

## IMPLEMENTASI ALGORITMA SVM DAN OPTIMASI MENGGUNAKAN KOMPARASI N-GRAM DAN GLOVE PADA SENTIMEN PENGGUNA APLIKASI ACCESS BY KAI

Raka Kusuma Wardana, Nuri Cahyono, Uyock Anggoro Saputro, Arifiyanto Hadi Negoro, Nur Aini

Informatika, Universitas Amikom Yogyakarta  
Jl. Ring Road Utara, Daerah Istimewa Yogyakarta  
rakka.kusuma@students.amikom.ac.id

### ABSTRAK

Analisis sentimen memiliki peran penting dalam memahami opini pengguna, khususnya dalam meningkatkan kualitas layanan digital. Penelitian ini berfokus pada optimasi klasifikasi sentimen menggunakan Support Vector Machine (SVM) dengan pendekatan N-Gram dan Word Embedding GloVe terhadap ulasan pengguna aplikasi Access by KAI. Permasalahan utama yang dihadapi adalah bagaimana meningkatkan akurasi klasifikasi sentimen agar lebih optimal dalam mengidentifikasi keluhan dan preferensi pengguna. Dataset yang digunakan terdiri dari 1.718 ulasan pengguna yang telah dilabeli secara manual ke dalam kategori positif dan negatif. Metode penelitian mencakup pengumpulan data, preprocessing, ekstraksi fitur, pemodelan SVM, dan evaluasi kinerja. Hasil penelitian menunjukkan bahwa pendekatan N-Gram berbasis TF-IDF menghasilkan akurasi terbaik sebesar 84,88%, mengungguli metode GloVe yang hanya mencapai 79,65%. Selain itu, pemodelan SVM tanpa optimasi menghasilkan akurasi 83,43%. Perbandingan ini menegaskan bahwa pemilihan metode ekstraksi fitur yang tepat berpengaruh signifikan terhadap performa klasifikasi. Temuan penelitian juga mengungkap bahwa sentimen negatif mendominasi dengan persentase 51,7%, yang mencerminkan tantangan dalam performa dan fitur aplikasi.

**Kata kunci :** Analisis Sentimen, SVM, N-Gram, Glove, Optimasi Klasifikasi.

### 1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi (TIK) dalam beberapa tahun terakhir telah membawa perubahan besar dalam berbagai aspek kehidupan, termasuk di sektor transportasi. Salah satu bentuk pemanfaatan TIK di bidang ini adalah pengembangan aplikasi mobile yang bertujuan meningkatkan pelayanan kepada pengguna[1].

PT Kereta Api Indonesia (KAI) mengembangkan aplikasi Access by KAI untuk mempermudah pemesanan tiket, penjadwalan ulang, serta mendapatkan e-boarding pass[2].

Meskipun aplikasi ini dirancang untuk meningkatkan kenyamanan dan efisiensi, banyak pengguna yang mengeluhkan performanya, seperti akses yang lambat dan ketidakstabilan sistem, yang berdampak pada pengalaman pengguna secara keseluruhan. Oleh karena itu, analisis terhadap ulasan pengguna menjadi penting untuk memahami permasalahan yang ada dan mencari solusi yang tepat guna meningkatkan kualitas aplikasi[3].

Analisis sentimen merupakan salah satu metode yang dapat digunakan untuk mengkategorikan ulasan pengguna menjadi sentimen positif dan negatif. Dalam proses klasifikasi ini, algoritma Support Vector Machine (SVM) banyak digunakan karena kemampuannya dalam menangani data berukuran besar serta menghasilkan klasifikasi yang akurat dan stabil[4].

Namun, efektivitas klasifikasi sentimen juga sangat bergantung pada teknik representasi teks yang digunakan. Berbagai pendekatan dapat diterapkan

untuk mengubah teks menjadi fitur numerik yang dapat diproses oleh algoritma pembelajaran mesin. Oleh karena itu, diperlukan penelitian yang membandingkan teknik representasi teks yang berbeda untuk menentukan metode yang paling optimal dalam meningkatkan akurasi analisis sentimen[5].

Dalam penelitian ini, teknik N-Gram dan GloVe dikomparasikan dalam satu model untuk analisis sentimen pada aplikasi Access by KAI[6].

