

ANALISIS SENTIMEN PENGGUNA APLIKASI SATUSEHAT MOBILE MENGUNAKAN SVM DAN NAÏVE BAYES DENGAN OPTIMASI SMOTE

M. Rasyid Ridho, Muanai Khalifah Revindo, Danny Matthew Saputra

Teknik Informatika, Universitas Sriwijaya

Jl. Palembang – Prabumulih KM.32 Kabupaten Ogan Ilir, Sumatera Selatan, Indonesia

09021282328119@student.unsri.ac.id

ABSTRAK

Aplikasi SATUSEHAT Mobile telah mencapai lebih dari 50 juta unduhan sejak diluncurkan tahun 2022 tetapi masih mengalami masalah teknis, dibuktikan dengan rating rata-rata 3,4 di Google Play Store. Keluhan pengguna yang umum mencakup kesalahan sistem saat mengakses fitur utama, navigasi antarmuka yang menantang untuk pengguna awam, dan keterbatasan sinkronisasi data medis antar fasilitas kesehatan. Penelitian ini menganalisis sentimen pengguna menggunakan algoritma Naïve Bayes dan Support Vector Machine (SVM) terhadap 3.000 ulasan yang dikumpulkan. Metodologi CRISP-DM diterapkan dengan enam tahapan sistematis mulai dari pengumpulan data hingga evaluasi mendalam. Temuan penelitian mengindikasikan adanya perbedaan distribusi sentimen pengguna yang cukup mencolok sebelum dilakukan optimasi sehingga diperlukan penerapan teknik SMOTE agar lebih seimbang. Hasil pengujian menunjukkan SVM mencapai akurasi 94% dengan *precision* 94%, *recall* 94%, dan *F1-score* 94% yang konsisten pada semua kategori sentimen. Naïve Bayes memperoleh akurasi 88% namun menunjukkan kelemahan pada *recall* sentimen negatif (73% vs 84% SVM). Uji McNemar mengkonfirmasi perbedaan performa signifikan secara statistik ($p < 0.001$, $\chi^2 = 52.94$). Analisis WordCloud mengungkapkan kata dominan pada sentimen negatif seperti "error", "login", dan "tidak bisa" yang mengindikasikan permasalahan teknis fundamental. Penelitian ini memberikan kontribusi dalam memahami persepsi pengguna dan menyediakan landasan strategis bagi peningkatan berkelanjutan aplikasi SATUSEHAT Mobile yang lebih responsif dan *user-friendly*.

Kata kunci : Analisis Sentimen, SATUSEHAT Mobile, Naïve Bayes, Support Vector Machine, SMOTE.

1. PENDAHULUAN

Dalam beberapa tahun ke belakang, transformasi digital dalam sektor kesehatan telah menjadi fokus utama pemerintah Indonesia. Aplikasi SATUSEHAT Mobile, diluncurkan pada tahun 2022, merupakan inisiatif strategis untuk memfasilitasi akses layanan kesehatan terintegrasi bagi masyarakat. Aplikasi ini merupakan versi yang telah dikembangkan dari aplikasi sebelumnya, yakni PeduliLindungi, yang pertama kali diluncurkan pada tahun 2020 sebagai respons terhadap pandemi COVID-19. Transformasi dari PeduliLindungi ke SATUSEHAT mencerminkan pergeseran fokus dari pelacakan kasus COVID-19 ke layanan kesehatan digital yang lebih luas dan terintegrasi. SATUSEHAT Mobile dirancang untuk menyederhanakan proses pendaftaran fasilitas kesehatan, konsultasi daring, dan pemantauan rekam medis, sehingga meningkatkan efisiensi dan transparansi layanan. Perkembangan aplikasi pemerintah seperti SATUSEHAT Mobile mencerminkan upaya untuk mengakselerasi digitalisasi sektor publik, khususnya pasca-pandemi COVID-19, di mana kebutuhan akan layanan kesehatan digital semakin mendesak [1].

Aplikasi pemerintah seperti SATUSEHAT Mobile hadir untuk mempermudah akses layanan kesehatan dan mengurangi ketimpangan di daerah terpencil, namun sering mendapat keluhan terkait *error* sistem, fitur terbatas, dan navigasi yang rumit. Sebelum aplikasi digital, masyarakat mengandalkan layanan konvensional seperti antrian langsung atau

konsultasi via telepon. Dengan SATUSEHAT Mobile, interaksi pasien dan penyedia layanan menjadi lebih efisien melalui fitur seperti akses informasi kesehatan, jadwal pemeriksaan, dan notifikasi vaksinasi secara *real-time*. Sampai tahun 2025, aplikasi ini telah diunduh lebih dari 50 juta kali dengan perolehan *rating* rata-rata 3,4 di Google Play Store.

Pemilihan SVM dan Naïve Bayes didasarkan pada karakteristik data ulasan berbentuk teks tidak terstruktur yang berpotensi mengalami ketidakseimbangan kelas. SVM dipilih karena unggul menangani data berdimensi tinggi lewat kernel trick serta mampu membangun hyperplane dengan margin optimal, yang terbukti mencapai akurasi 96% pada analisis sentimen vaksinasi COVID-19 di Twitter, mengungguli Naïve Bayes dan KNN yang mencapai 94% dan 91% [2]. Untuk mengatasi dominasi sentimen positif dalam dataset, diterapkan teknik SMOTE yang terbukti meningkatkan akurasi hingga 7% [3] dan mampu mendongkrak *recall* kelas minoritas hingga 6% di sektor kesehatan [4]. Kombinasi kedua metode ini menjadi pilihan tepat guna memahami persepsi pengguna SATUSEHAT Mobile secara menyeluruh.

