

ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN APLIKASI TINDER DENGAN METODE SUPPORT VECTOR MACHINE

Fasya Mutiara Pagi¹, Nelly Indriani Widiastuti²

^{1,2} Teknik Informatika, Universitas Komputer Indonesia
Jl. Dipati Ukur No. 112-116 Coblong, Kota Bandung, Jawa Barat 40132

E-mail : nelly.indriani@email.unikom.ac.id²

Abstrak

Penelitian dilakukan untuk mengevaluasi performa *Support Vector Machine* (SVM) dalam mengkategorikan sentimen pada ulasan pengguna aplikasi Tinder, dengan fokus pada empat aspek: harga, fitur aplikasi, keamanan akun, dan nonaspek. Dari 1627 ulasan yang dianalisis, dibagi menjadi 1138 untuk pelatihan dan 489 untuk pengujian. Evaluasi kinerja dilakukan menggunakan matriks akurasi, *precision*, *recall*, dan *F1-score*. Hasil mengindikasikan kinerja SVM yang bervariasi tergantung pada aspek yang dianalisis. Tingkat akurasi tertinggi diperoleh pada aspek keamanan akun dengan nilai 0.944, sedangkan akurasi terendah ditemukan pada aspek fitur aplikasi yaitu 0.8998. *Precision* tertinggi untuk sentimen negatif dan netral ada pada aspek keamanan akun, sementara *precision* untuk sentimen positif paling rendah. Sebaliknya, *recall* tertinggi ditemukan pada sentimen netral, terutama pada fitur aplikasi dan keamanan akun. Namun, *recall* untuk sentimen positif pada aspek fitur aplikasi sangat rendah, menunjukkan kesulitan model dalam mendeteksi ulasan positif. Secara keseluruhan, SVM menunjukkan kinerja baik, terutama pada aspek keamanan akun dan sentimen netral, tetapi mengalami tantangan pada aspek fitur aplikasi dalam mengklasifikasikan sentimen positif.

Kata kunci : Analisis Sentimen, *Support Vector Machine*, Klasifikasi, Ulasan Tinder

Abstract

The research was conducted to evaluate the performance of the *Support Vector Machine* (SVM) in classifying sentiment in Tinder app user reviews, focusing on four aspects: price, app features, account security, and non-aspect. Of the 1627 reviews analyzed, 1138 were used for training and 489 for testing. Performance evaluation was carried out using accuracy, precision, recall, and F1-score metrics. The results indicated that SVM performance varied depending on the aspect analyzed. The highest accuracy was achieved in the account security aspect, with a score of 0.944, while the lowest accuracy was found in the app features aspect, at 0.8998. The highest precision for negative and neutral sentiments was observed in the account security aspect, while precision for positive sentiment was the lowest. Conversely, the highest recall was found for neutral sentiment, particularly in the app features and account security aspects. However, recall for positive sentiment in the app features aspect was very low, indicating the model's difficulty in detecting positive reviews. Overall, SVM demonstrated good performance, especially in the account security aspect and neutral sentiment, but faced challenges in classifying positive sentiment in the app features aspect.

Keywords : Sentiment Analysis, *Support Vector Machine*, Classification, Tinder Reviews

1. PENDAHULUAN

Tinder adalah aplikasi yang digunakan untuk mencari teman secara *online*, pertama kali diperkenalkan pada tahun 2012 oleh Justin Mateen, Sean Rad, Jonathan Badeen, dan rekan-rekan lainnya di Los Angeles, California, Amerika Serikat[1]. Aplikasi ini tersedia dalam versi berbayar dan gratis, di mana pengguna yang membayar mendapatkan fitur yang lebih lengkap, yang berdampak positif maupun negatif terhadap pengalaman mereka[2]. Tinder dapat diunduh melalui *Google Play Store*, tempat di mana pengguna juga bisa memberikan ulasan atau *rating* terhadap aplikasi tersebut[3]. Tinder sebagai aplikasi kencan *online* yang paling populer dengan pengguna yang terus bertambah setiap tahun, mengalami peningkatan ulasan

di *Google Play Store*[1]. Seiring bertambahnya unduhan Tinder mengakibatkan semakin beragamnya polaritas dalam ulasan tersebut.