Pemilihan kedua teknik ini didasarkan pada keunggulan masing-masing dalam menangkap pola teks. N-Gram efektif dalam mengidentifikasi hubungan lokal antar kata dan frasa yang sering muncul dalam ulasan pengguna, sementara GloVe mampu memahami hubungan semantik antar kata dengan memanfaatkan informasi dari keseluruhan korpus teks[7].

Dengan membandingkan kedua teknik ini dalam satu model, penelitian ini bertujuan mengevaluasi sejauh mana kombinasi atau dominasi salah satu metode dapat meningkatkan akurasi klasifikasi sentimen dalam konteks aplikasi transportasi digital[8].

Selain itu, penelitian ini juga mempertimbangkan efisiensi waktu pemrosesan guna memastikan solusi yang diusulkan tidak hanya akurat tetapi juga dapat diimplementasikan secara efektif dalam sistem nyata[9].

### 2. TINJAUAN PUSTAKA

Kandungan dalam tinjauan pustaka bisa berupa landasan teori yang relevan dengan topik skripsi yang

akan dibahas, di mana landasan teori ini menjadi panduan pertama dalam penyusunan skripsi.

## 2.1. Analisis Sentimen

Analisis sentimen adalah proses mengkategorikan opini atau sentimen dalam teks menjadi positif, negatif, atau netral. Proses ini mencakup pengumpulan data, preprocessing (seperti tokenisasi dan pembersihan teks), serta klasifikasi menggunakan algoritma machine learning. Teknik ini banyak digunakan untuk memahami persepsi publik terhadap suatu layanan atau produk[10].

## 2.2. Access by KAI

Access by KAI adalah aplikasi mobile dari PT KAI yang mempermudah pengguna dalam memesan tiket, mengecek jadwal, serta mendapatkan boarding pass digital. Fitur lain mencakup notifikasi perjalanan, layanan pelanggan, serta promosi tiket, yang meningkatkan kenyamanan pengguna dalam perjalanan kereta api[11].

## 2.3. Support Vector Machine (SVM)

SVM adalah metode klasifikasi berbasis supervised learning yang bekerja dengan membentuk hyperplane pemisah antara kelas positif dan negatif. SVM menggunakan berbagai jenis kernel, seperti linear, polynomial, RBF, dan sigmoid, yang memungkinkan pemisahan data baik secara linear maupun non-linear[12].

## 2.4. N-Gram

N-Gram adalah teknik dalam NLP untuk menganalisis teks berdasarkan urutan kata yang berdekatan, seperti unigram ( $n=1$ ), bigram ( $n=2$ ), dan trigram ( $n=3$ ). Metode ini membantu memahami pola teks dan digunakan dalam berbagai aplikasi seperti analisis sentimen dan pemodelan bahasa[13].

## 2.5. Word Embedding GloVe

GloVe adalah metode word embedding yang mengubah kata menjadi vektor numerik berdasarkan hubungan kata dalam korpus besar. Dengan menggunakan metode global matrix factorization, GloVe menangkap makna semantik antar kata untuk meningkatkan akurasi dalam analisis teks[14].

## 2.6. Evaluasi Akurasi Model Klasifikasi

Evaluasi model klasifikasi dilakukan menggunakan metrik seperti akurasi, presisi, recall, dan F1-score. Metrik ini mengukur seberapa baik model dalam mengklasifikasikan data secara akurat.

## 2.7. Penelitian Terdahulu

Penelitian terdahulu adalah kajian yang dilakukan sebelumnya dan digunakan sebagai bahan perbandingan serta referensi dalam penelitian ini. Berikut adalah beberapa penelitian yang relevan.

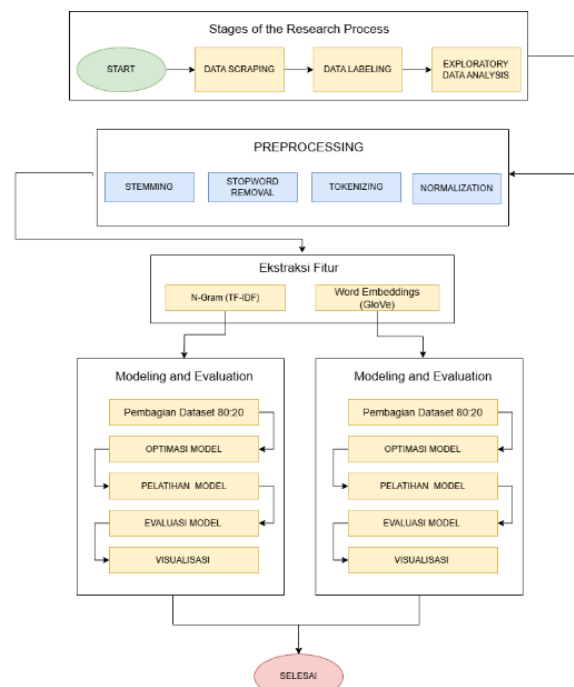
- Penelitian ini menggunakan metode Support Vector Machine (SVM) dan Random Forest untuk