Ulasan yang diberikan pengguna pada aplikasi pemerintah dapat dimanfaatkan sebagai bahan analisis sentimen untuk memahami persepsi masyarakat secara menyeluruh. Analisis sentimen sendiri adalah proses pengolahan data berbasis teks yang bertujuan mengidentifikasi pola opini, baik positif, netral, maupun negatif, yang terkandung dalam ulasan tersebut. Metode ini telah terbukti efektif dalam

konteks aplikasi kesehatan, Salah satu contoh adalah penelitian pada aplikasi PeduliLindungi yang berhasil mencapai akurasi 89% dengan memanfaatkan Support Vector Machine (SVM). Dalam studi ini, digunakan algoritma Naïve Bayes dan SVM untuk menganalisis 3.000 ulasan pengguna aplikasi SATUSEHAT Mobile. Kedua algoritma ini dipilih karena mampu menangani ketidakseimbangan data serta efektif dalam tugas klasifikasi teks [5].

Penelitian ini memiliki tujuan utama mengidentifikasi permasalahan teknis dan ekspektasi pengguna melalui analisis kata kunci yang dominan pada ulasan. Selain itu, penelitian ini juga bertujuan mengevaluasi kinerja Support Vector Machine (SVM) serta Naïve Bayes dalam mengelompokkan sentimen pengguna terhadap aplikasi SATUSEHAT Mobile. Kontribusi penelitian ini tidak hanya memperkaya kajian analisis sentimen di sektor publik, tetapi juga menyajikan wawasan praktis untuk pengembangan aplikasi pemerintah agar dapat meningkatkan layanan berdasarkan masukan dari pengguna. Lebih lanjut, hasil analisis ini diharapkan menjadi landasan bagi pemberian rekomendasi berbasis data demi peningkatan kualitas aplikasi pemerintah.

2. TINJAUAN PUSTAKA

2.1. Transformasi Digital SATUSEHAT

Sejak 2022, wajah layanan kesehatan digital Indonesia berubah cukup signifikan ketika PeduliLindungi digantikan oleh SATUSEHAT Mobile. Pergeseran ini bukan sekadar *rebranding*, tetapi fungsi aplikasinya ikut berkembang dari yang awalnya fokus pada pelacakan pandemi, kini merambah ke pengelolaan data kesehatan pribadi dan pemantauan mandiri pengguna. Platform ini dirancang untuk menyatukan seluruh data rekam medis pasien dari berbagai fasilitas kesehatan ke dalam satu sistem nasional yang terintegrasi dan saling terhubung [6]. Respons masyarakat yang terekam melalui ulasan di Google Play Store dapat menjadi salah satu cermin paling nyata untuk melihat seberapa efektif dan transparan sistem tersebut berjalan dalam praktik.

2.2. Analisis Sentimen

Analisis sentimen adalah teknik pemrosesan bahasa alami yang menentukan sentimen yang diungkapkan dalam data teks, mengklasifikasikannya sebagai positif, negatif, atau netral [7]. Dalam penelitian ini, ulasan pengguna dikategorikan menjadi sentimen positif, netral, dan negatif untuk membedakan antara kepuasan pengguna dan kendala teknis yang dialami. Proses ini memungkinkan penyedia layanan untuk memahami persepsi masyarakat secara masif dan objektif.

2.3. Algoritma Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah algoritma yang dikenal efektif dalam menghasilkan solusi optimal untuk tugas klasifikasi. Prinsip utama SVM adalah menemukan *hyperplane* optimal yang

mampu memisahkan titik data antar kelas dengan margin maksimal. Metode ini memiliki kemampuan generalisasi yang baik meskipun menggunakan data latih terbatas, serta mampu menangani data berdimensi tinggi. Namun, untuk data yang tidak bersifat linear, SVM perlu memanfaatkan fungsi kernel agar dapat bekerja secara optimal [8]. Keunggulan SVM terletak pada penggunaan *kernel trick* yang memungkinkannya menangani dimensi data yang besar tanpa memerlukan transformasi yang eksplisit, sehingga sangat efektif dalam mendeteksi nuansa sentimen pengguna.

2.4. Algoritma Naïve Bayes

Naïve Bayes merupakan metode klasifikasi berbasis statistik yang menerapkan Teori Bayes dengan asumsi independensi yang kuat antar fitur [9]. Meskipun asumsi tersebut jarang terpenuhi secara mutlak dalam bahasa alami, algoritma ini tetap populer karena efisiensi komputasinya yang sangat tinggi dan performa yang kompetitif pada *dataset* besar. Formula dasar Naïve Bayes menghitung probabilitas posterior dengan menggabungkan prior probability kelas dan likelihood dari fitur input yang ada.

2.5. SMOTE

Ketidakseimbangan distribusi label sering kali menjadi hambatan dalam pemodelan klasifikasi, di mana model cenderung berpihak pada kelas mayoritas. Synthetic Minority Over-sampling Technique (SMOTE) dapat digunakan menjadi solusi untuk membuat data sintesis baru pada kelas minoritas berdasarkan karakteristik fitur yang ada [4]. Berbeda dengan metode undersampling, SMOTE memastikan tidak ada informasi dari data asli yang hilang, sehingga membantu model memahami pola kelas minoritas secara lebih mendalam.

