



Analisis Komparasi Sentimen Ulasan Pembaruan BRImo 2024 Menggunakan Algoritma SVM

Luthfi Firmansyah

Informatika, Universitas Sebelas April Sumedang

*Penulis Korespondensi: Luthfifirmansyah719@gmail.com

Abstract. Significant transformations in the user interface and features of the BRImo application throughout 2024 have generated diverse responses from users. Monitoring changes in customer perceptions is essential for developers to evaluate the effectiveness of these updates and maintain the quality of digital banking services. This study aims to compare user sentiment before and after the BRImo application update using a text mining approach with the Support Vector Machine (SVM) algorithm. User reviews were collected from the Google Play Store through a web scraping technique. The collected data were processed through several text preprocessing stages, including cleaning, case folding, tokenization, stopword removal, and stemming. Furthermore, the Term Frequency-Inverse Document Frequency (TF-IDF) method was applied for feature weighting before classification using the SVM algorithm. The experimental results show that the SVM model achieved an accuracy, precision, recall, and F1-score of 92%, indicating its effectiveness in sentiment classification. Comparative analysis revealed an increase in negative sentiment after the application update, mainly related to login issues, authentication problems, adaptation to the new interface, and system stability. In contrast, positive sentiment remained associated with the application's comprehensive features and transaction convenience. In conclusion, technical stability after system updates has a significant influence on user satisfaction, while the SVM algorithm provides an effective automated approach for evaluating user feedback and supporting future application improvements.

Keywords: BRImo; Comparison; Sentiment Analysis; Support Vector Machine; TF-IDF

Abstrak. Transformasi signifikan pada antarmuka dan fitur aplikasi BRImo sepanjang tahun 2024 telah memunculkan beragam respons dari pengguna. Pemantauan terhadap perubahan persepsi nasabah menjadi aspek penting bagi pengembang untuk mengevaluasi efektivitas pembaruan serta menjaga kualitas layanan perbankan digital. Penelitian ini bertujuan untuk membandingkan sentimen pengguna terhadap aplikasi BRImo sebelum dan sesudah pembaruan menggunakan pendekatan text mining dengan algoritma Support Vector Machine (SVM). Data penelitian diperoleh dari ulasan pengguna di Google Play Store melalui teknik scraping. Data yang terkumpul kemudian diproses melalui tahapan pra-pemrosesan teks yang meliputi cleaning, case folding, tokenizing, stopword removal, dan stemming. Selanjutnya, metode Term Frequency-Inverse Document Frequency (TF-IDF) digunakan untuk pembobotan fitur sebelum dilakukan proses klasifikasi menggunakan algoritma SVM. Hasil penelitian menunjukkan bahwa model SVM memperoleh nilai akurasi, precision, recall, dan F1-score masing-masing sebesar 92%, sehingga terbukti efektif dalam melakukan klasifikasi sentimen. Analisis komparatif menunjukkan adanya peningkatan sentimen negatif setelah pembaruan aplikasi, yang terutama disebabkan oleh kendala login, proses autentikasi, adaptasi terhadap antarmuka baru, serta masalah stabilitas sistem. Di sisi lain, sentimen positif tetap didominasi oleh apresiasi pengguna terhadap kelengkapan fitur dan kemudahan transaksi. Kesimpulannya, stabilitas teknis setelah pembaruan berpengaruh terhadap tingkat kepuasan pengguna, sementara algoritma SVM efektif digunakan sebagai metode otomatis untuk mengevaluasi umpan balik pengguna dan mendukung pengembangan aplikasi BRImo di masa mendatang.

Kata kunci: Analisis Sentimen; BRImo; Komparasi; Support Vector Machine; TF-IDF

1. LATAR BELAKANG

Transformasi digital di sektor perbankan Indonesia telah menempatkan aplikasi mobile banking sebagai komponen vital dalam ekosistem keuangan nasional. Sebagai salah satu aktor utama dalam ekosistem keuangan nasional, PT Bank Rakyat Indonesia (Persero) Tbk terus berkomitmen melakukan inovasi digital melalui platform BRImo. Upaya adaptasi ini ditujukan untuk merespons kebutuhan nasabah yang semakin menekankan aspek kecepatan. Sejumlah

penelitian empiris menegaskan adanya hubungan positif yang signifikan antara kualitas layanan elektronik (*e-service quality*) dengan tingkat loyalitas pengguna aplikasi perbankan (Alloerung et al., 2025), (Adut et al., 2022). Dalam konteks tersebut, nasabah menuntut standar layanan yang tinggi, terutama pada aspek keandalan (*reliability*) dan daya tanggap (*responsiveness*) sistem (Kebumen, 2023).