menganalisis ulasan aplikasi KAI Access. Hasil penelitian menunjukkan bahwa SVM memberikan akurasi lebih tinggi (97%) dibandingkan dengan Random Forest (93%). Studi ini memperlihatkan pentingnya metode SVM dalam analisis sentimen terhadap aplikasi mobile[5].

- Penelitian ini meneliti penggunaan GloVe dalam meningkatkan akurasi model SVM untuk analisis sentimen pada proyek kereta cepat. Hasil penelitian menunjukkan bahwa integrasi GloVe meningkatkan akurasi klasifikasi sentimen dari 72.63% menjadi 77.72%. Ini menunjukkan bahwa GloVe dapat memperkaya representasi semantik dalam model SVM[15].

## 3. METODE PENELITIAN

Penelitian ini menggunakan metode *Support Vector Machine (SVM)* untuk menganalisis sentimen ulasan pengguna terhadap aplikasi *Access by KAI*. Fokus utama penelitian adalah mengoptimalkan klasifikasi SVM dengan pendekatan *N-Gram* dan *GloVe*. Alur penelitian dapat dilihat pada Gambar 1.



Gambar 1. Alur Penelitian  
Sumber : Raka Kusuma Wardana

### 3.1. Scraping Data

Data dalam penelitian ini diperoleh dari ulasan pengguna aplikasi Access by KAI di Google Playstore menggunakan Teknik scraping data [16], Peneliti menggunakan library 'google-play-scraper'. Pada proses ini peneliti mengumpulkan sebanyak 1718 ulasan yang dipilih secara acak.

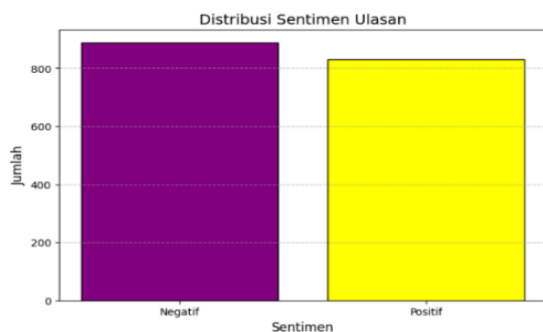
### 3.2. Labeling Data

Data yang diperoleh kemudian diberi label positif atau negatif. Peneliti melakukan pelabelan secara kolaboratif dengan bantuan beberapa pelabel guna

memaksimalkan akurasi dari hasil pelabelan[17]. Peneliti tidak mengandalkan score pada aplikasi karena ditemukan beberapa ulasan yang tidak relevan dengan score.

### 3.3. Exploratory Data Analysis (EDA)

Tahap Exploratory Data Analysis (EDA) bertujuan untuk memberikan pemahaman awal mengenai dataset yang akan digunakan. Analisis ini sangat penting untuk mengidentifikasi struktur dataset, pola distribusi data, serta potensi masalah seperti data yang hilang (*missing values*) atau ketidakseimbangan data. Untuk mempermudah pemahaman terhadap distribusi data, visualisasi dalam bentuk diagram pie dibuat menggunakan pustaka *matplotlib.pyplot*. Diagram batang ini menggambarkan proporsi antara ulasan positif dan negatif secara lebih jelas, di mana warna dan persentase ditampilkan untuk masing-masing kategori. Visualisasi ini tidak hanya memberikan wawasan intuitif mengenai dataset, tetapi juga menjadi dasar dalam memahami dinamika sentimen pengguna terhadap aplikasi Access by KAI [18].



Gambar 2. Distribusi Sentimen Ulasan  
Sumber : Raka Kusuma Wardana

### 3.4. Preprocessing Data

Pemrosesan data awal adalah langkah penting sebelum membangun model penambangan teks. Peneliti melakukan analisis singkat terhadap karakteristik data melalui beberapa tahapan preprocessing, seperti normalisasi, tokenisasi, penghapusan stopwords, dan stemming. [19].