2.6. TF-IDF

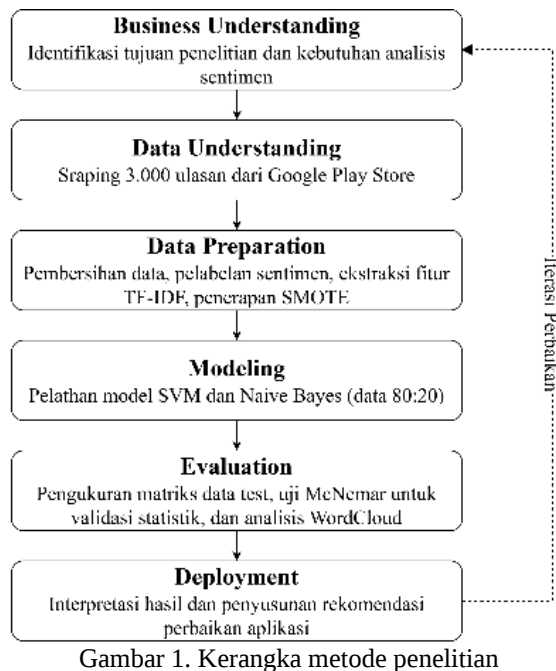
Untuk mentransformasi teks menjadi data numerik, digunakan metode Term Frequency-Inverse Document Frequency (TF-IDF). Teknik ini memberikan bobot lebih pada kata-kata yang sering muncul dalam satu dokumen tertentu namun jarang ditemukan di dokumen lain dalam satu korpus [10]. Dengan demikian, TF-IDF mampu menonjolkan kata-kata yang bersifat diskriminatif dan informatif, yang sangat membantu model dalam membedakan kategori dokumen berdasarkan distribusi katanya.

3. METODE PENELITIAN

3.1. Kerangka Kerja CRISP-DM

Penelitian ini dimulai dengan menggunakan metodologi CRISP-DM (Cross Industry Standard Process for Data Mining), yang terdiri atas enam tahap utama sebagaimana digambarkan pada Gambar 1, meliputi Business Understanding, Data Understanding melalui *web scraping* 3.000 ulasan Google Play Store, Data Preparation dengan pembersihan teks dan SMOTE, ekstraksi fitur menggunakan TF-IDF, Modeling dengan SVM dan Naïve Bayes, serta

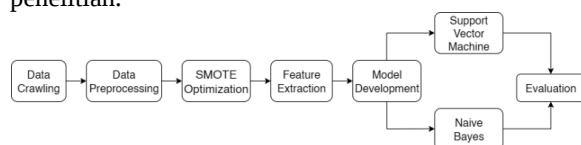
Evaluation dan Deployment untuk analisis dan rekomendasi perbaikan.



Gambar 1. Kerangka metode penelitian

3.2. Alur Proses Penelitian

Proses dimulai dengan pengumpulan data, yang disebut *data crawling*. Tahap pra-pemrosesan, yang disebut *data preprocessing*, diikuti oleh tahap ini. Ketidakseimbangan kelas pada data ditangani dengan teknik optimasi SMOTE. Setelah itu, ekstraksi fitur dilakukan untuk mengekstraksi atribut penting dari data yang telah diproses. Analisis sentimen kemudian memanfaatkan algoritma klasifikasi: Naïve Bayes dan Support Vector Machine (SVM). Berdasarkan evaluasi menggunakan berbagai metrik, model ini menunjukkan kinerja yang unggul dalam klasifikasi sentiment [11]. Gambar 2 menunjukkan alur proses penelitian.



Gambar 2. Alur proses penelitian

3.3. Data Crawling

Tahap awal penelitian ini dimulai dengan proses pengumpulan data menggunakan teknik *web crawling* untuk memperoleh kumpulan data yang representatif [12]. Data yang digunakan berupa *review* pengguna aplikasi SATUSEHAT Mobile yang diperoleh langsung dari halaman resmi aplikasi pada Google Play Store. Pengambilan data dilakukan dengan mengakses serta mengekstraksi bagian komentar atau ulasan yang berada di bawah deskripsi aplikasi. Rentang waktu pengumpulan data mencakup periode April 2020 hingga April 2025, dengan total sebanyak 3.000 ulasan yang dikumpulkan secara acak. Ulasan ini mencakup seluruh variasi skor bintang, mulai dari

1 hingga 5. Keberagaman skor ini dimaksudkan untuk memperoleh cakupan opini pengguna yang luas, baik yang bersifat negatif, netral, maupun positif, sehingga data yang diperoleh dapat merepresentasikan persepsi pengguna secara menyeluruh terhadap aplikasi tersebut.

3.4. Data Preprocessing

Preprocessing data merupakan tahap awal yang krusial dengan mengubah data mentah dari berbagai sumber menjadi informasi yang lebih bersih untuk dianalisis [13]. Tujuan utama dari tahapan ini adalah untuk menyederhanakan dan menyeragamkan format data sehingga dapat diolah secara lebih efisien dan akurat. Dengan data yang telah dipreproses, algoritma pengelompokan atau klasifikasi dapat bekerja secara optimal karena gangguan dari elemen-elemen tidak relevan atau inkonsisten telah diminimalkan.

Proses *preprocessing* meliputi beberapa tahap penting, di antaranya: *cleansing* untuk menghapus karakter atau elemen yang tidak relevan; *case folding* guna membuat seluruh huruf menjadi bentuk kecil; serta *filtering* dan *stopword removal* untuk menghapus kata-kata umum yang tidak memberikan makna signifikan dalam proses analisis; *tokenizing* untuk memecah kalimat menjadi potongan-potongan kata; *stemming* yang mengembalikan kata ke bentuk dasarnya; serta *labeling*, yaitu pemberian label pada data sesuai dengan kategori atau klasifikasi tertentu.

