



PENERAPAN METODE *SUPPORT VECTOR MACHINE* UNTUK ANALISIS SENTIMEN ULASAN APLIKASI *THREADS* DI *GOOGLE PLAY STORE*

Rijald Djonedhi Elimanafe¹, Sapriani Gustina², Selvi Dwi Hartiyani³, Agung Prayogo⁴

¹rijaldelimanafe@gmail.com, ²sagustina@up45.ac.id, ³selvihartiyani@up45.ac.id, ⁴agung.p@up45.ac.id
^{1,2,3,4}Program Studi Teknologi Informasi, Universitas Proklamasi 45 Yogyakarta

Abstrak

Pertumbuhan penggunaan media sosial mendorong meningkatnya jumlah ulasan pengguna yang mencakup pendapat dan pengalaman terhadap suatu aplikasi. Dengan menggunakan algoritma *Support Vector Machine* (SVM), penelitian ini menyelidiki persepsi pengguna aplikasi *Threads* di *Google Play Store*. Data diperoleh melalui teknik *scraping* sebanyak 5.000 ulasan berbahasa Indonesia. Data tersebut kemudian diproses melalui tahap prapemrosesan teks yang mencakup pembersihan data, *case folding*, normalisasi, tokenisasi, *stopword removal*, *stemming*, serta penghapusan data kosong. Proses pelabelan sentimen dilakukan secara semi-otomatis berdasarkan nilai *rating* dan dikelompokkan ke dalam tiga kategori, yaitu positif, netral, dan negatif. Metode *Frequency-Inverse Document Frequency* (TF-IDF) dipakai untuk membangun representasi fitur. Kemudian, klasifikasi dilakukan menggunakan model SVM dengan kernel *Radial Basis Function* (RBF), dan metrik akurasi dan *confusion matrix* digunakan untuk mengevaluasi. Hasil penelitian menunjukkan tingkat akurasi 96,29% pada data pelatihan dan 77,68% pada data pengujian, sehingga menunjukkan bahwa kombinasi TF-IDF dan SVM efektif dalam analisis sentimen ulasan aplikasi berbahasa Indonesia.

Kata kunci: Analisis Sentimen, *Support Vector Machine*, TF-IDF, *Threads*, *Google Play Store*.

Abstract

The rapid growth of social media platforms has increased the number of user reviews containing opinions and experiences about application usage. This study analyzes the sentiment of user reviews of the *Threads* application on the *Google Play Store* using the *Support Vector Machine* (SVM) method. A total of 5,000 Indonesian-language reviews were collected through scraping and processed using text preprocessing techniques, including cleaning, case folding, normalization, tokenization, stopword removal, and stemming. Sentiment labeling was conducted semi-automatically based on user ratings and grouped into positive, neutral, and negative categories. Text features were represented using the *Term Frequency-Inverse Document Frequency* (TF-IDF) method. The SVM model with a *Radial Basis Function* (RBF) kernel was evaluated using accuracy metrics. The results show an accuracy of 96.29% on training data and 77.68% on testing data, indicating that TF-IDF and SVM are effective for sentiment analysis of Indonesian-language application reviews.

Keywords: Sentiment Analysis, *Support Vector Machine*, TF-IDF, *Threads*, *Google Play Store*.

1. Pendahuluan

Kemajuan aplikasi media sosial berbasis teks mengalami peningkatan signifikan seiring dengan kebutuhan pengguna untuk berkomunikasi secara cepat dan terbuka di ruang digital. Salah satu aplikasi yang menarik perhatian publik adalah *Threads*, yaitu panggung media sosial yang terintegrasi dengan Instagram dan dirilis oleh Meta sebagai alternatif komunikasi berbasis teks. Sejak diluncurkan, aplikasi *Threads* memperoleh jumlah unduhan yang besar dan disertai dengan ribuan ulasan pengguna di *Google Play Store*. Tanggapan tersebut memuat beragam opini, baik positif, netral dan juga negatif, yang mencerminkan pengalaman pengguna terhadap fitur, performa, dan kegunaan aplikasi. Informasi ini juga sebagai sumber data yang penting untuk mengevaluasi kualitas pada aplikasi secara objektif berbasis persepsi pengguna [1].

Namun, banyaknya ulasan dalam bentuk teks tidak terstruktur menyebabkan proses analisis secara manual menjadi tidak efisien dan rentan terhadap subjektivitas. Oleh karena itu, diperlukan pendekatan otomatis berbasis *text mining* dan *machine learning* untuk mengklasifikasikan sentimen pengguna secara

sistematis. Analisis sentimen telah banyak diterapkan pada ulasan aplikasi di *Google Play Store* karena mampu memberikan gambaran kuantitatif mengenai tingkat kepuasan pengguna. Sejumlah penelitian sebelumnya menunjukkan bahwa metode *Support Vector Machine (SVM)* memiliki kinerja yang baik dalam klasifikasi sentimen ulasan aplikasi, seperti pada aplikasi *By.U*, *TikTok*, *iPusnas*, dan *SatuSehat* [2], [3], [4], [5]. Hasil penelitian tersebut menunjukkan bahwa *SVM* mampu menangani data teks berdimensi tinggi secara efektif.

Keberhasilan penerapan algoritma *Support Vector Machine (SVM)* dalam proses analisis sentimen sangat dipengaruhi oleh kualitas tahapan *feature extraction* yang digunakan untuk merepresentasikan informasi dari data teks. Data berupa teks pada umumnya termasuk ke dalam kategori data tidak teratur, akibatnya tidak dapat langsung diolah oleh algoritma *machine learning*. Oleh karena itu, diperlukan suatu proses transformasi yang mengubah teks menjadi representasi numerik supaya bisa diolah secara komputasional oleh sistem pembelajaran mesin. Salah satu pendekatan yang banyak digunakan dalam tahap tersebut adalah metode *Term Frequency-Inverse Document Frequency (TF-IDF)*. Teknik ini merupakan metode pembobotan kata yang bekerja dengan mempertimbangkan dua aspek utama, yaitu keseringan timbul suatu kata dalam dokumen tertentu serta tingkat kepentingan kata tersebut terhadap keseluruhan kumpulan dokumen. Dengan mekanisme tersebut, kata-kata yang punya kandungan pemberitahuan tinggi akan memperoleh nilai kualitas yang lebih besar berbanding dengan kata yang kerap muncul timbul rendah merepresentasikan isi dokumen. Pendekatan ini memungkinkan sistem untuk menangkap karakteristik teks secara lebih tepat dan representatif. Sejumlah penelitian terdahulu juga menunjukkan bahwa penggunaan *TF-IDF* yang dikombinasikan dengan algoritma *SVM* mampu memberikan peningkatan kinerja dalam proses klasifikasi sentimen. Kombinasi kedua metode ini terbukti efektif, terutama ketika diterapkan pada analisis ulasan pengguna pada berbagai aplikasi berbasis platform digital yang umumnya memiliki variasi kosakata yang cukup beragam. Dengan representasi fitur yang lebih informatif, proses pemisahan kelas sentimen dapat dilakukan dengan lebih optimal sehingga menghasilkan performa klasifikasi yang lebih baik [2], [6], [7]. Selain metode ekstraksi fitur, pemilihan kernel pada algoritma *SVM* juga memegang peranan penting dalam menentukan hasil klasifikasi. Kernel *Radial Basis Function (RBF)* sering digunakan karena kesanggupan dalam mengurus hubungan nonlinier pada data berdimensi tinggi, sehingga menghasilkan pemisahan kelas yang lebih optimal [8]. Pendekatan ini sejalan dengan prinsip dasar *machine learning* dalam pengolahan data teks, di mana representasi fitur yang tepat dan pemilihan model yang sesuai menjadi faktor utama dalam meningkatkan akurasi dan stabilitas hasil klasifikasi, sebagaimana dijelaskan dalam literatur mutakhir [9], [10], [11].