#### a. Normalisasi

Proses ini mencakup penghapusan elemen disruptif dari teks dan normalisasi kata-kata tidak baku menjadi bentuk standar untuk mempermudah analisis sentiment[20]. Contoh kata yang dinormalisasi dapat dilihat pada Tabel 1.

Tabel 1. Normalisasi Kata

| No. | Kata Tidak Baku              | Kata Baku  |
|-----|------------------------------|------------|
| 1.  | gk,g,gak,ngak,gx,tdk,ga,ngga | Tidak      |
| 2.  | gabisa,gbsa,gsb              | Tidak Bisa |
| 3.  | tp,tpi,tapi                  | Namun      |
| 4.  | ok,oke,okee,okk              | Baik       |
| 5.  | lemot,lmt,lmot               | Lambat     |
| 6.  | kerta,kreta,sepur            | Kereta     |
| 7.  | pake,pke,pakee               | Pakai      |
| 8.  | indo                         | Indonesia  |

#### b. Tokenizing

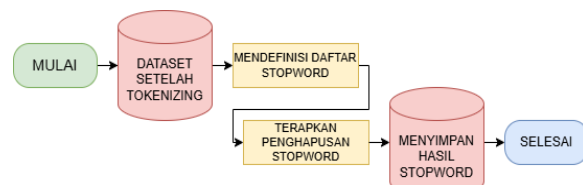
Proses ini memecah teks ulasan menjadi token atau kata-kata individu. Tokenization memudahkan analisis setiap kata dalam teks, terutama untuk metode seperti N-Gram[21]. Alur dari tokenisasi dapat dilihat pada gambar 3.



Gambar 3. Alur Tokenizing  
Sumber : Raka Kusuma Wardana

#### c. Stopword Removal

Menghapus kata-kata umum (stopwords) yang tidak memiliki kontribusi penting dalam analisis sentimen. Stopwords seperti “di”, “yang”, “dan”, sering kali tidak membawa makna signifikan dan dapat membingungkan model[22]. Alur Stopword Removal dapat dilihat pada Gambar 4.



Gambar 4. Alur Stopword Removal  
Sumber : Raka Kusuma Wardana

#### d. Stemming

Proses ini mengubah kata ke bentuk dasar atau akar katanya. Stemming membantu menyederhanakan variasi kata sehingga kata-kata dengan akar yang sama dikelompokkan menjadi satu[23]. Alur pada stemming dapat dilihat pada gambar 5.



Gambar 5. Alur Stemming  
Sumber : Raka Kusuma Wardana

### 3.5. Ekstraksi Fitur

Peneliti menggunakan pendekatan gabungan dalam ekstraksi fitur untuk analisis sentiment[24], dengan memanfaatkan dua teknik utama yaitu N-Gram dengan TF-IDF dan embedding Glove.

#### a. N-Gram dengan TF-IDF

N-Gram adalah metode untuk menangkap hubungan antar kata dalam teks, seperti unigram (satu kata), bigram (dua kata), atau trigram (tiga kata). Representasi N-Gram menggunakan TF-IDF, yang memberi bobot pada n-gram berdasarkan frekuensi kemunculannya dalam dokumen dan korpus, menggabungkan konsep Term Frequency (TF) dan Inverse Document Frequency (IDF).

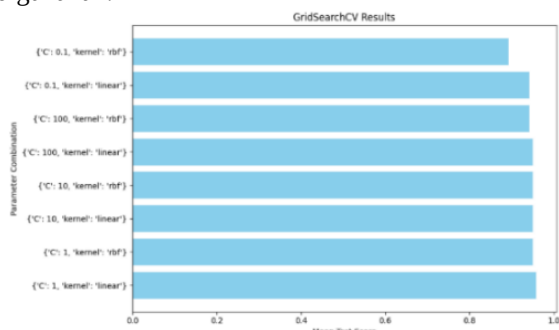
#### b. Embedding Glove

GloVe adalah metode embedding yang mengonversi kata-kata menjadi vektor berdimensi

tetap berdasarkan hubungan semantiknya dalam korpus. GloVe dilatih untuk menangkap makna kata dari konteks di mana kata tersebut muncul. Model ini memanfaatkan distribusi bersama dari kata-kata dalam korpus besar sehingga kata-kata yang sering muncul bersama akan memiliki vektor yang serupa. Pembelajaran GloVe berfokus pada meminimalkan perbedaan antara produk skalar dari dua kata dengan logaritma dari jumlah kemunculan mereka dalam konteks tertentu.