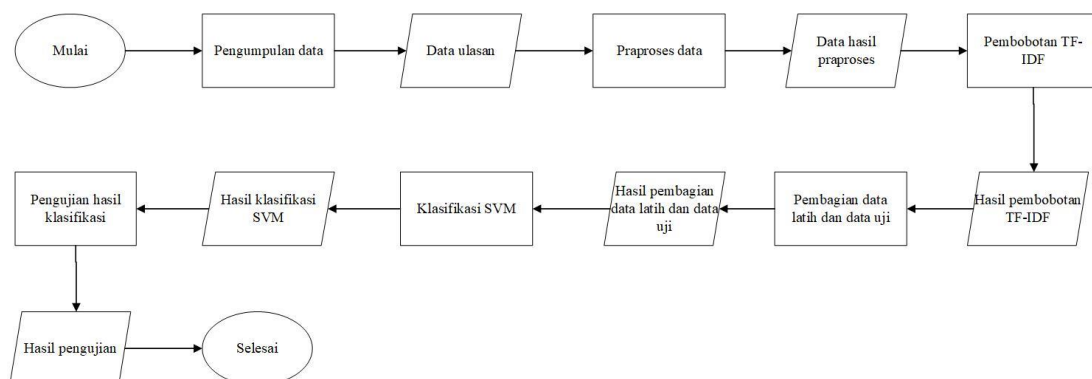
Analisis sentimen, atau yang sering disebut sebagai pemahaman opini, merupakan proses untuk mengidentifikasi dan mengambil informasi tentang perasaan atau pandangan seseorang dari data teks. [4]. Opini ini dapat dikategorikan sebagai positif, negatif, atau netral[5]. Analisis sentimen berbasis aspek atau ABSA merupakan bentuk lanjutan dari analisis sentimen yang berfokus pada kalimat. ABSA melibatkan dua tugas utama, yaitu ekstraksi aspek untuk mengidentifikasi aspek yang dibahas dan klasifikasi aspek untuk menentukan sentimen sebagai positif, negatif, atau netral. Dengan pendekatan ini, ABSA memungkinkan pengungkapan aspek atau atribut produk yang diekspresikan oleh pengguna[6], agar mengevaluasi kinerja SVM secara lebih mendalam dan detail memastikan bahwa model tidak hanya bekerja baik secara keseluruhan tetapi juga efektif dalam menangani kompleksitas yang ada di setiap aspek data ulasan.

Salah satu pendekatan yang diterapkan dalam analisis sentimen adalah *Support Vector Machine* (SVM). Memanfaatkan ruang model dengan fungsi linear pada dimensi tinggi yang dikembangkan melalui metode optimisasi, suatu metode pembelajaran dikenal sebagai SVM[7]. Dipresentasikan oleh sekelompok orang bernama Boser, Guyon, dan Vapnik di tahun 1992, SVM mengatasi masalah klasifikasi dengan menemukan *hyperplane* margin maksimum yang secara optimal memisahkan dua kelas, dikenal sebagai *Maximum Marginal Hyperplane*. SVM dipilih sebagai *classifier* dalam kajian ini karena dianggap sebagai salah satu teknik yang paling populer, efisien, dan akurat dalam klasifikasi[8]. Penelitian terkait *support vector machine* yang dibandingkan dengan metode lainnya seperti K-NN dan *random forest*, menunjukkan bahwa SVM berhasil menjadi metode klasifikasi terbaik dengan nilai akurasi, presisi, dan *recall* senilai 89,4%, 89,5%, dan 89,7%[9]. Selanjutnya, penelitian tentang analisis sentimen menggunakan pendekatan SVM dan KNN di suatu aplikasi menunjukkan bahwa SVM mengungguli KNN, dengan akurasi 87,98%, presisi 88,55%, dan *recall* 95,4%[10]. Pengembangan SVM lebih lanjut memungkinkan klasifikasi untuk kelas ganda disebut multikelas, menerapkan pendekatan umum seperti pendekatan *one-against-all* (OAA). Dalam pendekatan ini, setiap kelas diuji terhadap kelas lainnya dengan *hyperplane* yang berbeda dan kelas untuk data baru ditentukan berdasarkan nilai tertinggi dari *hyperplane* tersebut[11].

Penelitian ini dilakukan untuk menguji seberapa optimal performa metode *support vector machine* dalam mengklasifikasi teks ulasan aplikasi Tinder.