Meskipun penelitian terkait pengkajian sikap pada penilaian aplikasi di *Google Play Store* telah berkembang cukup pesat, studi yang secara spesifik mengkaji opini pengguna aplikasi *Threads* masih tergolong terbatas. Sebagian besar penelitian sebelumnya lebih banyak menyoroti berbagai platform media sosial dan layanan digital lainnya, seperti *TikTok*, *By.U*, *iPusnas*, dan *SatuSehat*. Akibatnya, karakteristik serta kecenderungan sentimen yang muncul pada ulasan pengguna *Threads* belum banyak memperoleh perhatian dalam kajian ilmiah. Oleh karena itu, penelitian ini dilakukan dengan mengimplementasikan kombinasi metode *Term Frequency-Inverse Document Frequency (TF-IDF)* dan algoritma *Support Vector Machine (SVM)* untuk mengklasifikasikan sentimen pada ulasan berbahasa Indonesia yang berasal dari aplikasi *Threads* di *Google Play Store*. Selain mengevaluasi performa metode yang digunakan, penelitian ini juga bertujuan memberikan pemahaman mengenai pola dan kecenderungan opini pengguna terhadap platform yang masih relatif baru tersebut.

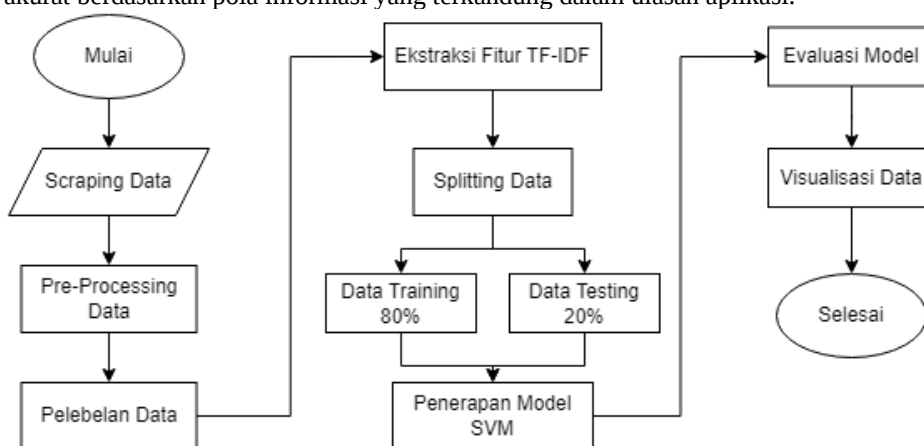
Karakteristik bahasa pada ulasan aplikasi *Threads* juga memiliki keunikan yang membedakannya dari aplikasi lain. Ulasan pengguna umumnya ditulis menggunakan bahasa yang bersifat informal, mengandung singkatan, kata tidak baku, campuran bahasa Indonesia dan Inggris (*code-mixing*), serta penggunaan istilah yang berkembang di media sosial. Selain itu, banyak ulasan yang disampaikan dalam bentuk kalimat singkat, ekspresi spontan, maupun penggunaan emoji yang merepresentasikan emosi pengguna. Variasi linguistik tersebut menyebabkan representasi teks menjadi lebih kompleks dan berpotensi memengaruhi proses klasifikasi sentimen, sehingga diperlukan tahapan prapemrosesan yang tepat serta metode klasifikasi yang mampu menangani data teks berdimensi tinggi.

Atas dasar permasalahan tersebut, penelitian ini dilakukan untuk mengidentifikasi dan memilah sentimen yang tercantum dalam ulasan pemakai aplikasi *Threads* pada *Google Play Store* dengan menggunakan tingkat dari *Term Frequency-Inverse Document Frequency (TF-IDF)* sebagai teknik pembobotan dan representasi fitur, serta algoritma *Support Vector Machine (SVM)* sebagai model klasifikasi. Melalui penelitian ini, diharapkan dapat diperoleh gambaran yang lebih jelas mengenai

tanggapan dan penilaian pengguna terhadap aplikasi *Threads*. Selain itu, hasil yang diperoleh diharapkan mampu memberikan kontribusi terhadap pengembangan kajian analisis sentimen pada platform media sosial modern, sekaligus menjadi sumber informasi yang bermanfaat bagi pengembang dalam melakukan evaluasi dan penyempurnaan kualitas layanan serta pengalaman pengguna.

2. Metode

Penelitian ini menerapkan pendekatan kuantitatif dengan memanfaatkan metode *text mining* untuk menganalisis sentimen pada ulasan pengguna aplikasi *Threads* yang tersaji di *Google Play Store*. Pendekatan tersebut digunakan karena mampu mengolah data teks secara sistematis sehingga pola sentimen yang terkandung di dalam ulasan dapat diidentifikasi secara terukur dan objektif. Proses penelitian terdiri dari beberapa tahapan yang saling berkaitan. Tahap awal penelitian dimulai dengan mengakumulasi data komentar dari pengguna yang kemudian melalui serangkaian proses persiapan untuk memastikan kualitas data sebelum dianalisis lebih lanjut. Setelah data dinyatakan siap, dilakukan penyaringan fitur memanfaatkan metode *Term Frequency-Inverse Document Frequency (TF-IDF)* guna mengganti data teks menjadi representasi numerik yang dapat diolah oleh sistem komputasi. Selanjutnya, proses klasifikasi sentimen dilakukan dengan menerapkan algoritma *Support Vector Machine (SVM)* untuk mengelompokkan ulasan berdasarkan kecenderungan sentimennya. Melalui penerapan kedua metode tersebut, diharapkan dapat dibangun model klasifikasi yang mampu mengenali dan membedakan sentimen pengguna secara akurat berdasarkan pola informasi yang terkandung dalam ulasan aplikasi.



Gambar 1. Diagram Alir Penelitian

2.1. Scraping Data

Data penelitian berupa komentar atau ulasan pengguna aplikasi *Threads* yang didapatkan dari *Google Play Store*. Tahap untuk mengumpulkan data dilakukan dengan memperhatikan kebijakan dan ketentuan pemanfaatan data publik yang ditetapkan oleh *Google Play Store* [1]. Ulasan yang dikumpulkan berbentuk teks opini pengguna yang mencerminkan pengalaman penggunaan aplikasi, baik berupa tanggapan positif maupun negatif. Data yang didapat kemudian disimpan dalam format *dataset* lalu akan diolah pada tahap berikutnya.

2.2. Pre-processing

Sebelum dilakukan proses analisis sentimen, data ulasan terlebih dahulu melalui tahapan persiapan untuk meningkatkan konsistensi dan kualitas data. Tahap ini mempunyai tujuan untuk mengubah teks mentah menjadi format yang lebih terstruktur sehingga lebih mudah untuk diolah pada proses berikutnya. Beberapa teknik yang diterapkan meliputi *case folding*, *tokenization*, *stopword removal*, dan *stemming*. Pada tahap *case folding*, seluruh karakter huruf diubah sebagai huruf kecil untuk menghilangkan perbedaan antara huruf kapital dan nonkapital. Selanjutnya, *tokenization* dilakukan dengan memecah teks menjadi sejumlah token atau kata yang berdiri sendiri sehingga memudahkan proses pengolahan. Tahap berikutnya adalah *stopword removal*, yaitu menghilangkan kata-kata yang banyak muncul namun memiliki kontribusi informasi yang rendah terhadap proses klasifikasi. Setelah itu, dilakukan *stemming* untuk mengonversi kata yang memiliki imbuhan ke bentuk kata dasarnya. Rangkaian proses tersebut berperan penting dalam

meminimalkan keberadaan data yang tidak relevan (*noise*), sehingga representasi fitur yang dihasilkan menjadi lebih informatif dan mampu mendukung peningkatan performa analisis sentimen [9], [11].