Penerapan pembaruan sistem pada aplikasi perbankan kerap menghadirkan sejumlah tantangan teknis yang berimplikasi langsung terhadap pengalaman pengguna. Saluran distribusi resmi seperti Google Play Store kerap menjadi wadah munculnya keluhan teknis, mulai dari kegagalan verifikasi wajah hingga hambatan akses login setelah pembaruan sistem (Pratama et al., n.d.), (Junanda & Alfresi, 2024). Ulasan pengguna tersebut merupakan sumber data tekstual yang memiliki nilai tinggi karena mampu mencerminkan persepsi publik secara langsung, sekaligus menjadi indikator penting dalam menilai kualitas layanan (Junanda & Alfresi, 2024). Kumpulan ulasan pengguna pada dasarnya merupakan aset data tekstual yang bernilai tinggi karena mampu merekam sentimen publik secara real-time, sehingga dapat dijadikan indikator valid dalam evaluasi kualitas layanan. Dengan mempertimbangkan besarnya volume data, pendekatan manual tidak lagi relevan, sehingga diperlukan metode analisis sentimen berbasis *Machine Learning* untuk mengekstraksi wawasan strategis secara otomatis (Putri et al., 2025), (Budi, 2024).

Pemilihan metode klasifikasi merupakan aspek mendasar dalam analisis data. *Support Vector Machine (SVM)* terbukti memiliki kemampuan unggul dalam mengolah data berdimensi tinggi. Algoritma ini mampu mencapai tingkat akurasi yang sangat tinggi, melampaui performa *Naive Bayes* dalam kasus analisis sentimen investasi (Maulana & Fahmi, 2024). Temuan lain turut menegaskan bahwa algoritma SVM mempunyai tingkat presisi lebih tinggi dibandingkannya *Naive Bayes* dalam pemetaan opini publik di media sosial, dengan perbedaan akurasi yang signifikan (Yudhistira, 2025). Meskipun demikian, ketimpangan distribusi data (*imbalanced data*) sering kali menurunkan performa model klasifikasi. Sebagai solusi, penerapan teknik *SMOTE* terbukti efektif dalam meningkatkannya sensitivitas sekaligus akurasi model, khususnya pada pengenalan kelas sentimen negatif (Hermawan et al., 2025).

Kinerja *Support Vector Machine (SVM)* tidak hanya dipengaruhi oleh penanganan data yang tidak seimbang, tetapi juga oleh pemilihan fungsi kernel yang tepat. Literatur menunjukkan adanya variasi temuan mengenai kernel yang dianggap paling optimal. Beberapa kajian mengindikasikan bahwa kernel Sigmoid dengan penyesuaian parameter dapat menghasilkan performa terbaik pada aplikasi perbankan digital (Wajidi & Cirua, 2025). Di sisi lain, beberapa penelitian menunjukkan bahwa kernel RBF lebih efektif digunakan untuk

analisa sentimen teks yang rumit dibandingkan pendekatan lainnya (Ernawati & Wati, 2024), (Rabbani et al., 2023). Di sisi lain, sejumlah kajian menunjukkan bahwa kernel Linear cukup efektif diterapkan dalam kasus analisa sentimen pada *software* BSI Mobile. (Meilani & Furqan, 2024). Perbedaan temuan studi menegaskan perlunya dilakukan perbandingan yang lebih spesifik pada dataset BRImo agar dapat diperoleh konfigurasi model yang presisi dan sesuai dengan sifat data.

Melihat pemaparannya, tujuannya studi ini guna melaksanakan analisa komparatif terhadap sentimen ulasan pembaruan BRImo 2024 dengan memanfaatkannya algoritma SVM. Fokus kajian diarahkan pada perbandingan kinerja empat jenis kernel SVM (Linear, RBF, Polynomial, dan Sigmoid) yang dikombinasikan dengan penerapan *Synthetic Minority Over-sampling Technique* (SMOTE) sebagai solusi atas permasalahan ketidakseimbangan data. Pendekatan ini diharapkan mampu menghasilkan model klasifikasi yang lebih presisi dalam mengevaluasi persepsi pengguna pasca-pembaruan aplikasi. Selain itu, penelitian ini ditujukan untuk memperkaya literatur yang sebelumnya masih banyak mengandalkan metode klasik seperti *Naive Bayes* (Umair & Sutanto, 2024) maupun varian lain dari *Support Vector* (Rosanti et al., 2024), sehingga dapat memberikan kontribusi baru bagi pengembangan analisis sentimen dalam sektor perbankan digital.

2. KAJIAN TEORITIS

Kumpulan ulasan pengguna merupakan sumber data tekstual yang memiliki nilai informasi tinggi karena memuat opini, pengalaman, kritik, dan saran yang mencerminkan persepsi pengguna terhadap kualitas suatu layanan. Informasi tersebut dapat dimanfaatkan sebagai indikator dalam mengevaluasi kualitas layanan sekaligus menjadi dasar bagi pengambilan keputusan. Namun, meningkatnya jumlah ulasan yang tersedia menyebabkan proses analisis secara manual menjadi kurang efektif dan efisien. Oleh karena itu, diperlukan penerapan analisis sentimen berbasis *Machine Learning* untuk mengolah data secara otomatis sehingga informasi yang terkandung dalam ulasan dapat diekstraksi dengan lebih cepat, akurat, dan konsisten (Putri et al., 2025; Budi, 2024).