3.5. SMOTE Optimization

Untuk mengatasi ketidakseimbangan kelas dalam data, SMOTE (Synthetic Minority Over-sampling Technique) digunakan untuk menghasilkan sampel sintesis untuk kelas minoritas dalam kumpulan data yang tidak seimbang, mengatasi masalah ketidakseimbangan kelas. Teknik ini meningkatkan kinerja model dengan meningkatkan representasi kelas minoritas, yang mengarah ke generalisasi yang lebih baik dan mengurangi bias dalam prediksi [14].

3.6. Ekstraksi Fitur

Pada tahap ekstraksi fitur, token hasil pra-pemrosesan dikonversi menjadi representasi numerik berupa vektor. Untuk mengukur tingkat kepentingan suatu kata dalam sebuah dokumen dibandingkan dengan seluruh korpus, digunakan metode Term Frequency–Inverse Document Frequency (TF-IDF) [15]. Pendekatan ini menyediakan bobot lebih pada kata yang sering muncul di satu dokumen tetapi jarang ditemukan di dokumen lain, sehingga kata tersebut dianggap relevan serta informatif untuk proses klasifikasi. Dengan demikian, TF-IDF membantu model membedakan kelas dokumen berdasarkan distribusi kata yang bersifat diskriminatif.

3.7. Pengembangan Model

Pada tahap pemodelan, diterapkan dua algoritma klasifikasi, yakni Support Vector Machine (SVM) dan Naïve Bayes, untuk menganalisis sentimen dari ulasan

pengguna aplikasi SATUSEHAT Mobile. Pemilihan kedua algoritma ini dilakukan guna membandingkan performa model dalam mengklasifikasikan sentimen berdasarkan fitur teks yang telah diekstraksi.

Untuk klasifikasi dan regresi, Support Vector Machine (SVM) menemukan *hyperplane* terbaik dengan memaksimalkan margin antara dua kelas data [16]. Dengan menggunakan teknik *kernel trick*, SVM dapat menangani data besar tanpa perlu melakukan transformasi yang jelas. Sebagian kecil data digunakan dalam pelatihan untuk membuat model. Pemilihan fungsi kernel dan pengaturan parameter seperti C dan gamma sangat memengaruhi kinerja SVM.

Salah satu teknik statistik yang digunakan untuk klasifikasi adalah estimasi kemungkinan bahwa suatu objek atau data akan termasuk ke dalam kelas tertentu [2]. Teknik ini berdasarkan pada teori Bayes yang kuat, yaitu setiap fitur dalam data saling independen antara satu sama lain terhadap kelas yang dimaksud. Meskipun asumsi ini jarang benar secara mutlak dalam praktik, metode naif ini membuat model ini sangat efisien dan ringan secara komputasional, sehingga cocok untuk *dataset* berukuran besar. Naïve Bayes telah terbukti memberikan performa klasifikasi yang kompetitif, sebanding dengan metode populer lainnya seperti decision tree dan *artificial neural networks*, terutama dalam domain seperti analisis teks dan klasifikasi dokumen [17].

$$P(h|D) = \frac{P(D|h)P(h)}{P(D)} \quad (1)$$

Persamaan (1) menunjukkan formula dasar dari Naïve Bayes yang menggambarkan bagaimana probabilitas posterior dihitung dengan menggabungkan probabilitas prior kelas dan *likelihood* dari fitur-fitur *input*. Probabilitas posterior ditulis sebagai $P(h|D)$, probabilitas *likelihood* sebagai $P(D|h)$, probabilitas prior sebagai $P(h)$, dan evidence atau probabilitas marginal sebagai $P(D)$.

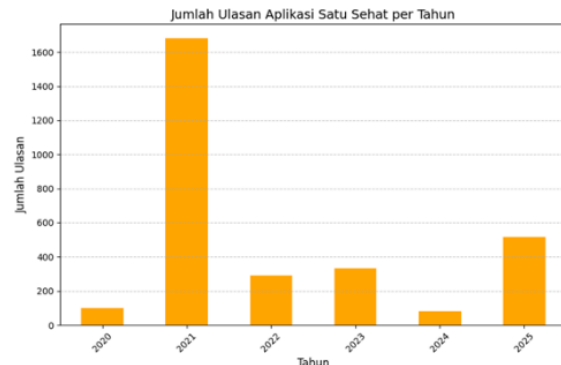
4. HASIL DAN PEMBAHASAN

4.1. Dataset

Dalam periode April 2020 hingga April 2025, terkumpul total 3.000 ulasan. Gambar 3 memperlihatkan distribusi jumlah ulasan pengguna aplikasi SATUSEHAT selama kurun waktu tersebut. Tahun 2021 menunjukkan lonjakan tajam dengan lebih dari 1.600 ulasan, yang kemungkinan mencerminkan peningkatan penggunaan aplikasi secara signifikan. Fenomena ini besar kemungkinan dipengaruhi oleh faktor eksternal, seperti pandemi COVID-19, yang mendorong masyarakat untuk lebih sering memanfaatkan layanan kesehatan digital [18].

Tahun 2022 dan 2023 menunjukkan penurunan drastis namun stabil dibanding 2021, dengan masing-masing sekitar 300-an ulasan. Ini mungkin menandakan bahwa setelah puncak penggunaan, minat pengguna mulai menurun seiring meredanya kebutuhan darurat. Namun menariknya, pada tahun

2024 jumlah ulasan kembali turun drastis ke titik terendah sejak 2020, sebelum kemudian meningkat lagi secara signifikan di 2025. Kenaikan ini bisa jadi pertanda adanya pembaruan aplikasi, kampanye pemasaran baru, atau integrasi fitur yang menarik perhatian pengguna kembali.