2.3. Pelabelan Data

Setelah proses *preprocessing* selesai, data ulasan selanjutnya melalui tahap pelabelan sentimen. Pada tahap ini, setiap ulasan dikelompokkan ke dalam tiga kategori, yaitu sentimen positif, netral, dan negatif, berdasarkan karakteristik opini yang terkandung dalam teks. Hasil pelabelan tersebut kemudian digunakan sebagai acuan dalam proses pelatihan dan pengujian model klasifikasi. Tahap ini berperan penting sebagai dasar dalam proses pelatihan model klasifikasi, karena label sentimen digunakan sebagai target (*class label*) dalam pembelajaran algoritma *machine learning*. Pendekatan ini termasuk semi-otomatis, mengapa demikian, karena mengombinasikan pelabelan otomatis dengan validasi manual. Meski efisien, pelabelan berbasis *Rating* tidak selalu mencerminkan makna sesungguhnya dari teks ulasan. Oleh karena itu, diperlukan verifikasi manual oleh anotator manusia untuk memastikan kesesuaian antara label sentimen dan isi komentar.

2.4. Ekstraksi Fitur TF-IDF

Proses Pemisahan fitur menerapkan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) adalah langkah berikutnya dalam penelitian ini. Metode ini digunakan untuk mengubah data teks yang telah melewati tahap *pre-processing* menjadi representasi numerik yang dapat diproses oleh algoritma komputasi. Metode ini memberi kualitas pada masing-masing kata dalam dokumen berlandaskan frekuensi kemunculannya dan tingkat kepentingannya dalam kumpulan dokumen secara totalitas. Kata-kata yang berkontribusi secara signifikan terhadap perasaan yang digambarkan dalam ulasan akan diprioritaskan oleh mekanisme ini dibandingkan dengan kata-kata yang kurang relevan. Oleh karena itu, metode TF-IDF dinilai efektif dalam menekankan kata-kata yang berperan besar dalam pembentukan sentimen dalam teks, sehingga banyak digunakan dalam beberapa penelitian yang berkaitan dengan analisis sentimen berbasis data teks [2], [6], [7].

2.5. Splitting Data

Setelah tahapan ekstraksi fitur selesai dilakukan, dataset selanjutnya dibagi menjadi dua kelompok, yaitu data pelatihan (*training data*) dan data pengujian (*testing data*). Sebanyak 80% dari total data digunakan untuk membangun dan melatih model klasifikasi, sedangkan 20% sisanya dimanfaatkan untuk mengevaluasi kekuatan pola dalam menggolongkan informasi yang tidak pernah dipahami terlebih dahulu. Pembagian ini dilakukan agar model tidak hanya mampu mempelajari pola yang terdapat pada data pelatihan, tetapi juga dapat diuji tingkat generalisasinya melalui data pengujian. Untuk menjaga kualitas proses pembelajaran, pembagian data dilakukan secara acak menggunakan metode *stratified sampling*, sehingga proporsi setiap kelas sentimen tetap terdistribusi secara seimbang pada kedua kelompok data. Dengan pendekatan tersebut, representasi masing-masing kelas dapat dipertahankan sehingga proses pelatihan dan evaluasi model menjadi lebih reliabel.

2.6. Penerapan Model SVM

Tahap kategorisasi pandangan pada penelitian ini dijalankan dengan menerapkan algoritma *Support Vector Machine* (SVM). Algoritma tersebut bertugas dengan menelusuri dan membentuk *hyperplane* optimal yang sanggup memecah data ke dalam kelas-kelas sentimen yang berbeda berdasarkan karakteristik fitur yang dimiliki. Pemilihan SVM didasarkan pada kapabilitas dalam mengolah data berukuran tinggi serta daya guna dalam menangani permasalahan klasifikasi teks. Berbagai penelitian terdahulu juga menunjukkan bahwa algoritma ini mampu membangun kinerja yang baik dalam proses identifikasi dan pengelompokan sentimen pada data ulasan berbasis teks [3].

Secara matematis, fungsi dasarnya dapat dituliskan sebagai berikut:

$$\text{Fungsi keputusan: } f(x) = \text{sign}(w \cdot x + b) \quad (1)$$

$$\text{Fungsi optimisasi: } \min_{w,b} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n \xi_i \quad (2)$$

Fungsi keputusan hanya menampilkan bentuk hubungan input dan output, sedangkan fungsi optimisasi menentukan bagaimana w dan b dicari menggunakan parameter C dan Gamma , dengan

meminimalkan kesalahan dan memaksimalkan margin, di mana w merupakan vektor bobot dan b adalah bias. Parameter C mengatur keseimbangan antara kesalahan klasifikasi dan lebar margin, sedangkan Γ mengontrol pengaruh setiap titik data terhadap *Decision Boundary* [8].

2.7. Evaluasi Model

Setelah model klasifikasi berhasil dibangun, tahap berikutnya adalah melakukan evaluasi untuk mengukur tingkat kinerjanya dalam mengelompokkan sentimen pada data ulasan. Evaluasi dilakukan menggunakan beberapa metrik pengukuran, yaitu akurasi, *precision*, *recall*, dan *confusion matrix*. Penggunaan metrik tersebut memungkinkan penilaian yang lebih komprehensif terhadap kemampuan model dalam mengenali serta membedakan setiap kategori sentimen. Melalui hasil evaluasi ini, dapat diketahui tingkat efektivitas algoritma *Support Vector Machine* (SVM) dalam melakukan klasifikasi sentimen ulasan pengguna aplikasi *Threads* secara konsisten dan akurat.

1. Akurasi (*Accuracy*): rasio antara prediksi benar dengan total data.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (3)$$

2. Presisi (*Precision*): proporsi data yang benar-benar positif dari seluruh prediksi positif.

$$Precision = \frac{TP}{TP+FP} \quad (4)$$

3. *Recall* (*Sensitivity*): rasio data positif yang berhasil dikenali dengan benar.

$$Recall = \frac{TP}{TP+FN} \quad (5)$$

4. *F1-Score*: rata-rata harmonis antara *precision* dan *recall*.

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (6)$$

2.8. Visualisasi Data

Tahap terakhir dalam metode penelitian ini adalah proses visualisasi data, yang bertujuan untuk menyediakan perolehan penelaahan sentimen dalam wujud tampilan grafis sehingga lebih gampang dimengerti dan diinterpretasikan. Visualisasi dilakukan dengan menampilkan hasil klasifikasi dalam bentuk grafik atau diagram yang mampu menggambarkan pola serta kecenderungan sentimen yang muncul dari ulasan pengguna. Melalui penyajian visual tersebut, distribusi sentimen pengguna terhadap aplikasi *Threads* dapat terlihat dengan lebih jelas berdasarkan hasil klasifikasi yang telah diperoleh sebelumnya.

3. Hasil dan Pembahasan

Fase ini menyampaikan hasil evaluasi sikap terhadap penilaian pemakai aplikasi *Threads* pada *Google Play Store* serta pembahasan terhadap temuan yang diperoleh. Hasil penelitian diperoleh melalui tahapan *scraping* data, *pre-processing*, pelabelan data, ekstraksi fitur menggunakan *TF-IDF*, penerapan model *Support Vector Machine* (SVM), dan evaluasi model sesuai dengan metode penelitian yang telah dijelaskan sebelumnya.

Tabel 1. Contoh Data Ulasan Aplikasi *Threads*.