### 3.6. Implementasi Model

Model klasifikasi yang digunakan adalah Support Vector Machine (SVM). SVM akan memisahkan data ulasan ke dalam dua kelas, yaitu sentimen positif dan negatif, dengan memanfaatkan kernel linear untuk meningkatkan akurasi pemisahan data[25]. Berikut visualisasi kombinasi Nilai C dan Jenis Kernel yang digunakan.



Gambar 6. Kombinasi Nilai C dan Jenis Kernel  
Sumber : Raka Kusuma Wardana

Pemilihan kernel linear didasarkan pada performa yang ditunjukkan selama proses optimasi parameter menggunakan GridSearchCV, di mana kernel linear dengan parameter  $C = 1$  menghasilkan akurasi tertinggi. Parameter SVM yang digunakan, seperti kernel dan margin, dioptimalkan untuk mendapatkan performa terbaik.

## 4. HASIL DAN PEMBAHASAN

Penelitian ini berfokus pada aplikasi Access by KAI, dengan data yang diambil dari 2000 komentar pengguna di Google Play Store selama periode 2023-2024. Data dikumpulkan menggunakan Google Play Scraper, yang memungkinkan pengunduhan ulasan berdasarkan package name aplikasi (com.kai.kaiticketing). Setelah proses penyaringan, 1.718 komentar valid dipilih untuk analisis akhir, yang mencakup kolom-kolom penting seperti username, content, score, thumbsUpCount, dan reviewCreatedVersion. Contoh data yang di ambil dapat dilihat pada Tabel 2.

Tabel 2. Contoh Dataset

| Username      | Score | Content                                                                                     |
|---------------|-------|---------------------------------------------------------------------------------------------|
| Jihad Asy'ari | 2     | Kenapa tidak bisa mencari stasiun tujuan? Jaringan stabil dan aplikasi sudah versi terbaru. |

| Username  | Score | Content                                                                     |
|-----------|-------|-----------------------------------------------------------------------------|
| Bro Djoni | 5     | Perjalanan yang sangat menyenangkan dan praktis dalam pelayanannya... ⭐⭐⭐⭐⭐ |

Setelah terkumpul jumlah data yang di dapat adalah 1717 data, Proses pelabelan dilakukan secara manual untuk mengkategorikan komentar menjadi sentimen positif (830 komentar) dan negatif (887 komentar). Hasil analisis menunjukkan bahwa 48,3% dari total komentar bersentimen positif, sedangkan 51,7% bersentimen negatif, menandakan adanya tantangan dalam pelayanan aplikasi. Pelabelan ini penting untuk memahami persepsi pengguna terhadap aplikasi dan untuk analisis lebih lanjut. Peneliti melakukan pelabelan secara manual. Contoh proses pelabelan data ditunjukkan pada Tabel 3.

Tabel 3. Label Sentimen

| Content                                                                                     | Sentimen |
|---------------------------------------------------------------------------------------------|----------|
| Kenapa tidak bisa mencari stasiun tujuan? Jaringan stabil dan aplikasi sudah versi terbaru. | Negatif  |
| Perjalanan yang sangat menyenangkan dan praktis dalam pelayanannya... ⭐⭐⭐⭐⭐                 | Positif  |

Setelah proses pelabelan, Exploratory Data Analysis (EDA) dilakukan untuk melakukan pemeriksaan awal terhadap kumpulan data. Pada tahap ini menunjukkan tidak ada nilai yang hilang atau null dalam kumpulan data. Dari 1717 data yang di kategorikan, 830 data adalah positif dan 887 data adalah negatif. Tahap selanjutnya melibatkan preprocessing data. Proses preprocessing meliputi normalisasi data, tokenisasi, penghapusan stopwords, dan stemming untuk menyiapkan data agar lebih terstruktur. Normalisasi menghilangkan tanda baca dan mengubah teks menjadi format seragam. Tokenisasi memecah teks menjadi unit kata, sedangkan penghapusan stopwords menghilangkan kata-kata yang tidak signifikan. Stemming mengubah kata-kata menjadi bentuk dasarnya, menghasilkan data yang lebih bersih dan siap untuk analisis lebih lanjut. Bagian ini merupakan langkah penting dalam pengembangan penelitian ini. Contoh preprocessing data dapat dilihat pada Tabel 4.