2. METODOLOGI

Penelitian ini melibatkan serangkaian proses dari awal hingga akhir. Tahap awal yaitu pengambilan ulasan aplikasi Tinder di *Google Play Store*, setelah itu akan ada praproses data yang melibatkan tujuh langkah, yaitu pembersihan data (*cleaning*), perubahan huruf menjadi kecil semua (*case folding*), pemecahan teks menjadi token (*tokenizing*), normalisasi, konversi negasi (*convert negation*), penghapusan kata umum (*stopword removal*), dan *stemming*. Setelah itu kata-kata ulasan akan diambil nilai bobotnya dengan teknik TF-IDF, lalu pembagian data dengan skenario 70% untuk pelatihan dan 30% untuk pengujian, klasifikasi dengan SVM dan diakhiri dengan evaluasi performa hasil klasifikasi. Proses penelitian secara runtun ditunjukkan Gambar 1.



Gambar 1. Alur penelitian

2.1 Data

Data masukan untuk penelitian adalah ulasan aplikasi Tinder yang diambil dari situs *Google Play Store* dengan total ulasan 1627 yang dikumpulkan antara bulan Juni 2023 hingga Desember 2023 melalui teknik *web scrapping* menggunakan *library google_play_scrapper*. Setiap ulasan telah dilabeli secara manual dan bersifat multikelas, mencakup empat aspek yang dianalisis yaitu harga, fitur aplikasi[2], keamanan akun dan nonaspek. Setiap aspek pada ulasan diberikan notasi polaritas sentimen, dengan "0" bagi sentimen netral, "1" bagi sentimen positif dan "-1" bagi sentimen negatif.

2.2 Preprocessing

Ekstraksi informasi dari sumber yang tidak terstruktur seringkali menghadapi tantangan dalam otomatisasi komputerisasi, sehingga diperlukan proses yang dapat mengubah dokumen tidak terstruktur menjadi format yang terstruktur. Proses ini dikenal sebagai *Text Preprocessing*, yang melibatkan konversi dokumen tidak terstruktur menjadi data dengan nilai-nilai numerik. Setelah data tersebut terstruktur dan diberi nilai numerik maka data tersebut siap untuk diproses ke proses selanjutnya[12]. Terdapat beberapa langkah *preprocessing* pada penelitian ini, yaitu:

- Cleaning*
Cleaning merupakan tahapan untuk mengidentifikasi dan menghapus data yang tidak valid atau tidak relevan seperti angka, tanda baca, *link*, atau simbol[12].
- Case folding*
Case folding adalah proses untuk mengganti tiap huruf dalam data atau ulasan menjadi huruf kecil semua[12].
- Tokenizing*
Tokenizing adalah proses untuk mengubah teks menjadi token-token atau kumpulan kata yang terpisah[12].
- Normalization*
Normalization merupakan proses untuk membetulkan token atau kata yang salah penulisan menjadi benar atau kata yang disingkat[13].
- Convert negation*
Convert negation adalah proses yang melibatkan penyatuan token negasi bersama token yang mengikutinya[14]. Token-token negasi yang dimaksud dalam penelitian ini seperti: tidak, bukan, jangan, belum, kurang, dan tanpa.
- Stopword removal*
Stopword removal yaitu tahapan untuk menghapus token-token yang dikiranya tidak ada makna atau tidak berpengaruh untuk proses selanjutnya[12]. Dalam penelitian ini, *stoplist* yang digunakan berasal dari *library nltk*.
- Stemming*
Stemming yaitu proses mengubah kata atau token yang memiliki imbuhan ke kata dasar[12]. Tahapan ini akan dilakukan dengan *library Sastrawi*.

2.3 Pembobotan Term Frequency-Inverse Document Frequency (TF-IDF)

Pembobotan Frekuensi Kemunculan Term-Frekuensi Invers Dokumen (TF-IDF) yaitu teknik untuk menilai signifikansi kata dalam sebuah dokumen dibandingkan dengan koleksi dokumen lainnya dengan menggabungkan frekuensi munculnya token di dokumen tertentu *Term Frequency* (TF) dan seberapa tidak umum token tersebut muncul di semua jumlah dokumen, yaitu *Inverse Document Frequency* (IDF), sehingga menghasilkan bobot yang digunakan dalam analisis sentimen untuk mengklasifikasikan teks berdasarkan sentimen[15]. Berikut merupakan rumus TF-IDF.