Username	Rating	Review	Date	Thumbs_Up
Ivi Yanti Fitaloka	1	GK ada guna nya	2025-11-28 16:29:13	0
Yosia Vemi	1	loadingnya lama banget padahal IG dan FB lancar	2025-11-28 13:07:30	0
Wan Ardiansyah	5	bagus. saya menyukai Threads	2025-11-28 11:52:24	0
pirusandie	5	sangat berguna dan baik	2025-11-28 11:41:16	0
Aranta Totota	5	tu admin, saya tidak bisa membalas utas orang lain	2025-11-28 11:34:23	0

Tabel ini berfungsi untuk memberikan gambaran umum mengenai bentuk data teks yang dianalisis serta hasil pelabelan sentimen sebelum dilakukan proses klasifikasi.

3.1. Karakteristik Data Ulasan

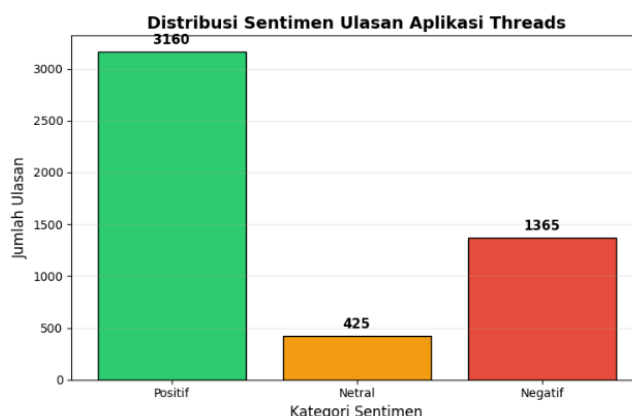
Studi ini menggunakan ulasan pengguna dari aplikasi *Threads* yang diperoleh dari *Google Play Store*. Data siap untuk tahap pelabelan dan klasifikasi sentimen setelah proses scraping dan *pre-processing* selesai. Tanggapan pengguna terhadap aplikasi dapat dilihat melalui ulasan yang dikumpulkan, yang menunjukkan kesan positif dan negatif.

Data yang dikumpulkan melalui proses pengambilan data otomatis yang dikenal sebagai *scraping*, termasuk ulasan pengguna aplikasi *Threads*, digunakan dalam penelitian ini. Metode ini dipilih karena memungkinkan pengumpulan data dalam volume yang lebih besar secara lebih cepat dan efisien daripada proses pengambilan data secara manual. Data ulasan dapat dikumpulkan secara sistematis dari sumber online yang tersedia dengan teknik ini. Proses scraping dilakukan pada halaman ulasan aplikasi *Threads* di *Google Play Store*. Ini menyaring ulasan dalam bahasa Indonesia sehingga hanya ulasan dalam bahasa ini yang digunakan sebagai data penelitian. Skrip dijalankan dengan pengaturan pengambilan ulasan terbaru (ulasan terkini), agar data yang diperoleh benar-benar mencerminkan kondisi terkini. Target jumlah data ditetapkan sebanyak 5.000 ulasan, dan proses pengambilan berhasil mencapai jumlah tersebut tanpa adanya kegagalan atau kehilangan data.

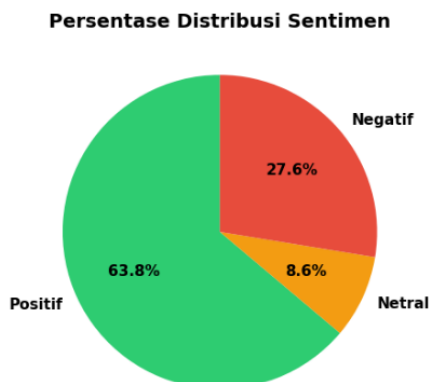
Berdasarkan keseluruhan data yang berhasil dikumpulkan, terdapat variasi penilaian yang cukup mencolok. Nilai *rating* 1 berjumlah 1.105 ulasan, *rating* 2 berjumlah 285 ulasan, *rating* 3 berjumlah 431 ulasan, *rating* 4 berjumlah 555 ulasan, dan *rating* 5 merupakan jumlah terbesar yaitu 2.624 ulasan. Perbedaan jumlah ini menunjukkan bahwa persepsi pengguna terhadap aplikasi *Threads* cukup beragam, mulai dari ulasan yang bersifat sangat negatif hingga sangat positif. Namun demikian, pada tahap hasil ini belum dibahas mengenai kecenderungan atau interpretasi sentimen, karena bagian ini hanya memaparkan kondisi dataset apa adanya.

3.2. Distribusi Sentimen Ulasan

Berdasarkan hasil proses pelabelan data, ulasan pengguna aplikasi *Threads* dibagi menjadi dua kelompok utama sentimen: sentimen positif dan sentimen negatif. Pembagian ke dalam kelompok-kelompok ini menunjukkan kecenderungan pendapat atau persepsi pengguna tentang aplikasi yang sedang diuji. Dengan mengetahui distribusi sentimen yang terbentuk, peneliti dapat memahami tingkat kepuasan maupun ketidakpuasan pengguna secara umum terhadap layanan aplikasi. Analisis terhadap distribusi sentimen ini menjadi tahap awal yang penting sebelum dilakukan proses klasifikasi menggunakan algoritma *machine learning*, karena dapat memberikan pemahaman dasar mengenai pola opini yang terdapat dalam data ulasan pengguna.



Gambar 2. Distribusi Sentimen Ulasan Aplikasi *Threads*.



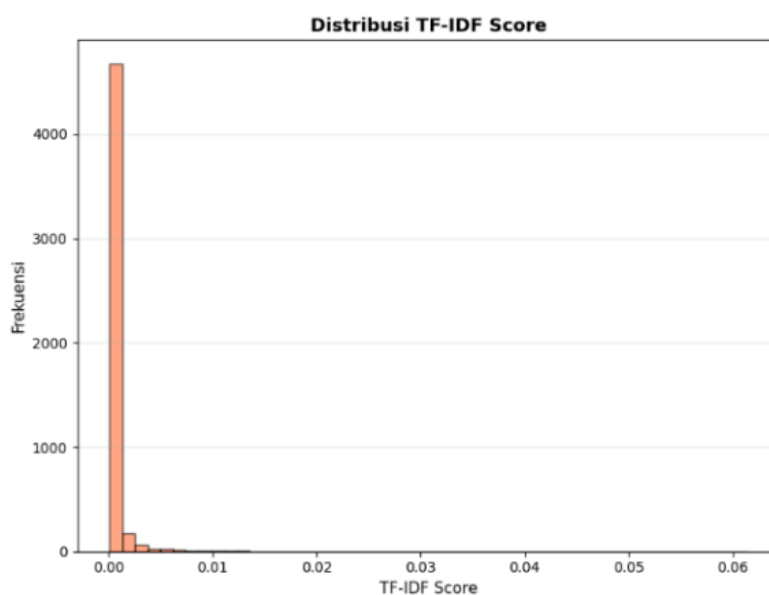
Gambar 3. Presentase Distribusi Sentimen.

Hasil distribusi sentimen yang ditunjukkan pada Gambar 2 dan Gambar 3 memperlihatkan pola yang konsisten dengan temuan pada berbagai penelitian terdahulu. Studi-studi tersebut mengungkapkan bahwa ulasan pengguna pada aplikasi yang ada di *Google Play Store* umumnya cenderung didominasi oleh kategori sentimen tertentu, yang dipengaruhi oleh kualitas layanan, stabilitas sistem, serta fitur-fitur yang ditawarkan oleh aplikasi.