Gambar 3. Grafik ulasan berdasarkan tahun

4.2. Data Preprocessing

Hasil *crawling* masih berupa *raw data* yang belum bisa dianalisis secara langsung. Oleh sebab itu, dibutuhkan sejumlah langkah *preprocessing* untuk mengonversinya menjadi format yang lebih terstruktur dan siap diolah secara sistematis. Proses *preprocessing* ini bertujuan untuk memastikan data berada dalam kondisi yang siap digunakan pada tahap analisis berikutnya. Adapun tahapan *preprocessing* yang dilakukan adalah sebagai berikut:

4.3. Cleansing

Tahap pembersihan bertujuan menghilangkan karakter yang tidak relevan untuk analisis, seperti angka, tanda baca, dan simbol khusus. Hanya huruf yang dipertahankan agar teks tetap berada dalam bentuk alfabet aslinya.

4.4. Case Folding

Dalam proses standarisasi huruf, *case folding* digunakan untuk mengonversi seluruh karakter menjadi huruf kecil. Langkah ini bertujuan menghindari perbedaan interpretasi akibat variasi penggunaan huruf kapital dan huruf kecil, sehingga analisis dapat berlangsung secara lebih konsisten.

4.5. Tokenizing

Pada proses *tokenizing*, teks yang sudah melalui tahap pembersihan dipisahkan menjadi unit-unit kecil berupa token atau kata. Masing-masing token berperan sebagai satuan kata yang menjadi landasan untuk analisis pada tahap berikutnya.

4.6. Filtering

Tahap *filtering* mencakup penghapusan elemen-elemen yang tidak relevan atau mengganggu, seperti URL, tag HTML, angka, serta kata-kata umum yang tidak memiliki makna signifikan dalam analisis. Tahapan ini turut meliputi proses transformasi

tambahan, misalnya mengonversi seluruh huruf menjadi huruf kecil guna mempertahankan konsistensi.

4.7. Stemming

Tujuan dari proses *stemming* adalah untuk mengurangi kata ke bentuk aslinya (*root word*). Dalam hal ini, kata-kata seperti "bermain", "bermainlah", dan "memainkan" akan diubah menjadi "main". Teknik ini lebih awal dikembangkan untuk digunakan dalam bahasa Inggris, tetapi jika diterapkan dalam bahasa Indonesia, hasil yang signifikan dapat dicapai jika menggunakan algoritma *stemming* yang tepat.

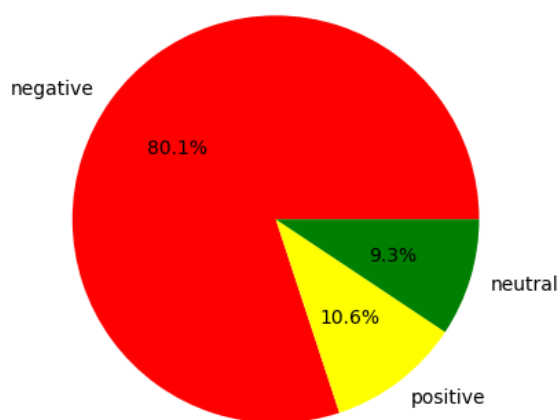
4.8. Labeling

Tahap berikutnya adalah *labeling*, yakni proses pemberian penanda atau klasifikasi emosional pada data sesuai dengan konteks yang dimilikinya. Berdasarkan skala bintang yang diberikan pengguna, penilaian dibagi menjadi tiga kategori: positif untuk skor 4–5, netral untuk skor 3, dan negatif untuk skor 1–2.

4.9. Pengujian dan Evaluasi

Hasil pengujian awal menunjukkan perbedaan yang cukup mencolok antara jumlah ulasan yang menunjukkan sentimen positif, netral, dan negatif. Dengan distribusi label terbanyak mulai dari 2.402 ulasan negative, 318 ulasan positif, dan 280 ulasan netral. Untuk mengatasi ketimpangan distribusi data ini, teknik Synthetic Minority Over-sampling Technique (SMOTE) digunakan. Jumlah ulasan positif ditingkatkan hingga setara dengan jumlah ulasan negatif. *Dataset* menjadi lebih seimbang setelah optimasi, karena kedua kategori sentimen memiliki jumlah ulasan yang sama, yaitu 2.402.

Dalam penelitian ini, sebanyak 5.765 data digunakan sebagai data latih, yang mencakup sekitar 80% dari total 7.206 data yang tersedia. Sementara itu, sisanya sebanyak 1.441 data (sekitar 20%) dialokasikan sebagai data uji.



Gambar 4. Grafik distribusi sentimen

Proporsi ini dipilih dengan pertimbangan bahwa model memerlukan volume data yang lebih besar pada

tahap pelatihan agar mampu mengenali pola dan karakteristik sentimen secara lebih menyeluruh. Dengan memberikan porsi data latih yang lebih dominan, diharapkan model dapat membangun pemahaman yang lebih kuat dan akurat terhadap data, sehingga performanya dalam mengklasifikasikan data baru menjadi lebih optimal.

4.10. Support Vector Machine

Optimalisasi SVM melalui metode SMOTE menghasilkan akurasi keseluruhan sebesar 94%. Pada sentimen positif, model mencatat presisi 94%, recall 99%, dan F1-score 97%, menunjukkan performa yang kuat. Untuk sentimen netral, diperoleh presisi 91%, recall 99%, dan F1-score 95%, yang menandakan keseimbangan kinerja. Sementara itu, pada kategori sentimen negatif, presisi mencapai 99%, recall 84%, dan F1-score 91%, mengindikasikan kemampuan deteksi yang baik meski terdapat beberapa kesalahan klasifikasi. Secara keseluruhan, hasil ini membuktikan bahwa penerapan SMOTE pada SVM mampu menghasilkan klasifikasi yang konsisten dan andal di semua kelas sentimen. Rincian lengkap ditampilkan pada Tabel 1.

