]. Penggunaan metrik tambahan (terutama *Recall* dan *F1-Score*) ini sangat krusial untuk mengevaluasi kinerja model secara lebih komprehensif, objektif, dan adil. Mengingat dataset yang digunakan mengalami ketidakseimbangan kelas (*class imbalance*) yang ekstrem, nilai *Accuracy* saja tidak cukup untuk menjadi tolok ukur keandalan klasifikasi. Hipotesis penelitian (H1) diterima jika akurasi mencapai atau melebihi 75%.

Perangkat lunak yang digunakan dalam penelitian ini dapat dilihat pada Tabel 2 sebagai berikut:

Tabel 2. Perangkat Lunak Yang Digunakan

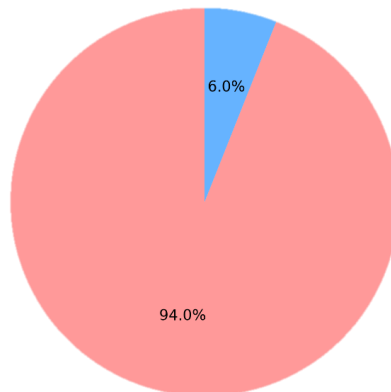
Perangkat Lunak	Fungsi
Python 3.13.5	Bahasa pemrograman utama implementasi tahapan KDD
google-play-scraper	Web scraping data ulasan dari Google Play Store
pandas	Manipulasi dan analisis data tabular
NLTK / Sastrawi	Stopword removal dan stemming teks Bahasa Indonesia
scikit-learn	Implementasi TF-IDF, SVM, dan evaluasi model
matplotlib / seaborn / WordCloud	Visualisasi data dan hasil evaluasi

Sumber : Hasil Penelitian (2026)

3. Hasil dan Pembahasan

3.1 Hasil Pengumpulan dan Praproses Data

Proses web scraping menghasilkan 1.000 catatan ulasan mentah untuk aplikasi MyRepublic dalam format CSV, yang mencakup atribut nama pengguna, peringkat, tanggal, dan teks ulasan, seperti yang ditunjukkan pada Gambar 2.



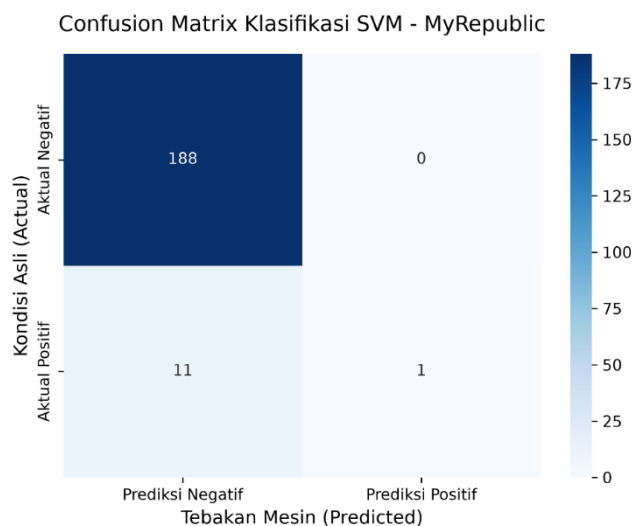
Sumber : Hasil Penelitian (2026)

Gambar 3. Distribusi Sentimen Ulasan Aplikasi MyRepublic

Distribusi data terbukti mengalami ketidakseimbangan kelas (*class imbalance*) yang sangat ekstrem. Terdapat 937 ulasan (94%) yang dikategorikan sebagai sentimen negatif, dan hanya 60 ulasan (6%) yang dikategorikan sebagai sentimen positif. Fakta tingginya ketimpangan jumlah data ini berdampak langsung pada bias tahap pemodelan, sehingga mengharuskan peneliti untuk menangani evaluasi algoritma secara lebih kritis menggunakan metrik tambahan seperti *Recall* dan *F1-Score*, bukan sekadar berpatokan pada *Accuracy* semata.

3.3 Hasil Klasifikasi dan Evaluasi Model SVM

Model SVM dengan *Linear Kernel* dilatih dan dievaluasi menggunakan data uji setelah proses ekstraksi fitur TF-IDF selesai dan data dibagi dengan perbandingan 80:20 (797 titik data pelatihan dan 200 titik data uji). Hasil pengujian ditampilkan dalam bentuk *Confusion Matrix* pada Gambar 5.



Sumber : Hasil Penelitian (2026)

Gambar 4. Confusion Matrix Klasifikasi Model SVM

Berdasarkan Gambar 5, model tersebut berhasil mengkategorikan 188 ulasan negatif (*True Negative*) dan 1 ulasan positif (*True Positive*) dari total 200 titik data uji. Sebanyak 11 ulasan yang ditandai sebagai positif salah dikategorikan sebagai negatif (*False Negative*) oleh

karena itu, tidak ada satu pun ulasan negatif yang salah dikategorikan sebagai positif (*False Positive* = 0).