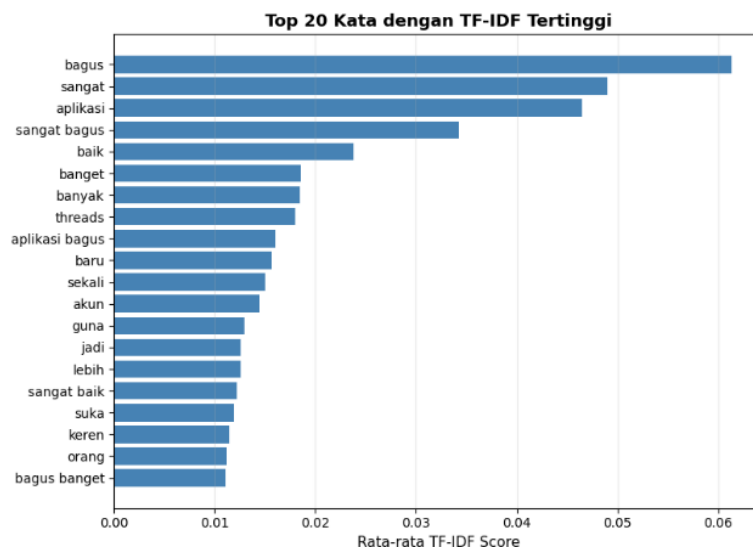
3.3. Hasil Ekstraksi Fitur *TF-IDF*

Ekstraksi fitur dilakukan terhadap 4.950 ulasan yang telah bersih dan siap diproses. Hasil transformasi berupa matriks sparse berukuran 4.950 baris dan 5.000 kolom, yang menunjukkan bahwa setiap dokumen direpresentasikan oleh 5.000 fitur unik yang menunjukkan bahwa sebagian besar nilai dalam matriks bernilai nol, hal ini merupakan karakteristik umum dari representasi teks menggunakan *TF-IDF*. Dengan jumlah fitur sebesar 5.000, model memiliki ruang representasi yang cukup luas untuk menangkap variasi kata yang muncul dalam ulasan.

Untuk memperjelas hasil ekstraksi fitur, visualisasi yang dihasilkan menunjukkan dua aspek utama: pertama, daftar dua puluh kata dengan nilai *TF-IDF* tertinggi yang menggambarkan kata-kata paling dominan dalam penilaian; kedua, distribusi keseluruhan nilai *TF-IDF* pada seluruh dokumen yang menunjukkan bahwa sebagian besar nilai berada pada kisaran rendah.



Gambar 4. Distribusi *TF-IDF* Score

Gambar 5. Top 20 Kata *TF-IDF* Tertinggi

Secara umum, proses ekstraksi fitur dengan metode *Term Frequency–Inverse Document Frequency* (*TF-IDF*) mampu mengonversi data ulasan yang semula berbentuk teks menjadi representasi numerik yang dapat diolah oleh model *machine learning*. Hasil transformasi tersebut menyediakan informasi yang relevan mengenai karakteristik setiap ulasan sehingga dapat dimanfaatkan pada tahap pembagian data maupun proses pelatihan model. Representasi fitur yang dihasilkan juga berperan penting dalam membantu algoritma klasifikasi mengidentifikasi pola dan kecenderungan sentimen yang terkandung dalam teks secara lebih efektif.

3.4. Hasil Klasifikasi Menggunakan *Support Vector Machine*

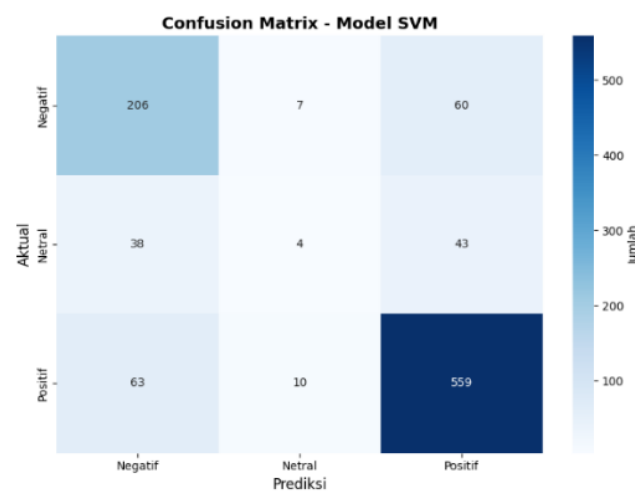
Pada penelitian ini, proses pengelompokan sentimen dilakukan dengan menerapkan algoritma *Support Vector Machine* (*SVM*) yang menggunakan fitur hasil transformasi *Term Frequency–Inverse Document Frequency* (*TF-IDF*) sebagai representasi data teks. Sebelum digunakan untuk melakukan prediksi, model terlebih dahulu melalui tahap pelatihan menggunakan *training data* agar mampu mengenali karakteristik serta pola sentimen yang terdapat pada kumpulan ulasan pengguna. Proses pembelajaran tersebut memungkinkan model membangun aturan klasifikasi yang dapat digunakan untuk membedakan kategori sentimen berdasarkan informasi yang terkandung dalam teks ulasan. Sesudah itu, pola yang telah dilatih diuji memakai data *testing* untuk menilai kemampuannya dalam mengklasifikasikan sentimen pada data yang belum pernah diproses sebelumnya. Berdasarkan hasil pengujian yang dilakukan, kombinasi metode *TF-IDF* dan algoritma *SVM* menunjukkan kemampuan yang baik dalam memisahkan ulasan dengan sentimen positif dan negatif secara efektif.

Pada tahap ini dilakukan proses pelatihan model dengan menggunakan algoritma *Support Vector Machine* (*SVM*) yang menerapkan fungsi *kernel Radial Basis Function* (*RBF*). Pemilihan *kernel RBF* didasarkan pada kemampuannya dalam memetakan data ke dalam ruang berdimensi yang lebih tinggi, sehingga model dapat mengenali pola hubungan yang bersifat nonlinier. Karakteristik ini sangat penting dalam pengolahan data teks, karena pola yang muncul pada data ulasan sering kali tidak mengikuti hubungan linier yang sederhana. Dengan memanfaatkan *kernel RBF*, model diharapkan mampu menangkap pola sentimen yang lebih kompleks sehingga kinerja sistem analisis sentimen dapat menjadi lebih optimal. Pelatihan dimulai dengan menyiapkan data *training* sebanyak 3960 sampel dan 5000 fitur hasil ekstraksi *TF-IDF*. Model *SVM* kemudian diinisialisasi menggunakan beberapa parameter penting, yaitu parameter *C* sebagai pengendali regulasi dan *gamma* sebagai pengatur pengaruh suatu titik data terhadap ruang keputusan. Selain itu, model juga diberi pengaturan *class weight balanced* untuk meminimalkan pengaruh ketidakseimbangan distribusi kelas.

Secara lebih rinci, proses pelatihan model klasifikasi sentimen pada penelitian ini dilakukan dengan memanfaatkan algoritma *Support Vector Machine* (*SVM*) yang menggunakan *kernel Radial Basis Function* (*RBF*). Dalam penerapannya, model dikonfigurasi dengan parameter *C* sebesar 1.0 yang berfungsi untuk

mengatur tingkat kompromi antara margin pemisah dan kesalahan klasifikasi pada data pelatihan. Selain itu, parameter *gamma* ditetapkan pada nilai *scale*, di mana nilainya dihitung secara otomatis berdasarkan variansi fitur yang terdapat pada data sehingga model dapat menyesuaikan tingkat pengaruh setiap titik data dalam proses pembelajaran. Untuk mengatasi potensi ketidakseimbangan jumlah data pada masing-masing kelas sentimen, digunakan pengaturan *class weight* dengan nilai *balanced*, sehingga model dapat memberikan bobot yang proporsional terhadap setiap kelas. Di samping itu, parameter *random state* ditetapkan pada nilai 42 dengan tujuan menjaga konsistensi hasil pelatihan model ketika proses dijalankan kembali. Pada tahap ekstraksi fitur, metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan dan menghasilkan sebanyak 5000 fitur yang merepresentasikan karakteristik teks. Sementara itu, jumlah data yang dimanfaatkan pada tahap pelatihan model terdiri dari 3960 ulasan pengguna.