Tabel 4. Hasil Preprocessing Data

| Proses           | Hasil Preprocessing                                                                         |
|------------------|---------------------------------------------------------------------------------------------|
| Data Awal        | Perjalanan yang sangat menyenangkan dan praktis dalam pelayanannya... ⭐⭐⭐⭐⭐                 |
| Normalisasi      | perjalanan yang sangat menyenangkan dan praktis dalam pelayanannya                          |
| Tokenizing       | ['perjalanan', 'yang', 'sangat', 'menyenangkan', 'dan', 'praktis', 'dalam', 'pelayanannya'] |
| Stopword Removal | [perjalanan, menyenangkan, praktis, pelayanannya]                                           |
| Stemming         | [jalan, senang, praktis, layan]                                                             |

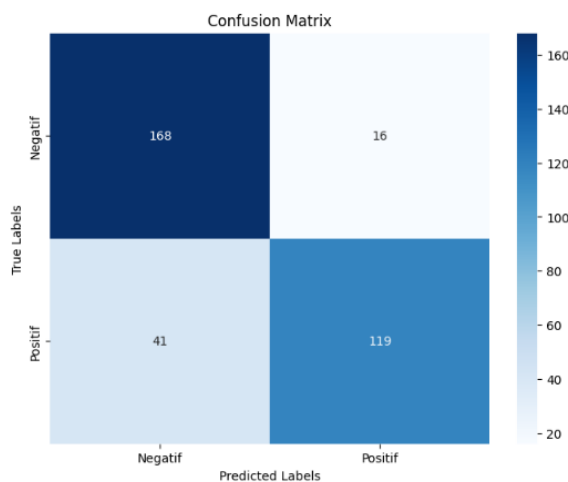
Selanjutnya penelitian ini membandingkan implementasi algoritma SVM tanpa optimasi dan dua pendekatan utama dalam analisis sentimen komentar aplikasi Access by KAI, yaitu N-Gram TF-IDF dan GloVe Embedding, yang masing-masing memiliki keunggulan dan kelemahan.

#### 4.1. Hasil SVM Tanpa Optimasi

Pada penelitian ini, model Support Vector Machine (SVM) tanpa optimasi diterapkan untuk mengklasifikasikan sentimen pengguna aplikasi Access by KAI berdasarkan data komentar yang tersedia. Evaluasi model dilakukan untuk memahami performa dasar dari algoritma ini tanpa menggunakan kernel khusus, seperti Linear atau Radial Basis Function (RBF).

Tabel 6. Hasil Klasifikasi SVM Tanpa Optimasi

| Kategori | Precision | Recall | F1-Score | Support |
|----------|-----------|--------|----------|---------|
| Negatif  | 0.80      | 0.91   | 0.85     | 184     |
| Positif  | 0.88      | 0.74   | 0.81     | 160     |
| Akurasi  | 0.83      | 0.83   | 0.83     | 344     |



Gambar 7. Confusion Matrix SVM Tanpa Optimasi  
Sumber: Raka Kusuma Wardana

Berdasarkan gambar 7 Confusion Matrix model Support Vector Machine (SVM) tanpa optimasi menunjukkan akurasi 83,43%. Kinerja model lebih baik dalam mengklasifikasikan komentar negatif, dengan precision 0,80 dan recall 0,91, menunjukkan kemampuan yang baik dalam mendeteksi sentimen negatif. Sebaliknya, untuk kelas positif, meskipun precision mencapai 0,88, recall hanya 0,74, menunjukkan adanya kesulitan dalam mendeteksi beberapa komentar positif, yang dapat menjadi fokus perbaikan. Dapat dilihat pada tabel 6 untuk hasil dari SVM tanpa optimasi dengan hasil sebagai berikut.