Salah satu algoritma yang banyak digunakan dalam analisis sentimen adalah *Support Vector Machine* (SVM) karena memiliki kemampuan yang baik dalam mengklasifikasikan data berdimensi tinggi, termasuk data teks. SVM bekerja dengan membentuk *hyperplane* optimal yang memaksimalkan jarak (*margin*) antar kelas sehingga mampu meningkatkan akurasi klasifikasi. Pada data ulasan yang umumnya memiliki pola non-linear, SVM memanfaatkan

teknik *Kernel Trick* untuk memetakan data ke ruang fitur berdimensi lebih tinggi sehingga proses pemisahan antar kelas menjadi lebih optimal. Beberapa fungsi kernel yang umum digunakan meliputi *Kernel Linear*, *Radial Basis Function (RBF)*, *Kernel Polynomial*, dan *Kernel Sigmoid*, yang masing-masing memiliki karakteristik berbeda sesuai dengan pola distribusi data (Ernawati & Wati, 2024; Rabbani et al., 2023; Wajidi & Cirua, 2025)

Sebelum proses klasifikasi dilakukan, data teks terlebih dahulu direpresentasikan ke dalam bentuk numerik menggunakan metode *Term Frequency–Inverse Document Frequency (TF–IDF)*. Metode ini memberikan bobot pada setiap kata berdasarkan tingkat kepentingannya dalam dokumen sehingga mampu meningkatkan representasi fitur yang digunakan oleh model klasifikasi. Selain itu, untuk mengatasi permasalahan ketidakseimbangan distribusi kelas sentimen, diterapkan metode *Synthetic Minority Over-sampling Technique (SMOTE)*. Metode ini menghasilkan data sintesis pada kelas minoritas menggunakan pendekatan *k-nearest neighbors* sehingga distribusi data pelatihan menjadi lebih seimbang. Dengan kombinasi TF–IDF, SMOTE, dan SVM, diharapkan model analisis sentimen dapat memberikan hasil klasifikasi yang lebih akurat, objektif, dan andal (Hermawan et al., 2025).

3. METODE PENELITIAN

Pengumpulan Data (Data Collection)

Sumber datanya studi dari ulasan publik aplikasi BRImo yang tersedianya secara terbuka di Google Play Store. Proses akuisisi dilakukan dengan metode *scraping*, yaitu teknik ekstraksi otomatis yang memungkinkan pengambilan data dari situs web tanpa memerlukan akses API berbayar. Dataset yang diperoleh mencakup atribut penting seperti teks komentar, rating bintang (skala 1–5), serta penanda waktu ulasan. Pengumpulan data difokuskan pada periode setelah pembaruan aplikasi tahun 2024 hingga waktu penelitian berlangsung, sehingga hasil analisis tetap relevan dengan kondisi fitur terkini.

Pra-pemrosesan Data

Mengingat data hasil *scraping* masih bersifat mentah dan mengandung banyak *noise*, diperlukan tahapan standarisasi yang dilakukan secara sistematis agar data siap dianalisis. Tahap awal berupa proses *cleaning* untuk menghapus elemen non-tekstual yang tidak relevan bagi analisis sentimen, seperti emoji, simbol, angka, tanda baca, maupun tautan URL. Setelah itu, dilakukan penyeragaman format huruf melalui *case folding*, yang mana keseluruhan teks dikonversinya jadi huruf kecil guna mencegah duplikasi fitur. Tahap selanjutnya adalah normalisasi, yakni proses perbaikan istilah tidak baku (seperti typo atau bahasa gaul) menjadi bentuk standar bahasa Indonesia agar mesin dapat memahami konteks kalimat secara tepat.

Sesudah itu dilaksanakan *stopword removal* dengan menghapus kata-katanya hubung umum yang tidak mempunyai makna sentimen, misalnya “dan”, “yang”, atau “di”. Pada tahap terakhir, teknik *stemming* diterapkan menggunakan pustaka Sastrawi guna mengembalikannya kata berimbuhan ke bentuk dasar (*root word*), sehingga variasi morfologis dapat disederhanakan tanpa mengubah makna inti. Tahapan akhir dari pra-pemrosesannya yakni *stemming*, tujuannya mereduksi variasi kata berimbuhan menjadi bentuk dasarnya (*root word*) dengan bantuan pustaka Sastrawi, sehingga kompleksitas fitur kata dapat disederhanakan tanpa mengurangi makna (Puspa et al., 2023).