Proses pelatihan memerlukan waktu komputasi sebesar 3,73 detik. Setelah proses selesai, model menghasilkan akurasi sebesar 96,29% pada data *training*. Nilai ini menerangkan bahwa model sanggup mendalami pola klasifikasi sikap dengan benar. Secara umum, hasil *training* memperlihatkan bahwa model SVM dengan kernel *RBF* bekerja efektif dalam membangun fungsi klasifikasi untuk data ulasan aplikasi, namun tetap memerlukan tahapan evaluasi lebih lanjut untuk mengetahui stabilitas performanya ketika diuji pada data yang tidak pernah dilihat sebelumnya.

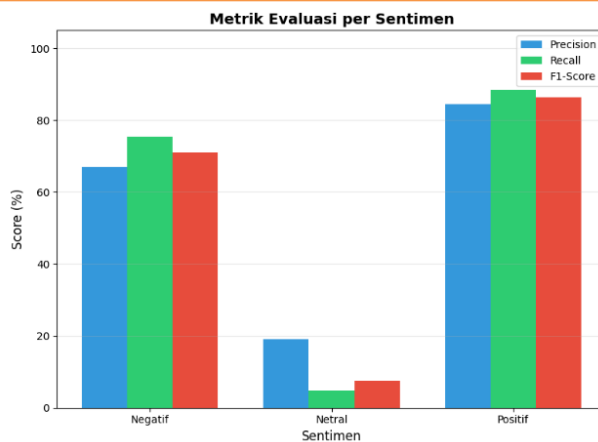


Gambar 6. *Confusion Matrix* Model SVM

Gambar 6 menjelaskan bahwa jumlah prediksi yang salah dan benar untuk masing-masing kelas sentimen ditunjukkan dalam *Confusion Matrix*. Visualisasi ini membantu dalam memahami performa model secara lebih rinci. Hasil *confusion matrix* menunjukkan seberapa sering model memprediksi kelas dengan benar maupun salah. Model cukup baik untuk mengidentifikasi perasaan positif, namun mengalami kesulitan dalam membedakan sentimen netral, yang tercermin dari rendahnya nilai *recall* pada kelas tersebut. Sebaliknya, performa untuk sentimen negatif berada pada tingkat menengah dengan nilai *precision* dan *recall* yang lebih stabil.

3.5. Evaluasi Performa Model

Kinerja model dievaluasi menggunakan beberapa metrik pengukuran, yaitu akurasi, *precision*, dan *recall*. Metrik akurasi dipakai untuk mengukur perbandingan perkiraan yang sesuai pada seluruh data yang diuji. Sementara itu, *precision* menggambarkan tingkat ketepatan model dalam memberikan prediksi pada setiap kategori sentimen, sedangkan *recall* menunjukkan kesanggupan model untuk menemukan data yang benar-benar termasuk dalam kelas sentimen. Berdasarkan hasil evaluasi yang diperoleh, algoritma *Support Vector Machine* (SVM) menunjukkan kemampuan yang baik dalam mengklasifikasikan sentimen pada ulasan pengguna aplikasi *Threads*.



Gambar 7. Metrik Evaluasi per Sentimen

Berdasarkan hasil evaluasi yang ditampilkan pada Gambar 7, performa model dalam mengklasifikasikan sentimen menunjukkan variasi pada setiap kategori. Kelas sentimen negatif memperoleh nilai *precision* sebesar 67,10%, *recall* 75,46%, dan *F1-score* 71,03%. Hasil tersebut menunjukkan bahwa model telah mampu mengidentifikasi ulasan bernada negatif dengan cukup baik, meskipun masih ditemukan sejumlah kesalahan klasifikasi. Pada kategori sentimen netral, nilai *precision* tercatat sebesar 19,05%, *recall* 4,71%, dan *F1-score* 7,55%. Nilai yang relatif rendah ini mengindikasikan bahwa model belum mampu mengenali ulasan netral secara optimal. Kondisi tersebut diduga dipengaruhi oleh lebih sedikit data netral dibandingkan kelas lainnya serta adanya kemiripan karakteristik linguistik dengan sentimen positif maupun negatif.

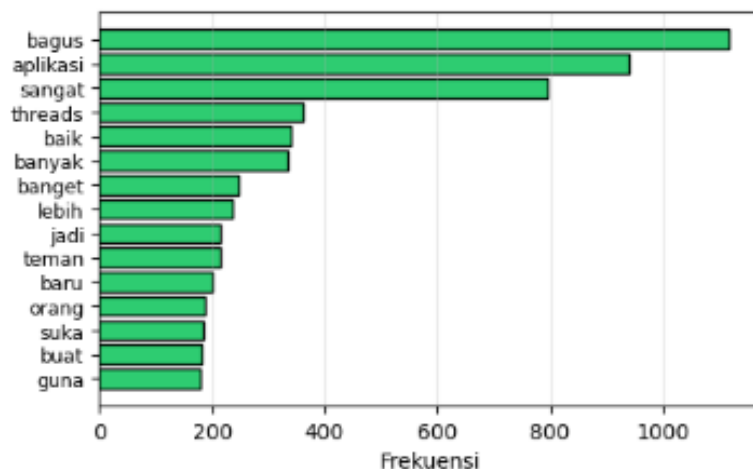
Di sisi lain, kategori sentimen positif menghasilkan kinerja yang paling baik dengan nilai *precision* sebesar 84,44%, *recall* 88,45%, dan *F1-score* 86,40%. Temuan ini menyatakan bahwa pola memiliki kemampuan yang tinggi dalam mengenali pola bahasa yang merepresentasikan opini positif pengguna. Jika ditinjau secara keseluruhan, nilai *weighted precision* sebesar 74,04%, *weighted recall* 77,68%, dan *weighted F1-score* 75,39% menunjukkan bahwa model *Support Vector Machine* (SVM) yang memanfaatkan fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) mampu memberikan hasil klasifikasi yang cukup memadai pada ulasan aplikasi *Threads* di *Google Play Store*.

Untuk memperjelas hasil evaluasi, kinerja model disajikan dalam bentuk *confusion matrix* serta grafik perbandingan nilai *precision*, *recall*, dan *F1-score* pada masing-masing kelas sentimen. Visualisasi tersebut memperlihatkan bahwa kategori sentimen positif memiliki tingkat performa tertinggi, diikuti oleh sentimen negatif, sedangkan sentimen netral masih menunjukkan hasil yang paling rendah. Fenomena ini umumnya terjadi pada dataset yang memiliki distribusi kelas tidak seimbang atau ketika suatu kategori memiliki karakteristik kata yang kurang spesifik sehingga lebih sulit dibedakan oleh model klasifikasi.

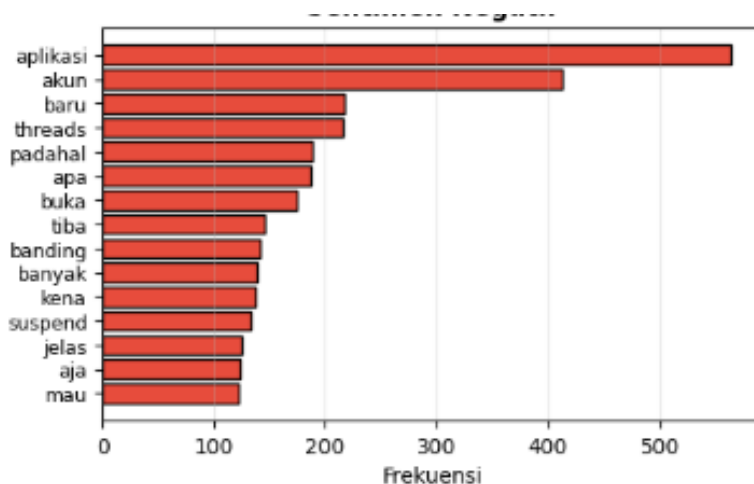
Meskipun demikian, hasil evaluasi menunjukkan bahwa model SVM dengan kernel *RBF* telah dapat beroperasi dengan cukup baik, terutama untuk sentimen positif dan negatif. Model masih dapat ditingkatkan melalui hyperparameter tuning, penyeimbangan data, atau peningkatan kualitas tahapan pra-proses untuk mengurangi ambiguitas pada sentimen netral.