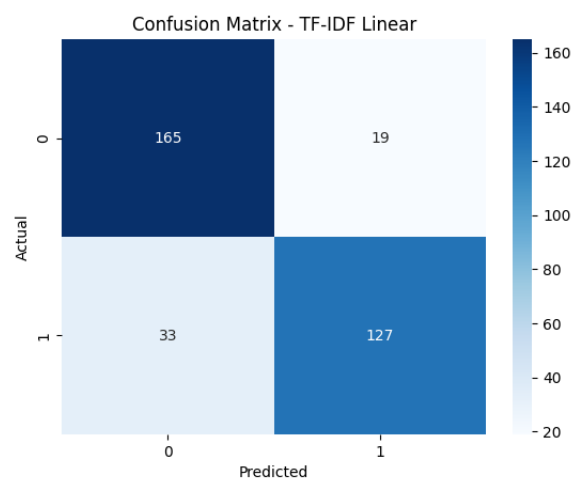
#### 4.2. Modeling N-Gram TF-IDF

Modeling analisis sentimen menggunakan N-Gram TF-IDF dipilih karena kemampuannya menangkap pola kata dan hubungan antar kata dalam komentar pengguna aplikasi Access by KAI. TF-IDF

memberi bobot lebih besar pada kata-kata informatif, sementara pendekatan N-Gram (1,3) mencakup unigram, bigram, dan trigram untuk memahami konteks lebih luas. Unigram mengenali kata tunggal, bigram menangkap frasa seperti "tidak baik", dan trigram membantu memahami konteks yang lebih kompleks, sehingga meningkatkan akurasi analisis sentimen terhadap ulasan yang lebih panjang dan mendetail. Berikut adalah rincian metrik evaluasi berdasarkan classification report dapat dilihat pada tabel 7.

Tabel 7. Hasil Model N-Gram TF-IDF

| Kategori | Precision | Recall | F1-Score | Support |
|----------|-----------|--------|----------|---------|
| Negatif  | 0.83      | 0.90   | 0.86     | 184     |
| Positif  | 0.87      | 0.79   | 0.83     | 160     |
| Akurasi  | 0.85      | 0.85   | 0.85     | 344     |



Gambar 8. Confusion Matrix N-Gram TF-IDF  
Sumber: Raka Kusuma Wardana

Berdasarkan gambar 8 Confusion matrix Model SVM dengan N-Gram TF-IDF menunjukkan hasil distribusi prediksi yang digunakan untuk meningkatkan akurasi model, yang mendapatkan hasil 84,88%. Model menunjukkan performa yang baik pada kelas negatif (precision 0,83, recall 0,90) dan positif (precision 0,87, recall 0,79), dengan keseimbangan yang lebih baik antara precision dan recall. Penggunaan N-Gram TF-IDF memungkinkan model untuk menangkap pola kata dan konteks lokal dengan lebih efektif, meningkatkan hasil analisis sentimen.

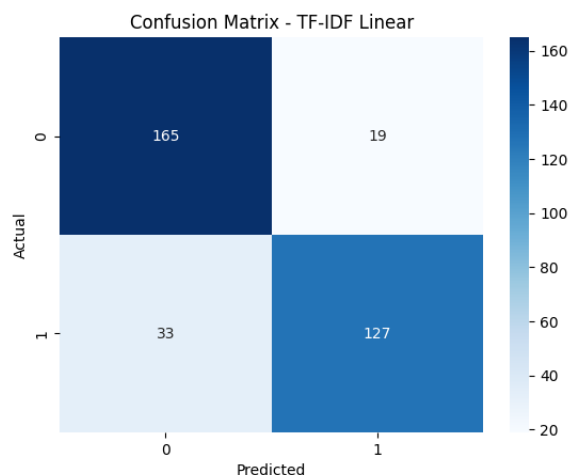
#### 4.3. Modeling Glove Embedding

GloVe (Global Vectors for Word Representation) adalah teknik embedding yang merepresentasikan kata dalam vektor 300 dimensi untuk menangkap hubungan semantik dalam teks. Dalam penelitian ini, GloVe digunakan untuk membantu model memahami konteks ulasan pengguna aplikasi Access by KAI secara lebih mendalam. Keunggulannya terletak pada kemampuannya merepresentasikan kata-kata dengan makna serupa dalam vektor yang berdekatan, seperti "baik" dan "bagus", sehingga lebih efektif

dibandingkan pendekatan berbasis frekuensi kata seperti TF-IDF dalam mengenali pola sentimen. Berikut adalah rincian metrik evaluasi berdasarkan classification report dapat dilihat pada tabel 8.