$$TF(t_d, d_t) = f(t_d, d_t) \quad (1)$$

$$IDF_t = \log_{10} \left(\frac{D}{df_t} \right) \quad (2)$$

$$W_{d,t} = TF_{d,t} \times IDF_{d,t} \quad (3)$$

Dimana:

- TF : jumlah kemunculan kata yang dicari dalam sebuah dokumen
- $f(t_d, d_t)$: umlah frekuensi kemunculan istilah dalam dokumen yang sama
- D : jumlah keseluruhan dokumen
- df_t : jumlah dokumen yang memuat istilah t
- IDF : Bobot invers dokumen

t : kata ke- t dari kata kunci
 W : bobot istilah ke- t dalam dokumen ke- d

2.4 Support Vector Machine (SVM)

Metode *Support Vector Machine* (SVM) adalah pembelajaran mesin untuk klasifikasi atau pengelompokan dengan cara mengklasifikasikan ruang model menjadi kelas biner. Dipresentasikan oleh sekelompok orang bernama Boser, Guyon, dan Vapnik di tahun 1992, SVM mengatasi masalah klasifikasi dengan menemukan *hyperplane* margin maksimum yang secara optimal memisahkan dua kelas, di mana *hyperplane* dengan margin lebih besar cenderung lebih akurat dalam mengklasifikasikan data. Margin adalah dua kali lipat jarak antara *hyperplane* dan *support vector*, di mana *support vector* adalah posisi yang paling mendekati *hyperplane*[8].

Persamaan 4 merupakan persamaan untuk pencarian lokasi *hyperplane*[16].

$$f(x) = \vec{w} \cdot \vec{x} + b \quad (4)$$

Dimana:

$\vec{w}$: vektor bobot
 $\vec{x}$: vektor fitur
 b : bias

Jarak margin dapat dioptimalkan dengan memanfaatkan *hyperplane* dan titik terdekatnya yang dihitung sebagai $\frac{1}{||w||}$. Untuk menemukan titik tersebut, diperlukan pemecahan persamaan *Quadratic Programming* (QP) *Problem*, yang melibatkan pencarian nilai minimum [16] pada persamaan 5.

$$\min r(w) = \frac{1}{2} ||w||^2 \quad (5)$$

Dimana:

$r(w)$: fungsi objektif yang diminimalkan
 w : vektor bobot

Penentuan *constraint* [16] menggunakan persamaan 6.

$$y_i (w \cdot x_i + b) - 1 \geq 0, (i = 1, \dots, n) \quad (6)$$

Dimana:

y_i : label kelas dari data latih ke- i
 $w \cdot x_i$: hasil dot product antara vektor bobot w dan vektor fitur x_i
 b : bias
 n : jumlah total data latih

Permasalahan ini dapat diatasi dengan berbagai teknik komputasi, salah satu pendekatan yang lebih sederhana adalah dengan mengubah persamaan 5 menjadi fungsi *Lagrangian* [16] seperti yang ditunjukkan dalam persamaan 7 berikut.

$$Lp = \frac{1}{2} ||w||^2 - \sum_{i=1}^n a_i y_i (x_i \cdot w + b) - 1 \quad (7)$$

Dimana:

Lp : fungsi Lagrangian
 w : vektor bobot
 a_i : *Lagrange multiplier*
 y_i : label kelas
 x_i : vektor fitur
 b : bias

Karena nilai Pengali *Lagrange* (α) tidak dapat ditentukan, persamaan tersebut tidak bisa langsung diatasi untuk menentukan w dan b . Dalam rangka mengatasi hal ini, persamaan 7 diubah menjadi masalah

maksimasi dengan menerapkan kondisi optimalitas dualitas menggunakan batasan Karush-Kuhn-Tucker (KKT)[16] sebagai berikut:

Aturan 1 :

$$\alpha_i[y_i(\mathbf{w} \cdot \mathbf{x}_i + b) - 1] = 0 \quad (8)$$

Dimana:

α_i : Lagrange multiplier
 y_i : label kelas
 $\mathbf{w}$: vektor bobot
 $\mathbf{x}_i$: vektor fitur
 b : bias

Aturan 2 :

$$\alpha_i > 0, i = 1, 2, \dots, N \quad (9)$$

Dimana:

α_i : Lagrange multiplier
 i : index
 N : jumlah total titik data dalam dataset

Persoalan optimasi tetap menantang untuk diselesaikan karena jumlah parameter yang banyak ($\mathbf{w}$, b , dan α_i). Untuk mempermudah penyelesaiannya, persamaan optimasi 7 perlu diubah menjadi fungsi *Lagrange Multiplier* (yang dikenal sebagai dualitas masalah). Persamaan *Lagrange Multiplier* [16] dapat dijabarkan sebagai persamaan 10.

$$Lp = \frac{1}{2} \|\mathbf{w}\|^2 - \sum_{i=1}^n \alpha_i y_i (\mathbf{w} \cdot \mathbf{x}_i) - b \sum_{i=1}^n \alpha_i y_i + \sum_{i=1}^n \alpha_i \quad (10)$$

Dimana:

Lp : fungsi Lagrangian
 $\mathbf{w}$: vektor bobot
 α_i : Lagrange multiplier
 y_i : label kelas
 $\mathbf{x}_i$: vektor fitur
 b : bias

2.6 Confusion Matrix

Matriks evaluasi adalah sebuah tabel untuk menyajikan hasil klasifikasi dari sistem, baik berdasarkan data sebenarnya maupun prediksi, dan berguna untuk menilai performa model klasifikasi. Dalam matriks ini terdapat beberapa istilah penting: negatif yang akurat (TN) merujuk pada data yang aktualnya negatif dan juga diprediksi sebagai negatif; positif yang salah (FP) merujuk data yang aktualnya positif namun diprediksi sebagai negatif; negatif yang salah (FN) merujuk pada data yang aktualnya negatif tetapi diprediksi sebagai positif; dan positif yang akurat (TP) menunjukkan data yang aktualnya positif dan juga diprediksi sebagai positif[17]. Tabel 1 berikut merupakan tabel matriks evaluasi.

Tabel 1. Confusion Matrix

Tanggapan	Peramalan Data	
	Negatif	Positif
Negatif	TN (<i>True Negative</i>)	FN (<i>False Negative</i>)
Positif	FP (<i>False Positive</i>)	TP (<i>True Positive</i>)

Akurasi merupakan tingkat kesesuaian antara hasil klasifikasi model dan nilai sesungguhnya. Ini menggambarkan proporsi kategorisasi yang akurat yang dihasilkan melalui model[17].

$$Accuracy = \frac{\text{Jumlah aspek atau kategori yang benar dideteksi}}{\text{Jumlah aspek atau kategori yang beranotasi}} \quad (11)$$

Precision adalah tingkat efektif model dalam mendeteksi sampel positif yang sebenarnya [17].

$$Precision = \frac{TP}{TP+FP} \quad (12)$$

Recall menilai seberapa efektif model dalam menemukan sampel yang benar-benar positif. Ini dihitung dengan membandingkan jumlah data positif yang berhasil diidentifikasi oleh model dibandingkan dengan total data positif yang sesungguhnya ada [17].

$$Recall = \frac{TP}{TP+FN} \quad (13)$$

F-measure diperoleh melalui cara menyatukan tingkat *recall* dengan *precision* [17].

$$f - measure = 2 \times \frac{precision \times recall}{precision + recall} \quad (14)$$

3. HASIL DAN PEMBAHASAN

Kumpulan data yang diterapkan dalam pengujian ini meliputi empat aspek yaitu harga, fitur aplikasi, keamanan akun, dan nonaspek. Keseluruhan data berjumlah 1627, dengan pembagian polaritas data untuk aspek harga adalah 142 data negatif, 1468 data netral dan 16 data positif, untuk aspek fitur aplikasi adalah 294 data negatif, 1318 data netral dan 15 data positif, untuk aspek keamanan akun adalah 318 data negatif, 1298 data netral dan 11 data positif dan nonaspek adalah 193 data negatif, 796 data netral dan 638 data positif.

Pengujian akurasi pada keempat aspek menggunakan data 1138 pelatihan dan 489 pengujian dengan skenario pembagian data 70% pelatihan dan 30% pengujian, akan disajikan dalam bentuk *confusion matrix*. Pada *confusion matrix* untuk klasifikasi aspek terdapat 2 kelas atau 2x2, yaitu kelas yang termasuk aspek dengan label "1" dan kelas yang tidak termasuk aspek dengan label "-1". Sementara itu, untuk klasifikasi sentimen, *confusion matrix* memiliki 3 kelas atau 3x3 karena terdapat 3 kategori sentimen yang berbeda: kelas positif dengan label "1", kelas negatif dengan label "-1", dan kelas netral dengan label "0".