Tabel 1. Laporan klasifikasi SVM

Support Vector Machine			
	Precision	Recall	F1-score
Positif	94%	99%	97%
Netral	91%	99%	95%
Negatif	99%	84%	91%

4.11. Naïve Bayes

Tabel 2 menampilkan *classification report* model setelah penerapan teknik SMOTE, dengan akurasi total sebesar 88%. Pada kategori sentimen negatif, nilai presisi mencapai 94%, recall berada di angka 73%, dan F1-score sebesar 82%. Hasil ini menunjukkan kalau model memiliki ketepatan yang sudah baik untuk mengenali ulasan negatif, meskipun masih ada keterbatasan dalam menangkap seluruh sampel negatif yang sesungguhnya. Pada klasifikasi sentimen netral, diperoleh presisi sebesar 83%, recall 95%, dan F1-score 89%, yang menandakan kemampuan model yang kuat dalam mengenali sebagian besar ulasan netral dengan tingkat kesalahan yang relatif rendah. Untuk kategori sentimen positif, hasil yang dicapai tergolong paling tinggi, yakni presisi 90%, recall 97%, dan F1-score 93%. Angka tersebut menunjukkan keseimbangan yang sangat baik dalam proses identifikasi sekaligus pengklasifikasian ulasan positif. Secara keseluruhan, hasil ini mengindikasikan bahwa model bekerja secara efektif pasca penyeimbangan data, meskipun masih terdapat ruang untuk perbaikan terutama dalam mendeteksi sentimen negatif secara lebih komprehensif.

Data hasil dari proses *crawling* masih berada dalam kondisi mentah dan belum siap untuk dianalisis secara langsung. Oleh sebab itu, diperlukan serangkaian langkah *preprocessing* guna mengonversi

data ke dalam format yang lebih terstruktur sehingga dapat diolah secara sistematis. Tujuan dari proses *preprocessing* ini adalah untuk memastikan bahwa data siap untuk digunakan dalam tahap analisis berikutnya. Berikut adalah penjelasan tentang proses *preprocessing*:

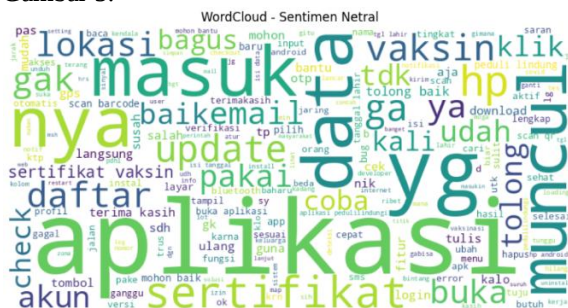
Tabel 2. Laporan klasifikasi Naïve Bayes

Naïve Bayes			
	Precision	Recall	F1-score
Positif	90%	97%	93%
Netral	83%	95%	89%
Negatif	94%	73%	82%

4.12. Visualisasi

Hasil WordCloud yang menunjukkan sentimen netral tentang aplikasi SATUSEHAT Mobile menunjukkan bahwa kata-kata "aplikasi" dan "masuk" mendominasi, menunjukkan bahwa fokus utama adalah pada platform dan proses akses. Faktor teknis seperti manajemen dan verifikasi data kesehatan ditampilkan dengan kata "data". Jumlah besar dari kata "vaksin" dan "sertifikat" menunjukkan bahwa banyak pengguna menggunakan aplikasi ini untuk melacak vaksinasi, mungkin merujuk pada kebutuhan sertifikat vaksin selama pandemi. Selanjutnya, kata "*update*" menunjukkan bahwa diskusi cenderung membahas pembaruan aplikasi dan fiturnya, kata "daftar" menunjukkan proses registrasi SATUSEHAT Mobile, dan kata "klik" dan "buka" menunjukkan interaksi teknis dengan antarmuka aplikasi. Secara keseluruhan, analisis ini membahas berbagai masalah yang dihadapi oleh pengguna SATUSEHAT Mobile, mulai dari pengetahuan teknis hingga persyaratan penting untuk dokumentasi kesehatan.

Kata-kata netral ini mungkin tampak "tenang", tetapi jika dilihat dari perspektif kontekstual, mereka menjadi titik penting dalam pengalaman pengguna. Misalnya, analisis berbasis frekuensi tidak cukup untuk memahami valensi emosi dalam kalimat pengguna; kata "masuk" dapat bermakna frustrasi tergantung pada konteksnya, seperti "sudah coba masuk tapi gagal terus." Visualisasi WordCloud Aplikasi SATUSEHAT Mobile dapat dilihat pada Gambar 5.



Gambar 5. Wordcloud sentimen netral

Dalam analisis WordCloud sentimen positif terhadap aplikasi SATUSEHAT Mobile, kata-kata yang paling sering muncul adalah "aplikasi", "bagus", "data", "hasil", "vaksin", dan "terima kasih". Hal ini

mencerminkan apresiasi pengguna terhadap fungsi aplikasi, terutama dalam pengelolaan data kesehatan dan vaksin. Kata-kata seperti "bagus" dan "mudah" menunjukkan pengalaman positif terhadap antarmuka dan kemudahan penggunaan aplikasi ini.



Gambar 6. Wordcloud sentimen positif

Kemunculan kata "terima kasih" mengindikasikan kepuasan dan rasa syukur dari pengguna atas layanan yang diberikan. Meskipun kata "coba" juga muncul, ini bisa menunjukkan ketertarikan awal atau kesan positif setelah penggunaan pertama. Untuk pemahaman yang lebih lengkap, analisis lanjutan terhadap konteks kata – kata tersebut tetap dibutuhkan. Gambar 6 menampilkan visualisasi WordCloud dari sentimen positif ini.