Perhitungan akurasi keseluruhan model dirumuskan sebagai berikut:

$$\text{Accuracy} = \frac{TP + TN}{\text{Total}} = \frac{1+188}{200} = \frac{189}{200} = 0.945 = 94,5 \%$$

Ringkasan metrik evaluasi pendukung (*Precision*, *Recall*, dan *F1-Score*) untuk masing-masing kelas sentimen ditunjukkan pada Tabel 3. Analisis dan interpretasi lebih mendalam mengenai tingginya ketimpangan nilai *Recall* antara kelas negatif dan positif ini akan diuraikan lebih lanjut pada bagian Pembahasan.

Tabel 3. Ringkasan Metrik Evaluasi Model SVM

Kelas	Precision	Recall	F1-Score
Positif	100,00%	8,33%	15,30%
Negatif	94,40%	100,00%	97,10%

Sumber : Hasil Penelitian (2026)

3.4 Pembahasan

Nilai akurasi sebesar 94,50%, yang secara signifikan lebih tinggi daripada ambang batas hipotesis sebesar 75%, hipotesis H1 diterima dan H0 ditolak. Hasil ini mendukung keandalan algoritma SVM dengan Linear Kernel untuk mengklasifikasikan emosi teks ulasan berbahasa Indonesia, sejalan dengan beberapa studi serupa yang menyatakan kinerja SVM berkisar antara 82%–88% untuk studi kasus ulasan aplikasi lainnya [5], [6], [7], [8]. Hasil dalam penelitian ini bahkan sedikit lebih baik, kemungkinan besar karena karakteristik dataset yang sangat condong ke kelas negatif sehingga memudahkan model untuk mengenali pola kelas mayoritas.

Hal ini terlihat jelas dari ketimpangan hasil *Recall*, di mana kelas negatif mencapai skor sempurna 100%, sedangkan kelas positif hanya 8,33%. Ketidakseimbangan data (*data imbalance*) pada tahap pelatihan secara langsung menyebabkan model jauh lebih mahir dan sensitif dalam mengidentifikasi pola kata-kata keluhan (seperti “lambat”, “terputus”, “buruk”) dibandingkan pola kata-kata pujian. Meskipun model belum cukup berimbang (handal) dalam memprediksi kelas minoritas, hasil klasifikasi ini tetap relevan dalam menjawab tujuan penelitian, yaitu mengetahui kecenderungan opini pengguna terhadap layanan MyRepublic. Mengingat kecenderungan opini pengguna sangat didominasi oleh sentimen negatif (94%), kemampuan model yang memiliki sensitivitas tinggi dalam mendeteksi kelas negatif ini pada akhirnya memberikan nilai praktis bagi pihak pengembang untuk mengidentifikasi dan memprioritaskan penanganan keluhan layanan secara cepat.

Sebagai bentuk referensi evaluasi layanan yang konkret, hasil *WordCloud* yang memunculkan dominasi kata 'gangguan', 'koneksi', dan 'respon' dapat ditranslasikan menjadi rekomendasi perbaikan *User Experience* (UX) pada antarmuka aplikasi MyRepublic. Pertama, meskipun saat ini aplikasi telah memiliki indikator 'Modem Status' di halaman beranda, pengembang disarankan untuk memperluasnya menjadi fitur *Live Area Network Status*. Jika terjadi gangguan massal atau pemeliharaan jaringan, sistem dapat memunculkan *banner* peringatan (*alert*) tepat di bawah menu utama. Transparansi ini akan mencegah kebingungan pengguna yang berujung pada pemberian ulasan negatif di *Google Play Store*. Kedua, tingginya keluhan terkait 'respon' dan 'tugas' (CS) mengindikasikan bahwa menu 'Bantuan MyRep Care' yang saat ini masih banyak mengarahkan pengguna ke saluran eksternal seperti WhatsApp atau *E-mail* perlu dioptimalkan. Pengembang perlu memperkuat sistem pelaporan (*in-app ticketing*) dan *live chat* terintegrasi dengan agen manusia secara langsung di dalam ekosistem aplikasi agar penyelesaian masalah terasa lebih intuitif dan cepat.

4. Kesimpulan

Berdasarkan hasil studi dan hasil pengujian yang disebutkan di atas, dapat disimpulkan bahwa kerangka kerja *Knowledge Discovery in Databases* (KDD), yang menggabungkan

pendekatan *Text Mining* dan algoritma *Support Vector Machine* (SVM), secara efektif mengkategorikan 997 ulasan pengguna terhadap aplikasi MyRepublic di *Google Play Store* ke dalam kelas sentimen positif dan negatif. Model SVM dengan *Linear Kernel*, yang diuji pada rasio data pelatihan terhadap data uji sebesar 80:20 (797:200), secara matematis mencapai akurasi sebesar 94,50%, melebihi ambang batas hipotesis sebesar 75%. Namun, terdapat perbedaan yang sangat ekstrem dalam kinerja kedua kelas tersebut: kelas negatif memiliki *recall* sebesar 100%, sedangkan kelas positif hanya 8,33%. Hal ini menyimpulkan bahwa tingginya akurasi tersebut didorong oleh bias data mayoritas (94% ulasan bernada negatif). Model terbukti sangat sensitif dan efektif dalam mengidentifikasi keluhan konsumen, tetapi belum cukup handal untuk mengklasifikasikan komentar positif secara tepat.