Pelabelan Data

Setelah proses standardisasi, dataset dibagi pada dua kategorinya sentimen utama, yakni positif serta negatif. Pelabelan dilaksanakan dengan pendekatan hibrida berbasis leksikon yang divalidasi menggunakan skor rating pengguna. Ulasannya rating bintang 4–5 ditetapkan sebagai sentimen positif, sedangkan rating bintang 1–2 dimasukkan ke dalam kategori negatif (Pratama et al., n.d.). Dengan cara ini, setiap ulasan tidak hanya dianalisis dari sisi teks, tetapi juga diperkuat oleh indikator numerik berupa rating, sehingga hasil pelabelan menjadi lebih konsisten dan dapat diandalkan untuk tahap analisis berikutnya.

Penyeimbangan Data

Ketidakseimbangan distribusi antara jumlah ulasan positif dan negatif berpotensi menimbulkan bias pada hasil prediksi, karena algoritma cenderung lebih fokus pada kelas mayoritas. Dalam mengatasinya permasalahan tersebut, studi ini memakai SMOTE. Metodologi ini menghasilkan data sintetis pada kelas minoritas dengan memanfaatkan prinsip *k-nearest neighbors*. Dengan cara ini, proporsi data latih menjadi lebih seimbang, hingga modelnya bisa mempelajari pola dari kedua kelas sentimen secara adil dan akurat.

Pembobotan Kata (TF-IDF)

Sistem penimbangan *Term Frequency–Inverse Document Frequency* (TF-IDF) digunakan untuk mengubah data teks menjadi vektor numerik. Pendekatan ini dipilih karena memberikan bobot tinggi pada frasa-frasa tertentu yang umum ditemukan dalam beberapa publikasi tetapi jarang muncul secara umum. Pendekatan ini efektif dalam menurunkan bobot kata umum yang tidak relevan sekaligus menonjolkan kata-kata yang bersifat diskriminatif. Secara matematis, bobot TF-IDF dapat dirumuskan sebagai:

$$w_{t,d} = TF_{t,d} \times \log\left(\frac{N}{df_t}\right) \quad (1)$$

di mana $w_{t,d}$ merepresentasikan bobot suatu kata dalam dokumen, $TF_{t,d}$ adalah frekuensi kemunculan kata, N ialah totalnya dokumen, serta df_t ialah jumlahnya dokumen berisikan kata tersebut. Dengan formula ini, tiap katanya pada ulasan pengguna dapat diubahnya jadi representasi numerik lebih informatif, sehingga memudahkan proses analisis sentimen pada tahap klasifikasi berikutnya.

Klasifikasi dengan Support Vector Machine (SVM)

Proses klasifikasinya sentimen pada studi ini memakai algoritma SVM yang bekerja berdasarkan prinsip pencarian *hyperplane* optimal. Tujuan utama pendekatan ini adalah memperbesar margin pemisah antara dua kelas data. Mengingat pola ulasan pengguna cenderung kompleks dan tidak selalu linear, penelitian ini mengimplementasi tekniknya *Kernel Trick* guna memproyeksikannya data ke ruang fitur berdimensi lebih tinggi, hingga pemisahannya kelas non-linear bisa dilaksanakan dengan optimal.

Untuk memperoleh performa terbaik, penelitian ini membandingkan empat jenis fungsi kernel yang berbeda. Perbandingan ini bertujuan untuk menemukan model klasifikasi yang paling sesuai dengan karakteristik data ulasan BRImo, sehingga hasil analisis sentimen dapat lebih akurat dan representatif (Wajidi & Cirua, 2025), (Rabbani et al., 2023), (Ernawati & Wati, 2024).

1. Linear Kernel: Kernel ini sesuai untuk data dengan dimensi fitur tinggi yang dapat dipisahkan secara linear. Secara matematis dirumuskan sebagai:

$$K(x_i, x) = x_i^T x \quad (2)$$

2. RBF (Radial Basis Function) Kernel: Kernel ini efektif dalam memetakan data ke ruang berdimensi tak hingga, sehingga mampu menangani pola non-linear yang kompleks. Rumusnya adalah:

$$K(x_i, x) = \exp(-\gamma ||x_i - x||^2) \quad (3)$$

3. Polynomial Kernel: Kernel ini memproyeksikan data ke ruang fitur polinomial dengan derajat tertentu, sehingga dapat menangkap hubungan non-linear antar fitur. Rumusnya dituliskan sebagai:

$$K(x_i, x) = (\gamma x_i^T x + r)^d \quad (4)$$

4. Sigmoid Kernel: Kernel ini bekerja mirip dengan fungsi aktivasi pada jaringan saraf tiruan, sehingga dapat digunakan untuk memodelkan hubungan non-linear. Rumusnya adalah:

$$K(x_i, x) = \tanh(\gamma x_i^T x + r) \quad (5)$$

Evaluasi Model

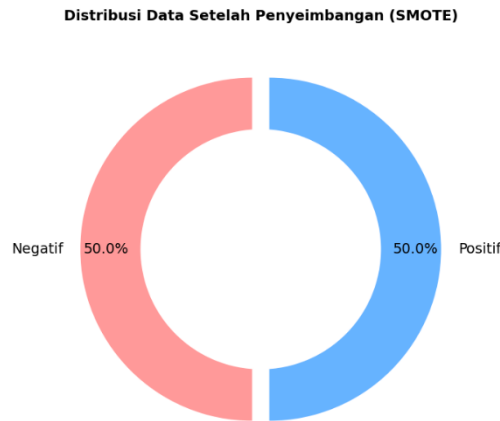
Validitas model diuji menggunakan matriks kebingungan (*Confusion Matrix*) yang memetakan hasil prediksi pada empat kuadran: TP, TN, FP, serta FN (Meilani & Furqan, 2024). Berdasarkan komponen tersebut, kinerja model diukur melalui parameter Akurasi, Presisi, *Recall*, dan *F1-Score*. Selain itu, analisa kurva ROC serta nilainya AUC dipakai guna menilainya keseimbangan sensitivitas model pada berbagai ambang batas klasifikasi..

Untuk menjamin keandalan model dan mencegah terjadinya *overfitting*, diterapkan metode validasi silang *K-Fold Cross Validation* dengan nilainya $k=10$. Mekanisme ini mempartisi dataset menjadi sepuluh segmen, di mana setiap segmen secara bergantian berperan sebagai data uji dan data latih, sehingga evaluasi yang dihasilkan lebih objektif dan stabil (Budi, 2024).

4. HASIL DAN PEMBAHASAN

Statistik dan Distribusi Data

Data ulasannya dihimpun dari Google Play Store dengan periodenya yang dipilih secara relevan untuk menangkap respons pengguna pasca-pembaruan aplikasi. Pada tahap eksplorasi awal, teridentifikasi adanya ketidakseimbangan distribusi sentimen, di mana ulasan positif jauh lebih dominan dibandingkan ulasan negatif. Kondisi ini berpotensi menimbulkan bias karena model cenderung memprediksi kelas mayoritas dan mengabaikan karakteristik kelas minoritas. Untuk mengatasi hal tersebut, penelitian ini menerapkan *Synthetic Minority Over-sampling Technique* (SMOTE). Teknik ini menghasilkan data sintesis pada kelas minoritas sehingga jumlah ulasan negatif dapat diseimbangkan dengan ulasan positif. Dengan demikian, distribusi data latih menjadi lebih proporsional dan memungkinkan model mempelajari pola kedua kelas secara adil. Visualisasi distribusi sebelum dan sesudah penerapan SMOTE ditampilkan pada Gambar 1, yang menegaskan efektivitas langkah penyeimbangan dalam mendukung kualitas pelatihan model.



Gambar 1. Distribusi Data Setelah Penyeimbangan (SMOTE)

Berdasarkan Gambar 1, terlihat bahwa distribusinya kelas pada dataset telah berhasil diseimbangkan melalui penerapannya SMOTE. Proporsi data kini terbagi merata, masing-masing 50% untuk kelas negatif dan 50% bagi kelas positif. Hal ini menegaskan bahwa permasalahan *imbalanced data* yang sebelumnya muncul telah sepenuhnya teratasi. Keseimbangan data memiliki peran krusial dalam proses pelatihan model *machine learning*, karena dengan distribusi yang setara model tidak lagi bias terhadap kelas mayoritas, melainkan terdorong untuk mempelajari karakteristik kedua kelas secara seimbang. Kondisi ini menjadikan proses pembelajaran lebih adil dan komprehensif, sehingga model mampu menangkap pola sentimen dengan tingkat akurasi yang lebih tinggi. Selain itu, keseimbangan data juga berkontribusi signifikan terhadap peningkatan performa model, khususnya dalam aspek akurasi dan presisi. Sebagaimana ditunjukkan pada hasil *confusion matrix*, model dapat melakukan klasifikasi dengan ketepatan yang tinggi. Dengan demikian, distribusi data yang seimbang tidak hanya menyelesaikan masalah teknis ketimpangan, tetapi juga memperkuat fondasi model untuk menghasilkan prediksi yang andal dan representatif terhadap persepsi pengguna.

Hasil Pra-pemrosesan Data

Pada tahapan pra-pemrosesannya, data ulasan teks diolah untuk mengurangi *noise* sekaligus menstandarisasi format agar lebih mudah diproses oleh algoritma klasifikasi. Prosedur yang dilakukan meliputi pembersihan karakter non-alfabet yang tidak relevan, penghapusan angka yang tidak berkontribusi terhadap analisis sentimen, normalisasi kata tidak baku agar sesuai dengan bentuk standar, serta penerapan *stemming* untuk mengembalikan kata berimbuhan ke bentuk dasar. Melalui tahapan ini, teks mentah yang semula beragam dan tidak terstruktur dapat diubah menjadi data yang lebih bersih, ringkas, dan siap digunakan dalam analisis berikutnya. Contoh hasil transformasi dari teks mentah menjadi teks bersih ditampilkan

pada Gambar 2, yang menunjukkan perbedaan signifikan antara kondisi awal dan hasil akhir setelah pra-pemrosesan.