Menurut analisis WordCloud, kata-kata seperti "error", "gagal", "salah", dan "susah" menunjukkan masalah teknis yang dihadapi pengguna SATUSEHAT Mobile, mungkin terkait dengan masalah saat mengakses atau menggunakan aplikasi. Kata "sudah" yang paling umum menunjukkan frustrasi pengguna saat mencoba melakukan proses tetapi tidak berhasil. Selain itu, kemunculan istilah "data" mengindikasikan adanya ketidakpuasan pengguna terhadap pengelolaan atau proses pengisian data. Sementara itu, kata "tolong" mencerminkan adanya permintaan bantuan atau rasa tidak puas yang berkaitan dengan berbagai permasalahan yang mereka alami. Selain itu, pengguna sering menggunakan istilah "buka" dan "aplikasi" untuk menunjukkan bahwa mereka menghadapi kesulitan untuk mengakses aplikasi atau fiturnya. Analisis ini menggambarkan kesulitan teknis yang dihadapi pengguna dan menekankan bahwa pengalaman pengguna harus diperbaiki untuk menyelesaikan masalah terutama tentang stabilitas sistem dan kemudahan akses.



Gambar 7. WordCloud sentimen negative

Penggunaan kata-kata seperti "tolong" dan "sudah" dalam situasi negatif membawa beban emosional yang signifikan, yang dapat menyebabkan pola frustrasi jika dipetakan secara tekstual. Ini merupakan sinyal penting untuk melakukan perbaikan pada fitur paling penting, seperti akses sertifikat, sinkronisasi data, dan *login*. Gambar 7 memberikan visualisasi yang jelas tentang aplikasi SATUSEHAT Mobile sentimen negatif.

4.13. Interpretasi Hasil

Hasil interpretasi menunjukkan bahwa pengujian pada aplikasi SATUSEHAT Mobile menghasilkan kinerja yang baik pada kedua algoritma, yakni Support Vector Machine (SVM) dan Naïve Bayes, setelah penerapan teknik SMOTE. SVM memperoleh akurasi sebesar 94%, sedangkan Naïve Bayes mencapai 88%. SVM juga menunjukkan keseimbangan recall, presisi, dan F1-score pada klasifikasi sentimen positif, netral, dan negatif. Uji McNemar menghasilkan nilai $p < 0.001$ ($\chi^2 = 52.94$), yang menandakan bahwa perbedaan performa kedua model signifikan secara statistik. Hal ini memperlihatkan bahwa SVM secara konsisten mengungguli Naïve Bayes dalam jumlah prediksi yang benar. Meski demikian, pemilihan akhir model tetap perlu mempertimbangkan konteks aplikasi dan tujuan analisis. Temuan ini konsisten dengan studi terdahulu yang juga menegaskan keunggulan SVM di berbagai skenario analisis sentimen.

5. KESIMPULAN DAN SARAN

Hasil penelitian menunjukkan bahwa kedua algoritma, Support Vector Machine (SVM) dan Naïve Bayes berfungsi dengan baik setelah penerapan metode SMOTE. SVM memperoleh akurasi sebesar 94% dengan F1-score yang seimbang pada seluruh kategori sentimen, sedangkan Naïve Bayes menunjukkan kelemahan dalam mengidentifikasi sentimen negatif (recall 73% dibandingkan 84%). Perbedaan kinerja tersebut terbukti signifikan secara statistik melalui hasil uji McNemar ($p < 0.001$). SVM secara konsisten menghasilkan prediksi yang lebih tepat. Hasil ini mendukung pilihan SVM sebagai model pemahaman persepsi pengguna utama dan sebagai dasar untuk meningkatkan kualitas aplikasi SATUSEHAT Mobile.

Analisis WordCloud menunjukkan bahwa sentimen netral berfokus pada masalah teknis seperti akses dan data vaksin, sedangkan sentimen positif menunjukkan kemudahan penggunaan, dan sentimen negatif menunjukkan masalah teknis. Temuan ini menunjukkan partisipasi digital masyarakat secara langsung, mengungkap pola frustrasi kolektif akibat hambatan teknis, dan menegaskan bahwa fitur yang bermanfaat dihargai. Jika dibiarkan, masalah *login*, verifikasi, dan data yang tidak sinkron dapat mengikis kepercayaan publik. Ini bukan hanya masalah teknis. Sebaliknya, ulasan positif tentang kemudahan penggunaan dan sertifikat vaksin menunjukkan bahwa persepsi publik dapat dipengaruhi langsung oleh

perbaikan layanan. Dari sisi lain, data ini membantu pengembang aplikasi dan pembuat kebijakan menemukan kesalahan sistem, memperbaiki desain, dan memprioritaskan aspek yang paling relevan bagi masyarakat.

Penelitian ini berkontribusi dalam memahami persepsi pengguna terhadap aplikasi SATUSEHAT Mobile serta menyajikan metode analisis sentimen yang efektif. Temuan ini dapat dijadikan dasar untuk meningkatkan kualitas aplikasi, khususnya dalam menyelesaikan masalah teknis, memperkuat fitur-fitur yang mendapat respons positif dari pengguna, dan juga menawarkan kerangka kerja strategis untuk memastikan bahwa perubahan layanan digital pemerintah berjalan dengan cara yang responsif, inklusif, dan berkelanjutan.