[illegible]

Wordcloud pada kategori kata-kata yang paling digunakan dalam sentimen positif, di antaranya seperti aplikasi, bagus, sangat, baik, dan mudah. Kemunculan kata-kata tersebut dengan frekuensi yang tinggi menandakan adanya kecenderungan pemakai untuk menyampaikan penilaian yang bersifat apresiatif terhadap aplikasi yang digunakan. Dominasi istilah tersebut juga mengindikasikan bahwa banyak pengguna menilai aplikasi memiliki kualitas yang baik, mudah digunakan, serta mampu memberikan pengalaman penggunaan yang menyenangkan ketika mengakses berbagai fitur yang tersedia. Pada kategori netral, kata yang banyak muncul antara lain *threads*, baru, akun, posting, dan komentar. Pola ini menunjukkan bahwa ulasan dengan sentimen netral umumnya berisi pernyataan informatif, deskriptif, atau naratif tanpa memberikan penilaian emosional yang jelas terhadap aplikasi. *Wordcloud* sentimen negatif menampilkan dominasi kata seperti aplikasi, akun, buka, padahal, *login*, dan masalah. Kemunculan istilah tersebut mencerminkan bahwa ulasan negatif umumnya berkaitan dengan keluhan teknis, kesulitan akses, permasalahan *login*, serta kendala fungsional. *Wordcloud* pada setiap kategori memberikan pemahaman tambahan mengenai konteks penggunaan bahasa dan aspek aplikasi yang paling sering menjadi fokus pengguna.



Gambar 10. Top 15 Kata Berdasarkan Frekuensi Kemunculan pada Sentimen Positif



Gambar 11. Top 15 Kata Berdasarkan Frekuensi Kemunculan pada Sentimen Negatif

Untuk melengkapi interpretasi visualisasi *wordcloud*, penelitian ini juga menyajikan tiga puluh kata dengan frekuensi kemunculan tertinggi, yang terdiri atas lima belas kata pada sentimen positif dan lima belas kata pada sentimen negatif. Penyajian frekuensi kata ini bertujuan untuk memberikan informasi kuantitatif mengenai kosakata yang paling dominan muncul pada masing-masing kategori sentimen setelah seluruh tahapan *pre-processing* selesai dilakukan.

Berdasarkan hasil analisis, kata-kata dengan sentimen positif paling sering muncul antara lain bagus, aplikasi, sangat, *threads*, baik, banyak, banget, lebih, jadi, teman, baru, orang, suka, buat, dan guna. Dominasi kata-kata tersebut menunjukkan bahwa sebagian besar pengguna memberikan apresiasi terhadap kualitas aplikasi, kemudahan penggunaan, serta pengalaman yang diperoleh selama menggunakan *Threads*. Kehadiran kata seperti bagus, baik, dan suka mengindikasikan adanya kecenderungan pengguna untuk memberikan penilaian yang positif terhadap fitur maupun layanan yang tersedia.

Sementara itu, pada opini buruk, kata-kata yang memiliki frekuensi kemunculan tertinggi meliputi aplikasi, akun, baru, *threads*, padahal, apa, buka, tiba, banding, banyak, kena, suspend, jelas, aja, dan mau. Kemunculan kata-kata tersebut menunjukkan bahwa ulasan negatif terkait dengan masalah teknis lebih banyak, seperti permasalahan akses akun, kesulitan membuka aplikasi, serta isu yang berhubungan dengan pembatasan akun (*suspend*). Selain itu, munculnya kata padahal dan apa mengindikasikan adanya ekspresi kekecewaan atau ketidakpuasan pengguna terhadap pengalaman yang mereka alami.

Hasil analisis frekuensi kata tersebut sejalan dengan visualisasi *wordcloud*, di mana kata yang memiliki lebih banyak kemunculan ditampilkan dalam ukuran yang lebih besar. Dengan demikian,

kombinasi antara *wordcloud* dan grafik frekuensi kata memberikan gambaran yang lebih komprehensif mengenai karakteristik kosakata yang mendominasi masing-masing kategori sentimen. Penyajian ini tidak hanya memudahkan interpretasi visual, tetapi juga memberikan dasar kuantitatif yang memperkuat analisis terhadap pola opini pengguna aplikasi *Threads* pada *Google Play Store*.

3.6. Pembahasan Hasil

Hasil penelitian ini menunjukkan bahwa penerapan metode *Term Frequency–Inverse Document Frequency* (TF-IDF) yang dipadukan dengan algoritma *Support Vector Machine* (SVM) mampu menghasilkan kinerja klasifikasi pandangan yang cukup baik pada penilaian konsumen aplikasi *Threads* di *Google Play Store*. Tingginya performa model pada kategori sentimen positif mengindikasikan bahwa pola bahasa yang digunakan dalam ulasan bernada positif cenderung lebih konsisten, sehingga lebih mudah dikenali dan dipelajari oleh model klasifikasi. Temuan ini sejalan dengan berbagai penelitian sebelumnya dalam bidang analisis sentimen aplikasi digital, yang juga membuktikan bahwa pola klasifikasi sering kali punya kinerja yang lebih dominan pada kelas sentimen positif dibandingkan dengan kelas lainnya [2], [3], [5]. Kondisi tersebut mengindikasikan bahwa metode TF-IDF memiliki kemampuan yang cukup efektif dalam merepresentasikan kata-kata yang mempunyai arti penting dan sering muncul dalam komentar yang bersifat apresiatif, sehingga membantu model dalam mengidentifikasi pola sentimen secara lebih akurat [7].

Sebaliknya, performa model pada kategori sentimen negatif yang relatif lebih rendah dibandingkan dengan sentimen positif menunjukkan adanya kemungkinan tumpang tindih karakteristik teks antara ulasan negatif dan kategori lainnya. Kondisi tersebut dapat terjadi karena ungkapan dalam ulasan bernada negatif sering kali disampaikan dengan variasi kosakata yang lebih beragam serta konteks kalimat yang lebih kompleks. Temuan ini sejalan dengan beberapa penelitian sebelumnya bahwa ulasan negatif umumnya memiliki keragaman bahasa yang lebih luas, sehingga proses pemisahan kelas oleh model *Support Vector Machine* (SVM) menjadi lebih menantang [4], [8]. Di sisi lain, penggunaan *kernel Radial Basis Function* (RBF) dalam model SVM turut memberikan kontribusi penting karena kemampuannya dalam menangkap hubungan pola yang bersifat nonlinier pada data teks berdimensi tinggi, sebagaimana dijelaskan dalam berbagai kajian mengenai performa SVM pada tugas klasifikasi teks [8].

Performa yang paling rendah pada kategori sentimen netral menunjukkan bahwa model cukup mengalami kesulitan dalam membedakan komentar yang bersifat netral dari ulasan yang bernada positif maupun negatif. Kondisi ini dapat terjadi karena karakteristik teks netral umumnya lebih singkat serta tidak banyak mengandung kata-kata yang mengekspresikan emosi secara jelas. Akibatnya, pola linguistik pada ulasan netral menjadi kurang menonjol sehingga lebih sulit dikenali oleh model klasifikasi. Temuan ini sejalan dengan beberapa penelitian dalam bidang analisis sentimen ulasan aplikasi berbasis *machine learning* yang menyebutkan bahwa kelas netral sering kali menjadi kategori yang paling sulit diprediksi secara akurat [6], [7]. Selain itu, meskipun pengaturan *class weight* dengan nilai *balanced* telah diterapkan untuk mengurangi pengaruh ketidakseimbangan data, jumlah data netral yang relatif lebih sedikit tetap menjadi faktor yang membatasi kemampuan model dalam mempelajari pola sentimen yang benar-benar representatif.