Contoh Transformasi Data pada Tahap Preprocessing

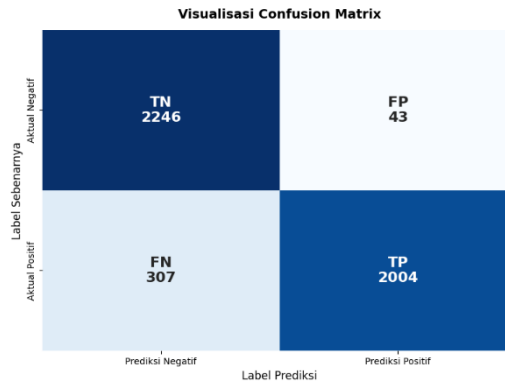
Teks Asli (Raw Data)	Hasil Proses (Clean Data)
aplikasi ap gimana nih, udh nyampe pembuatan akun tahap ke 7, gk bisa buat username	aplikasi ap gimana nih udh nyampe pembuatan akun tahap gk username
sangat membantu memudahkan saya	membantu memudahkan
login susah banget padahal sinyal bagus!!!	login susah banget padahal sinyal bagus
kecewa... verifikasi wajah gagal terus	kecewa verifikasi wajah gagal terus

Gambar 2. Contoh Transformasi Data Pada Tahap *preprocessing*

Gambar tersebut menunjukkan bahwa kalimat ulasan yang mengandung angka, tanda baca berlebihan, serta kata sambung yang tidak relevan telah dihapus pada tahap pra-pemrosesan. Proses ini dilakukan untuk memastikan teks yang dianalisis berfokus pada kata-kata mempunyai makna esensial bagi identifikasi sentimen. Setelah pembersihan, data dikonversi ke dalam representasi vektor numerik memakai metodologi pembobotan TF-IDF. Teknik ini memberikan bobot lebih tinggi pada kata-kata unik yang muncul dalam ulasan, sehingga dapat berfungsi sebagai fitur diskriminatif dalam klasifikasi. Dengan demikian, hasil pra-pemrosesan tidak hanya menghasilkan teks yang lebih ringkas dan terstruktur, tetapi juga menyiapkan representasi data yang lebih informatif bagi algoritma pembelajaran mesin.

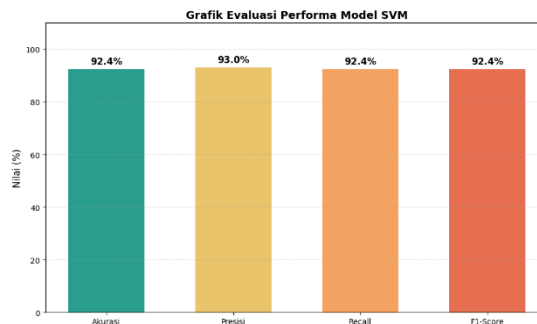
Hasil Klasifikasi dan Evaluasi Model

Klasifikasinya sentimen pada studi ini dilaksanakan dengan memanfaatkan algoritma SVM dengan kernel linear. Setelah tahap pelatihan selesai, model diuji menggunakan *testing set* untuk menilai reliabilitas prediksi yang dihasilkan. Proses pengujian ini krusial untuk mengetahui sejauh mana modelnya bisa mengenali pola sentimen secara konsisten dan akurat. Evaluasi performa kemudian divisualisasikannya pada bentuk *Confusion Matrix*, memberi deskripsi menyeluruh terkait distribusi prediksinya benar maupun salah. Visualisasinya tersebut tidak hanya menampilkan jumlah klasifikasi yang tepat, tetapi juga memperlihatkan kesalahan prediksi, sehingga dapat dijadikan dasar dalam menilai kekuatan dan keterbatasan model secara komprehensif.



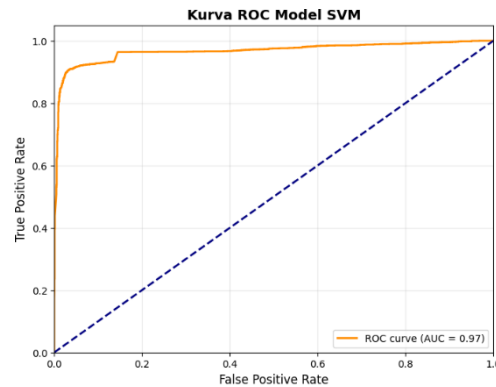
Gambar 3. Visualisasi Confusion Matrix

Berdasarkan Gambar 3, model mampu memprediksi dengan benar 2.246 data negatif (*True Negative*) dan 2.004 data positif (*True Positive*). Warna biru gelap pada diagonal utama menegaskan dominasi prediksi yang tepat, sementara kesalahan *False Positive* relatif kecil, yakni hanya 43 data. Temuan ini menunjukkan bahwa model memiliki tingkat presisi yang tinggi dan jarang salah dalam mengklasifikasikan keluhan sebagai pujian. Secara kuantitatif, performa model dievaluasi menggunakan empat metrik utama, yang divisualisasikan dalam bentuk grafik batang untuk memberikan gambaran menyeluruh.