3.1 Pengujian SVM

Pengujian akurasi adalah tahap yang dilakukan untuk mengukur sejauh mana sistem yang sudah dikembangkan dapat beroperasi dengan benar. Proses ini dikerjakan menggunakan metode *confusion matrix* untuk menguji kinerja metode yang diterapkan dalam analisis sentimen berbasis aspek terhadap data yang tersedia.

Tabel 2 berikut merupakan hasil pengujian keempat aspek yaitu harga, fitur aplikasi, keamanan akun, dan nonaspek.

Tabel 2. Hasil Pengujian Aspek

Aspek	Sentimen	Accuracy	Precision	Recall	F1-Score
Harga	Aspek	0.959	0.937	0.625	0.749
	Bukan Aspek		0.960	0.996	0.977
Fitur Aplikasi	Aspek	0.9038	0.9259	0.537	0.679
	Bukan Aspek		0.9011	0.9898	0.9433
Keamanan akun	Aspek	0.9509	0.974	0.78	0.8662
	Bukan Aspek		0.9463	0.994	0.969
Nonaspek	Aspek	0.9509	0.974	0.928	0.9504
	Bukan Aspek		0.928	0.974	0.9504

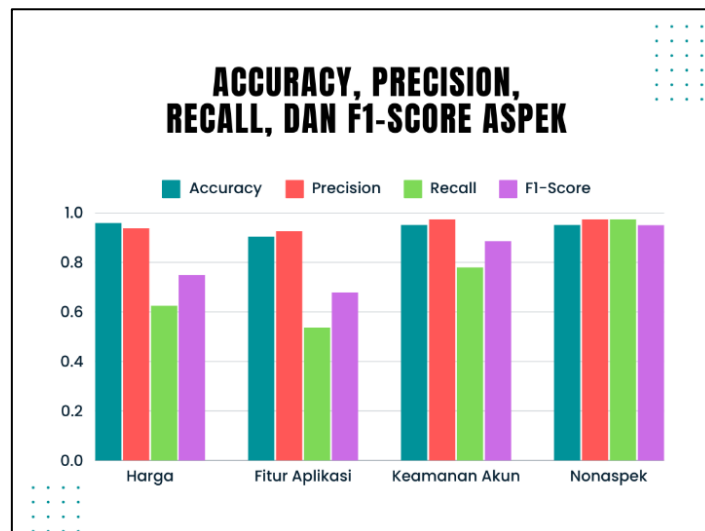
Tabel 3 berikut merupakan hasil pengujian sentimen keempat aspek yaitu harga, fitur aplikasi, keamanan akun, dan nonaspek.

Tabel 3. Hasil Pengujian Sentimen Berdasarkan Aspek

Aspek	Sentimen	Accuracy	Precision	Recall	F1-Score
Harga	Negatif	0.9550	0.9230	0.558	0.6955
	Netral		0.9564	0.9955	0.9755
	Positif		1	0.8	0.8889
Fitur Aplikasi	Negatif	0.8998	0.9230	0.5454	0.6856
	Netral		0.8970	0.9949	0.9434
	Positif		1	0	0
Keamanan akun	Negatif	0.944	0.972	0.755	0.849
	Netral		0.941	1	0.969
	Positif		0.5	0.33	0.397
Nonaspek	Negatif	0.93	0.913	0.734	0.81
	Netral		0.890	0.983	0.934
	Positif		0.988	0.921	0.953