Secara keseluruhan, nilai *weighted precision*, *weighted recall*, dan *weighted F1-score* memaparkan bahwa model *Support Vector Machine* (SVM) yang memanfaatkan fitur *Term Frequency–Inverse Document Frequency* (TF-IDF) mampu menghasilkan kinerja yang relatif stabil serta kompetitif apabila dibandingkan dengan berbagai penelitian sebelumnya yang mengkaji ulasan aplikasi pada *Google Play Store* [2], [3], [5]. Hasil tersebut mengindikasikan bahwa pendekatan yang menggabungkan metode TF-IDF dengan algoritma SVM masih menjadi strategi yang efektif dalam melakukan analisis sentimen berbasis data teks. Temuan ini sekaligus memperkuat pandangan bahwa kombinasi kedua metode tersebut relevan digunakan untuk mengidentifikasi dan mengevaluasi persepsi pengguna terhadap aplikasi *Threads* sebagai salah satu platform media sosial berbasis teks.

4. Kesimpulan dan Saran

4.1. Kesimpulan

Penelitian ini berhasil mengkaji pendapat pengguna tentang aplikasi *Threads* yang tersedia di *Google Play Store* dengan memanfaatkan kombinasi metode *Term Frequency–Inverse Document Frequency* (TF-IDF) dan algoritma *Support Vector Machine* (SVM). Berdasarkan hasil analisis yang diperoleh, kategori sentimen positif menjadi sentimen yang paling dominan dalam ulasan pengguna. Model

klasifikasi yang dibangun menggunakan kernel *Radial Basis Function (RBF)* menghasilkan tingkat akurasi sebesar 96,29% pada data pelatihan dan 77,68% pada data pengujian. Hasil tersebut menunjukkan bahwa penerapan *TF-IDF* sebagai metode perwakilan atribut dan *SVM* selaku langkah-langkah klasifikasi mampu mengidentifikasi sentimen pada ulasan berbahasa Indonesia dengan performa yang cukup baik. Meskipun demikian, perbedaan nilai akurasi antara data pelatihan dan data pengujian mengindikasikan bahwa masih terdapat ruang untuk meningkatkan kemampuan generalisasi model pada data yang belum pernah diproses sebelumnya.

Temuan penelitian ini menunjukkan bahwa metode ini dapat digunakan sebagai dasar dalam peningkatan sistem penggalian opini untuk mengevaluasi resensi aplikasi media sosial, karena penerapan kombinasi *TF-IDF* dan *SVM* untuk menganalisis ulasan aplikasi *Threads*, yang masih relatif sedikit dikaji dibandingkan dengan aplikasi media sosial lainnya.. Selain itu, penelitian ini memberikan gambaran tentang karakteristik sentimen pengguna berdasarkan data ulasan terbaru di *Google Play Store*.

Penelitian ini juga memberikan informasi yang dapat dipakai oleh pengembang aplikasi untuk mengidentifikasi kecenderungan opini pengguna sebagai bahan evaluasi untuk meningkatkan kualitas layanan dan pengalaman pengguna. Selain itu, temuan penelitian ini dapat menjadi acuan untuk pengembangan penelitian tentang analisis sentimen pada aplikasi media sosial yang memiliki karakteristik bahasa yang beragam dan dinamis.

4.2. Saran

Berdasarkan hasil penelitian yang telah diwujudkan, beberapa saran untuk penelitian yang akan dilanjutkan adalah sebagai berikut.

1. Meningkatkan keseimbangan distribusi data pada setiap kelas sentimen, terutama pada kelas sentimen netral, sehingga model dapat mempelajari karakteristik setiap kelas secara lebih representatif dan menghasilkan performa klasifikasi yang lebih seimbang.
2. Mengembangkan proses *pre-processing* dengan memperkaya kamus normalisasi bahasa Indonesia, termasuk penanganan bahasa tidak baku, singkatan, bahasa gaul, *code-mixing*, serta bahasa yang berbeda yang digunakan di media sosial, sehingga kualitas representasi teks dapat ditingkatkan.
3. Membandingkan metode yang digunakan dengan pendekatan lain, seperti *word embedding (Word2Vec, FastText, atau GloVe)* maupun model berbasis transformer (*IndoBERT* atau *IndoBERTweet*), untuk menyadari metode yang paling ampuh dalam mengelompokkan pandangan ulasan berbahasa Indonesia.
4. Mengembangkan model analisis sentimen menjadi sistem pemantauan opini pengguna secara otomatis sehingga hasil analisis dapat dimanfaatkan oleh pengembang aplikasi *Threads* untuk mengidentifikasi permasalahan yang sering dilaporkan pengguna, mengevaluasi kualitas layanan, serta memantau perubahan persepsi pengguna pada *Google Play Store* secara berkala.
5. Melakukan pembaruan dataset dan pelatihan ulang model secara berkala mengingat karakteristik bahasa pada media sosial terus berkembang seiring munculnya kosakata, istilah, dan pola komunikasi baru yang dapat memengaruhi kinerja model klasifikasi.

Daftar Pustaka

- [1] G. Developers, "Google Play Store Review Dataset and Policies," 2024. [Online]. Available: <https://developer.android.com>
- [2] S. Fransiska, R. Rianto, and A. I. Gufroni, "Sentiment Analysis Provider By.U on Google Play Store Reviews with TF-IDF and Support Vector Machine (SVM) Method," *UNP J. Sustain. Data Sci.*, vol. 1, no. 2, pp. 45–53, 2020.
- [3] N. P. Husain, S. Sukirman, and S. Sajiah, "Analisis Sentimen Ulasan Pengguna TikTok pada Google Play Store Berbasis TF-IDF dan Support Vector Machine," *J. Syst. Comput. Eng.*, vol. 5, no. 2, pp. 88–97, 2024.
- [4] R. Lestari, A. Wijaya, and D. Prasetyo, "Sentiment Analysis of iPusnas Application Reviews on Google Play Using Support Vector Machine," *J. Smart Comput. Eng.*, vol. 3, no. 1, pp. 25–33, 2022.
- [5] M. Imanuddin, S. Adi, and R. Gernowo, "Sentiment Analysis on SatuSehat Application Using Support Vector Machine," *J. Inf. Syst. Res.*, vol. 6, no. 3, pp. 110–118, 2023.

-
- [6] M. Sakhdiah, A. Salma, and D. Permana, "Sentiment Analysis Using Support Vector Machine (SVM) of ChatGPT Application Users in Play Store," *Indones. J. Technol. Innov.*, vol. 4, no. 1, pp. 12–20, 2024.
 - [7] P. Guleria, J. Frnda, and P. N. Srinivasu, "Improving Text Classification Performance Using TF-IDF and Machine Learning Algorithms," *Appl. Sci.*, vol. 15, no. 2, pp. 1–14, 2025.
 - [8] A. Tanveer, M. Z. Khan, and S. Ahmad, "Performance Analysis of Support Vector Machine with RBF Kernel for Text Classification," *Int. J. Adv. Comput. Sci. Appl.*, vol. 13, no. 4, pp. 421–428, 2022.
 - [9] E. Alpaydin, *Introduction to Machine Learning*, 4th ed. Cambridge, MA, USA: MIT Press, 2020.
 - [10] A. Géron, *Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow*, 2nd ed. Sebastopol, CA, USA: O'Reilly Media, 2022.
 - [11] J. Leskovec, A. Rajaraman, and J. D. Ullman, *Mining of Massive Datasets*, 3rd ed. Cambridge, UK: Cambridge University Press, 2020.