Gambar 4. Grafik Evaluasi Performa Model SVM

Gambar 4 menunjukkan bahwa algoritma SVM mampu mempertahankan performa yang konsisten pada seluruh metrik evaluasi. Nilai akurasi tercatat sebesar 92,4%, dengan presisi 93,0% dan recall 92,4%. Keseimbangan antara presisi dan recall tercermin pada nilai F1-Score yang juga mencapai 92,4%. Temuan ini memperlihatkan modelnya tidak hanya mempunyai akurasi tinggi secara keseluruhan, namun cukup sensitif dalam mendeteksi keluhan pengguna dengan tingkat kesalahan prediksi yang rendah. Selain itu, kualitas model divalidasi lebih lanjut melalui analisis kurva ROC (*Receiver Operating Characteristic*).



Gambar 5. Kurva ROC Model SVM

Kurva pada Gambar 5 yang melengkung tajam ke sudut kiri atas dengan area di bawah kurva yang luas menunjukkan kemampuan pemisahan kelas (*separability*) yang sangat baik dari model. Bentuk kurva tersebut juga menegaskan bahwa tingkat kesalahan prediksi relatif rendah, sehingga model mampu membedakan sentimen positif dan negatif secara konsisten. Visualisasi ini menjadi bukti tambahan atas keandalan model dalam menghasilkan klasifikasi yang akurat..

Pembahasan

Berdasarkan hasil pengujian, dapat dianalisis efektivitas metode yang diterapkan. Tingkat akurasi sebesar 92,4% menunjukkan bahwa algoritma *Support Vector Machine* (SVM) dengan kernel linear efektif dalam mengklasifikasikan sentimen pada data ulasan perbankan berdimensi tinggi. Keberhasilan ini tidak terlepas dari penerapan teknik SMOTE yang menyeimbangkan proporsi data latih, sehingga model mampu mengenali pola sentimen positif maupun negatif secara lebih adil. Temuan ini juga menegaskan bahwa kombinasi pra-pemrosesan yang ketat dengan SVM linear pada dataset terbaru menghasilkan model klasifikasi yang presisi dan reliabel untuk kasus pembaruan aplikasi *mobile banking*. Selain itu, nilai *Recall* yang tinggi (92,4%) memiliki implikasi praktis penting, karena model mampu menangkap hampir seluruh keluhan pengguna dengan tingkat kelolosan yang rendah. Bagi pengembang aplikasi, kemampuan ini krusial untuk mendeteksi bug atau ketidakpuasan pengguna secara cepat pasca-pembaruan sistem, sehingga langkah perbaikan dapat segera dilakukan demi menjaga kepercayaan nasabah.

5. KESIMPULAN DAN SARAN

Kesimpulan

Berdasarkan hasil penelitian, dapat disimpulkan bahwa algoritma Support Vector Machine (SVM) dengan kernel linear yang dikombinasikan dengan metode Synthetic Minority Over-sampling Technique (SMOTE) mampu menghasilkan performa klasifikasi yang sangat baik, dengan nilai akurasi sebesar 92,4%, presisi 93,0%, recall 92,4%, dan F1-Score 92,4%. Hasil analisis sentimen menunjukkan bahwa sentimen negatif didominasi oleh kata-kata seperti "gagal", "login", "verifikasi", dan "susah", yang mengindikasikan bahwa permasalahan utama pengguna berkaitan dengan proses autentikasi dan akses masuk ke aplikasi. Sementara itu, sentimen positif didominasi oleh kata-kata seperti "mudah", "cepat", "praktis", dan "membantu", yang menunjukkan bahwa pembaruan aplikasi berhasil meningkatkan kemudahan penggunaan serta kecepatan transaksi bagi para pengguna.

Saran

Berdasarkan hasil penelitian, beberapa saran dapat diberikan untuk pengembangan selanjutnya. Tim pengembang BRImo disarankan memprioritaskan peningkatan stabilitas pada fitur login dan verifikasi wajah guna meningkatkan kepuasan serta kenyamanan pengguna dalam mengakses aplikasi. Selain itu, penelitian selanjutnya dapat membandingkan performa algoritma Support Vector Machine (SVM) dengan metode deep learning, seperti Long Short-Term Memory (LSTM) atau Bidirectional Encoder Representations from Transformers (BERT), untuk memperoleh hasil klasifikasi dengan tingkat akurasi yang lebih optimal. Penelitian berikutnya juga disarankan menerapkan Aspect-Based Sentiment Analysis (ABSA) agar analisis sentimen dapat dilakukan secara lebih spesifik terhadap setiap fitur layanan, seperti transfer, QRIS, dan investasi, sehingga menghasilkan informasi yang lebih mendalam sebagai dasar pengembangan layanan digital.