Temuan kategorisasi dan ekstraksi kata ini dapat dimanfaatkan oleh pengembang aplikasi MyRepublic sebagai referensi evaluasi layanan yang konkret. Berdasarkan analisis sentimen, direkomendasikan perbaikan *User Experience* (UX) secara spesifik, seperti perluasan fitur *Modem Status* menjadi indikator informasi gangguan jaringan area secara *real-time*, serta optimalisasi antarmuka pada menu *Bantuan MyRep Care* agar pengguna tidak terlalu sering diarahkan ke platform eksternal. Penanganan keluhan yang terpusat dan responsif di dalam ekosistem aplikasi diharapkan dapat menekan angka ulasan negatif publik dan meningkatkan kualitas layanan secara keseluruhan.

Referensi

- [1] M. Bagas Dikal Putra and E. Setiawan, "Metode Lexicon Based Untuk Analisis Sentimen Pengguna Twitter Terhadap Kinerja Isp," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 6, pp. 12079–12087, 2024, doi: 10.36040/jati.v8i6.11765.
- [2] A. Adihsyah and Wahyudin, "Analisis Sentimen Pada Ulasan Aplikasi Home Credit Dengan Metode SVM dan K-NN," *J. Komput. Antart.*, vol. 1, no. 4, pp. 174–181, 2023, doi: 10.70052/jka.v1i4.50.
- [3] D. Safryda Putri and T. Ridwan, "ANALISIS SENTIMEN ULASAN APLIKASI POSPAY DENGAN ALGORITMA SUPPORT VECTOR MACHINE," *J. Ilm. Inform.*, vol. 11, no. 01, pp. 32–40, 2023, doi: 10.33884/jif.v11i01.6611.
- [4] G. Wahyudi, Rizki, Kusumawardhana, "Analisis Sentimen pada review Aplikasi Grab di Google Play Store Menggunakan Support Vector Machine," *J. Inform.*, vol. 8, pp. 1–8, 2021.
- [5] M. D. Hendriyanto, A. A. Ridha, and U. Enri, "Analisis Sentimen Ulasan Aplikasi Mola Pada Google Play Store Menggunakan Algoritma Support Vector Machine," *INTECOMS J. Inf. Technol. Comput. Sci.*, vol. 5, no. 1, pp. 1–7, Apr. 2022, doi: 10.31539/intecom.v5i1.3708.
- [6] M. M. Maarif and N. Setiyawati, "Analisis Sentimen Review Aplikasi LinkedIn di Google Play Store Menggunakan Support Vector Machine," *Progresif J. Ilm. Komput.*, vol. 20, no. 1, p. 454, Feb. 2024, doi: 10.35889/progresif.v20i1.1614.
- [7] Y. A. Nugroho and S. Sudarno, "Analisis Sentimen Aplikasi Gojek Pada Ulasan Pengguna di Google Play Store Menggunakan Metode Support Vector Machine," *J. Inf. Syst. Res.*, vol. 6, no. 4, pp. 2171–2179, 2025, doi: 10.47065/josh.v6i4.7891.
- [8] T. Setyani, K. Sari, H. N. Heidy, and 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," vol. 6, no. January, pp. 373–384, 2026.
- [9] I. Y. Yudo Bismo Utomo, Iin Kurniasari, "Penerapan Knowledge Discovery in Database," *J. Tek. Inform. Kaputama*, vol. 7, no. 1, 2023.
- [10] Google, "Google Play Store." Accessed: Jun. 18, 2026. [Online]. Available: <https://play.google.com/store>
- [11] Acte.in, "Support Vector Machine (SVM) Algorithm - Machine Learning | Everything You Need to Know | Updated 2026." Accessed: Jun. 18, 2026. [Online]. Available: <https://www.acte.in/support-vector-machine-algorithm-machine-learning-article>
- [12] H. Firda, P. Putra, N. R. Oktadini, P. E. Sevtiyuni, and A. Meiriza, "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

- using SVM (Case,” *Sist. J. Sist. Inf.* , vol. 14, pp. 516–528, 2025, [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [13] V. Fitriyana, Lutfi Hakim, Dian Candra Rini Novitasari, and Ahmad Hanif Asyhar, “Analisis Sentimen Ulasan Aplikasi Jamsostek Mobile Menggunakan Metode Support Vector Machine,” *J. Buana Inform.*, vol. 14, no. 01, pp. 40–49, 2023, doi: 10.24002/jbi.v14i01.6909.
- [14] Rina, “Memahami Confusion Matrix: Accuracy, Precision, Recall, Specificity, dan F1-Score untuk Evaluasi Model Klasifikasi.” Accessed: Jun. 19, 2026. [Online]. Available: <https://esairina.medium.com/memahami-confusion-matrix-accuracy-precision-recall-specificity-dan-f1-score-610d4f0db7cf>
- [15] B. Irawan, O. Nurdiawan, P. Studi, and T. Informatika, “Naive Bayes Dan Wordcloud Untuk Analisis Sentimen Wisata,” vol. 9, no. 1, pp. 236–242, 2023.