3.2 Kesimpulan Pengujian

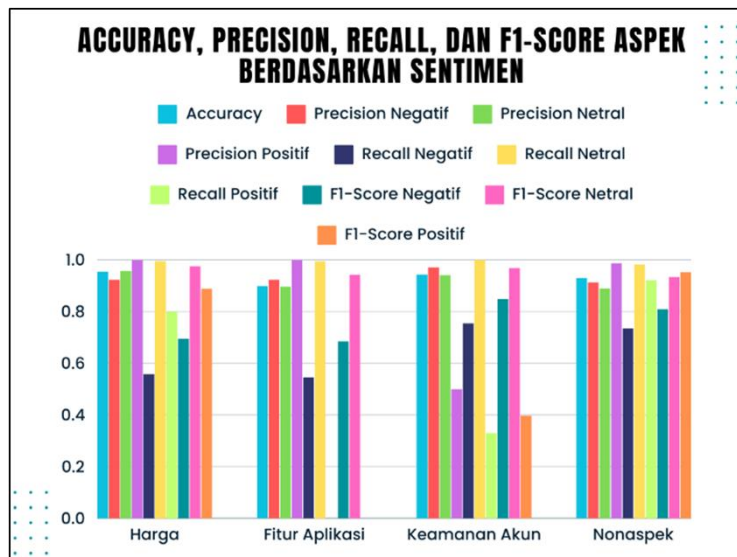
Pengujian performa dilakukan dengan menggunakan 1627 ulasan yang mencakup empat aspek berbeda. Nilai akurasi tertinggi dicapai pada aspek harga yaitu 0.959, menunjukkan bahwa model SVM dapat secara konsisten mengidentifikasi ulasan dengan tingkat kesalahan yang rendah. *Precision* tertinggi ditemukan pada aspek keamanan akun dan nonaspek yaitu 0.974, menunjukkan bahwa model jarang memberikan prediksi positif palsu untuk kedua aspek ini. *Recall* tertinggi ditemukan pada nonaspek yaitu 0.928, menunjukkan bahwa model berhasil menemukan hampir semua ulasan yang termasuk dalam kategori ini. *F1-score* tertinggi dicapai oleh nonaspek yaitu 0.9504, menunjukkan keseimbangan yang optimal antara *precision* dan *recall*. Secara keseluruhan, model menunjukkan kinerja yang baik dalam mengklasifikasikan aspek, terutama pada aspek nonaspek dan keamanan akun.



Gambar 2. Perbandingan Hasil Pengujian Aspek

Hasil analisis akurasi dalam klasifikasi sentimen berbasis aspek menghasilkan akurasi tertinggi pada aspek harga yaitu 0.955, yang mengindikasikan bahwa model hampir selalu tepat dalam mengklasifikasikan sentimen pada aspek ini. Pada *precision*, aspek keamanan akun menunjukkan presisi tertinggi pada sentimen negatif yaitu 0.972 dan netral yaitu 0.941, tetapi *precision* untuk sentimen positif paling rendah yaitu 0.5, menunjukkan kesulitan model dalam membedakan sentimen positif dari yang lainnya pada aspek ini. *Recall* tertinggi tercatat pada sentimen netral di semua aspek, terutama pada aspek keamanan akun yaitu 1, menunjukkan kemampuan model untuk mendeteksi hampir semua ulasan netral. Namun, *recall* pada sentimen positif untuk aspek fitur aplikasi paling rendah yaitu 0, menunjukkan kegagalan model dalam mendeteksi ulasan dengan sentimen positif pada aspek ini. *F1-score* tertinggi ditemukan pada sentimen

netral di aspek harga yaitu 0.9755, menunjukkan performa keseluruhan yang kuat dalam klasifikasi sentimen netral. Namun, *f1-score* untuk sentimen positif di aspek fitur aplikasi adalah 0, yang mengindikasikan bahwa model tidak mampu memprediksi ulasan positif dengan baik pada aspek ini.



Gambar 3. Perbandingan Hasil Pengujian Aspek Berdasarkan Sentimen

4. PENUTUP

Dari hasil keseluruhan proses pembangunan sistem analisis sentimen berbasis aspek yang memanfaatkan metode *support vector machine* terhadap tanggapan aplikasi Tinder, dapat disimpulkan bahwa aspek harga menunjukkan akurasi tertinggi sementara aspek fitur aplikasi memiliki akurasi terendah. Meskipun akurasi menjadi metrik utama dalam evaluasi *training*, nilai *recall* yang rendah, seperti pada aspek fitur aplikasi yang memiliki nilai *recall* 0 untuk sentimen positif, tetap perlu diperhatikan. Secara keseluruhan, metode SVM terbukti efektif dalam mengklasifikasikan aspek dan sentimen, namun aspek fitur aplikasi memerlukan peningkatan lebih lanjut, terutama dalam nilai *recall* dan *f1-score* untuk sentimen positif. Adapun saran untuk pengembang selanjutnya adalah perluas penelitian dengan jumlah ulasan yang lebih besar dan *balance* pada setiap sentimen di berbagai aspek.