DAFTAR REFERENSI

- Adut, F. A., Ramadhan, V. P., Informasi, S., Malang, U. M., Candi, P., & Malang, K. (2022). *Pengaruh Kualitas Layanan Terhadap Kepuasan dan Loyalitas Pengguna Brimo Menggunakan Model E-Service Quality*. 204, 219–227.
- Allorerung, A., Pakiding, D. L., Tahendrika, A., Atma, U., & Makassar, J. (2025). *PENGARUH KUALITAS LAYANAN BRIMO TERHADAP KEPUASAN NASABAH PT BANK RAKYAT INDONESIA (PERSERO) TBK*. 07(01), 95–107.
- Budi, I. (2024). *The Indonesian Journal of Computer Science*. 13(4), 6533–6547.
- Ernawati, S., & Wati, R. (2024). *Evaluasi Performa Kernel SVM dalam Analisis Sentimen Review Aplikasi ChatGPT Menggunakan Hyperparameter dan VADER Lexicon*. 15(April), 40–49.

- Hermawan, M. A., Faqih, A., Dwilestari, G., Informatika, T., & Informasi, S. (2025). *IMPLEMENTASI AKURASI MODEL NAIVE BAYES MENGGUNAKAN SMOTE DALAM ANALISIS SENTIMEN PENGGUNA APLIKASI BRIMO*. 13(1).
- Junanda, O., & Alfresi, A. I. (2024). *BRI KCP Sudirman*. 410–418.
- Kebumen, N. B. R. I. (2023). *Pengaruh e-service quality, kepercayaan dan kemudahan terhadap keputusan penggunaan bri mobile (brimo) pada nasabah bri kebumen*. 2(1).
- Maulana, B. A., & Fahmi, M. J. (2024). *Sentiment Analysis of Pluang Applications With Naive Bayes and Support Vector Machine (SVM) Algorithm Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes dan Support Vector Machine (SVM)*. 4(April), 375–384.
- Meilani, N., & Furqan, M. (2024). *Analisis sentimen pengguna aplikasi BSI mobile akibat ransomware menggunakan algoritma support vector machine Sentiment analysis user application of BSI mobile due to ransomware using the support vector machine algorithm*. 5, 42–51. <https://doi.org/10.37373/infotech.v5i1.1102>
- Pratama, M. R., Ramadhan, Y. R., & Komara, M. A. (n.d.). *Analisis Sentimen BRIimo dan BCA Mobile Menggunakan Support Vector Machine dan Lexicon Based*.
- Puspa, T., Sanjaya, R., Fauzi, A., Fitri, A., & Masruriyah, N. (2023). *Analisis sentimen ulasan pada e-commerce shopee menggunakan algoritma naive bayes dan support vector machine Analysis of review sentiment on shopee e-commerce using the naive bayes algorithm and support vector machine*. 4, 16–26. <https://doi.org/10.37373/infotech.v4i1.422>
- Putri, S. A., Tania, K. D., & Pasemah, K. (2025). *Knowledge Discovery Through Sentiment Analysis and Topic Modeling of BCA Mobile and MyBCA*. 8(2), 669–682.
- Rabbani, S., Safitri, D., & Rahmadhani, N. (2023). *Comparative Evaluation of SVM Kernels for Sentiment Classification in Fuel Price Increase Analysis Perbandingan Evaluasi Kernel SVM untuk Klasifikasi Sentimen dalam Analisis Kenaikan Harga BBM*. 3(October), 153–160.
- Rosanti, C., Artanto, F. A., Saputra, R. E., Syariah, E., Muhammadiyah, U., & Pekalongan, P. (2024). *Analisis Sentiment Pengguna Aplikasi Mobile Banking Pada Bank Syariah Dengan Support Vector Regression Sentiment Analysis of Mobile Banking Application User in Sharia Bank Using Support Vector Regression*. 4(8), 341–347.
- Umair, M., & Sutanto, E. R. (2024). *Analisis Sentimen Ulasan Pengguna Pada Aplikasi BRIimo BRI Menggunakan Metode Klasifikasi Algoritma Naive Bayes*. 8(April), 1149–1159. <https://doi.org/10.30865/mib.v8i2.7381>
- Wajidi, F., & Cirua, A. A. A. (2025). *ANALISIS SENTIMEN ULASAN APLIKASI WONDR BY BNI MENGGUNAKAN ALGORITMA SVM DENGAN*. 4(2), 69–81.
- Yudhistira, A. (2025). *Analisis Sentimen Petani Milenial Pada Media Sosial X Menggunakan Algoritma Support Vector Machine (SVM) Fakultas Teknik dan Ilmu Komputer , Universitas Teknokrat Indonesia , Indonesia Sentiment Analysis of Millennial Farmers on Social Media X Using the Support Vector Machine (SVM) Algorithm*. 5(3), 845–857.