DAFTAR PUSTAKA

- [1] R. Komalasari, "Manfaat teknologi informasi dan komunikasi di masa pandemi COVID-19," *TEMATIK*, vol. 7, no. 1, p. 38, Jun. 2020, doi: 10.38204/tematik.v7i1.369.
- [2] Munawaroh and A. Alamsyah, "Performance comparison of SVM, Naïve Bayes, and KNN algorithms for analysis of public opinion sentiment against COVID-19 vaccination on Twitter," *Journal of Advances in Information Systems and Technology*, vol. 4, no. 2, p. 113, Mar. 2023, doi: 10.15294/jaist.v4i2.59493.
- [3] N. Matondang and N. Surantha, "Effects of Oversampling SMOTE in the Classification of Hypertensive Dataset," *Advances in Science, Technology and Engineering Systems Journal*, vol. 5, no. 4, pp. 432–437, Aug. 2020, doi: 10.25046/AJ050451.
- [4] P. Budi Utomo, M. Faruqziddan, E. Herdika Septa Aulia, and S. Dini Azzahra, "Perbandingan Skenario Balancing Oversampling dan Undersampling dalam Klasifikasi Resiko Kambuh Kanker Tiroid menggunakan Algoritma SVM Linear", *jami*, vol. 5, no. 2, pp. 172 - 182, Dec. 2024.
- [5] U. H. Laksono and R. R. Suryono, "Sentiment analysis of online dating apps using Support Vector Machine and Naïve Bayes algorithms," *J. Tek. Inform. (JUTIF)*, vol. 6, no. 1, pp. 229–238, Feb. 2025.
- [6] P. Purwanto, M. Merry, Y. Hartanti, M. Sudiarman, dan E. Marlina, "Kebijakan Pemerintah Terhadap Digitalisasi Layanan Kesehatan Primer: Studi Kasus Platform SATUSEHAT", *JUMPA BHAKTI*, vol. 1, no. 2, hlm. 70–79, Agu 2025.
- [7] R. Verma and R. Singh, "Key Aspects of Sentiment Analysis," Jun. 2025, doi: 10.5281/zenodo.15666991.
- [8] T. M. Permata Aulia, N. Arifin, and R. Mayasari, "Perbandingan Kernel Support Vector Machine (Svm) Dalam Penerapan Analisis Sentimen

- Vaksinisasi Covid-19”, *SINTECH Journal*, vol. 4, no. 2, pp. 139–145, Oct. 2021.
- [9] R. Kumar, B. Goswami, S. M. Mhatre, and S. Agrawal, “Naive Bayes in Focus: A Thorough Examination of its Algorithmic Foundations and Use Cases,” *International journal of innovative science and research technology*, pp. 2078–2081, Jun. 2024, doi: 10.38124/ijisrt/ijisrt24may1438.
- [10] D. Septiani and I. Isabela, “Analisis Term Frequency Inverse Document Frequency (Tf Idf) Dalam Temu Kembali Informasi Pada Dokumen Teks,” *Jurnal SINTESIA*, vol. 1, no. 2, Mar. 2022.
- [11] Y. A. Singgalen, “Sentiment Classification of S.E.A Aquarium Singapore Reviews through CRISP-DM using DT and SVM with SMOTE,” *Building of Informatics, Technology and Science (BITS)*, vol. 5, no. 3, Dec. 2023, doi: 10.47065/bits.v5i4.4703.
- [12] G. G. S. Putra, W. Swastika, and P. L. T. Irawan, “Perbandingan Particle Swarm Optimization dengan Genetic Algorithm dalam feature selection untuk analisis sentimen pada Permendikbudristek PPKS-LPT,” *Jurnal Edukasi dan Penelitian Informatika (JEPIN)*, vol. 8, no. 3, p. 412, Dec. 2022, doi: 10.26418/jp.v8i3.57300.
- [13] M. Mashuri, A. E. P. Haryanto, and Y. Manik, “Dashboard Pre-Processing Data (DPD) as Data Analysis System with Technological Innovation to Perform Pre-Processing Quantitative Data,” *IPTEK: The Journal of Engineering*, vol. 9, no. 1, p. 35, Apr. 2023, doi: 10.12962/j23378557.v9i1.a13050.
- [14] Y.-L. He, X. Lu, P. Fournier-Viger, and J. Z. Huang, “A novel overlapping minimization SMOTE algorithm for imbalanced classification,” *Sep.* 2024, doi: 10.1631/fitee.2300278.
- [15] A. F. AlShammari, “Implementation of Keyword Extraction using Term Frequency-Inverse Document Frequency (TF-IDF) in Python,” *International Journal of Computer Applications*, Sep. 2023, doi: 10.5120/ijca2023923137.
- [16] H. A. Parhusip, B. Susanto, L. Linawati, S. Trihandaru, and Y. Sardjono, “Pembelajaran Vektor Untuk Klasifikasi Data Pada Bidang,” *Journal on Mathematics Education*, vol. 4, no. 2, Jul. 2020, doi: 10.35706/SJME.V4I2.3515.
- [17] R. Kumar, B. Goswami, S. M. Mhatre, and S. Agrawal, “Naive Bayes in Focus: A Thorough Examination of its Algorithmic Foundations and Use Cases,” *International journal of innovative science and research technology*, pp. 2078–2081, Jun. 2024, doi: 10.38124/ijisrt/ijisrt24may1438.
- [18] E. Prihantoro, F. R. Rakhman, and R. W. Ramadhani, “Digital movement of opinion mobilization: SNA study on #Dirumahaja vs. #Pakaimasker,” *Jurnal ASPIKOM*, vol. 6, no. 1, p. 77, Jan. 2021, doi: 10.24329/aspikom.v6i1.