DAFTAR PUSTAKA

- [1] W. C. P. Basel, N. W. Sitasari, and S. Safitri, "Bagaimana Self Disclosure dan Cyber Violence pada Pengguna Aplikasi Kencan Online Tinder Dewasa Awal di Jakarta," *Nalar: Jurnal Psikologi*, vol. 20, no. 2, Dec. 2022, doi: 10.47007/jpsi.v20i2.267.
- [2] A. Paramitha, S. Tanuwijaya, and S. Natakoesoemah, "Analisis Motif dan Dampak Penggunaan Aplikasi Tinder Berbayar," *Jurnal Psikologi dan Komunikasi*, vol. 5, 2021.
- [3] Lintang Aura, S. Dwi, R. Latifa, S. A. Kurniawan, and S. A. Nilasari, "Analisis Dataset Google Play Store Menggunakan Metode Exploratory Data Analysis (EDA)," 2022, doi: 10.13140/RG.2.2.14192.
- [4] A. Asrumi, D. Suharijadi, A. D. Setiara, and D. P. Wulanda, *Analisis Sentimen dan Penggalan Opini*, CV. Eurika Media Aksara, 2022.
- [5] Y. Wang, G. Shen, and L. Hu, "Importance Evaluation of Movie Aspects: Aspect-Based Sentiment Analysis," in 2020 5th International Conference on Mechanical, Control and Computer Engineering (ICMCCE), Harbin, China: IEEE, Dec. 2020, pp. 2444–2448. doi: 10.1109/ICMCCE51767.2020.00527.
- [6] D. Purnamasari, A. B. Aji, D. W. A. P., F. A. Reza, M. S. O., N. Yanda, and U. Hidayat, *Pengantar Metode Analisis Sentimen*, Gunadarma, 2023.
- [7] N. I. Widiastuti, E. Rainarli, and K. E. Dewi, "Peringkasan dan Support Vector Machine pada Klasifikasi Dokumen," *INFOTEL*, vol. 9, no. 4, p. 416, Nov. 2017, doi: 10.20895/infotel.v9i4.312.

- [8] N. Cristianini and J. Shawe-Taylor, *An Introduction to Support Vector Machines and Other Kernel-Based Learning Methods*, Cambridge University Press, 2000.
- [9] S. Watmah, S. Suryanto, and M. Martias, "Komparasi metode K-NN, Support Vector Machine dan Random Forest pada e-commerce Shopee," *INSTK*, vol. 2, no. 1, pp. 15–21, 2021.
- [10] M. N. Muttaqin and I. Kharisudin, "Analisis sentimen aplikasi Gojek menggunakan Support Vector Machine dan K Nearest Neighbor," *UNNES Journal of Mathematics*, 2021.
- [11] S. Abe, "Variants of Support Vector Machines," in *Support Vector Machines for Pattern Classification*, Advances in Pattern Recognition, London: Springer London, 2010, pp. 163–226. doi: 10.1007/978-1-84996-098-4_4.
- [12] A. Nurkholis, D. Alita, and A. Munandar, "Comparison of Kernel Support Vector Machine Multi-Class in PPKM Sentiment Analysis on Twitter," *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 6, no. 2, pp. 227–233, Apr. 2022, doi: 10.29207/resti.v6i2.3906.
- [13] M. Rahardi, A. Aminuddin, F. F. Abdulloh, and R. A. Nugroho, "Sentiment Analysis of Covid-19 Vaccination using Support Vector Machine in Indonesia," *Int. J. Advanced Computer Science and Applications (IJACSA)*, vol. 13, no. 6, 2022, doi: 10.14569/IJACSA.2022.0130665.
- [14] S. Mehta, "When to Use Negation Handling in Sentiment Analysis?," [Online]. Available: <https://analyticsindiamag.com/ai-mysteries/when-to-use-negation-handling-in-sentiment-analysis/>. [Accessed: Sep. 2, 2024, 12:50].
- [15] F. Karabiber, "TF-IDF — Term Frequency-Inverse Document Frequency," [Online]. Available: <https://www.learndatasci.com/glossary/tf-idf-term-frequency-inverse-document-frequency/>. [Accessed: Sep. 2, 2024, 17:14].
- [16] B. Schölkopf and A. J. Smola, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*, MIT Press, 2002.
- [17] D. Musfiroh, U. Khaira, P. E. P. Utomo, and T. Suratno, "Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon: Sentiment Analysis of Online Lectures in Indonesia from Twitter Dataset Using InSet Lexicon," *MALCOM*, vol. 1, no. 1, pp. 24–33, Mar. 2021, doi: 10.57152/malcom.v1i1.20.