Tabel 8. Hasil Model GloVe Embedding

| Kategori | Precision | Recall | F1-Score | Support |
|----------|-----------|--------|----------|---------|
| Negatif  | 0.78      | 0.86   | 0.82     | 184     |
| Positif  | 0.82      | 0.72   | 0.77     | 160     |
| Akurasi  | 0.80      | 0.80   | 0.80     | 344     |

Gambar 9. Confusion Matrix N-Gram TF-IDF  
Sumber: Raka Kusuma Wardana

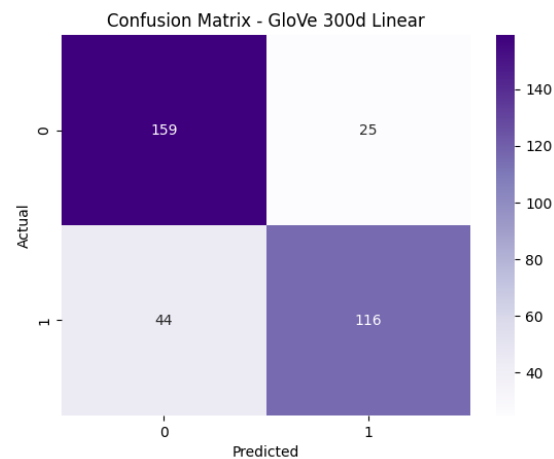
Berdasarkan gambar 9 Confusion matrix Model SVM dengan GloVe 300d menunjukkan akurasi 79,65%. Kinerja pada kelas negatif cukup baik (precision 0,78, recall 0,86), tetapi kelas positif menunjukkan recall yang lebih rendah (0,72), yang mengindikasikan beberapa komentar positif tidak terdeteksi dengan baik. Meskipun GloVe memiliki keunggulan dalam menangkap hubungan semantik, hasil ini menunjukkan bahwa pendekatan ini kurang optimal untuk dataset ini dibandingkan dengan N-Gram TF-IDF.

#### 4.4. Perbandingan Hasil

Hasil penelitian menunjukkan bahwa N-Gram TF-IDF memberikan performa lebih baik dibandingkan embedding GloVe dalam klasifikasi sentimen ulasan pengguna aplikasi Access by KAI. Berikut adalah rincian perbandingan evaluasi berdasarkan classification report dapat dilihat pada tabel 9.

Tabel 9. Perbandingan Hasil

| Model              | Precision | Recall | F1-Score | Accuracy |
|--------------------|-----------|--------|----------|----------|
| N-Gram TF-IDF      | 0.87      | 0.86   | 0.82     | 84,88%   |
| Glove Embedding    | 0.78      | 0.72   | 0.77     | 79,65%   |
| SVM Tanpa Optimasi | 0.80      | 0.88   | 0.71     | 83.43%   |



Berdasarkan tabel 9 menunjukkan bahwa perbandingan hasil evaluasi model klasifikasi sentimen menggunakan tiga pendekatan berbeda yaitu N-Gram TF-IDF, GloVe Embedding, dan SVM tanpa optimasi. Perbandingan dilakukan berdasarkan empat metrik utama, yaitu precision, recall, F1-score, dan accuracy.

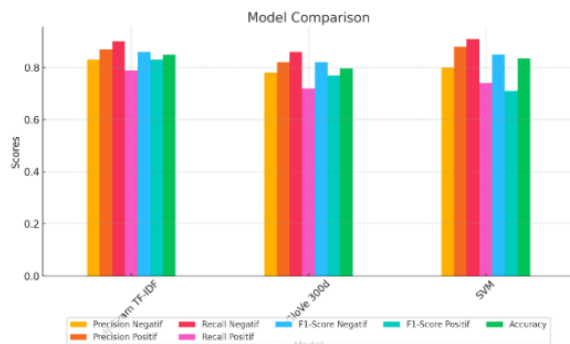
Dari tabel tersebut, terlihat bahwa model dengan N-Gram TF-IDF menghasilkan akurasi tertinggi, yaitu 84.88%, lebih tinggi dibandingkan model dengan GloVe Embedding yang hanya mencapai 79.65%. Model SVM tanpa optimasi juga memiliki akurasi lebih rendah dibandingkan N-Gram TF-IDF, meskipun recall-nya lebih tinggi (0.88), yang menunjukkan bahwa model ini lebih cenderung mengklasifikasikan lebih banyak data ke kelas positif.

Keunggulan N-Gram TF-IDF dapat dijelaskan oleh kemampuannya dalam menangkap pola kata secara lokal melalui unigram, bigram, dan trigram. Teknik ini mampu mengenali susunan kata yang spesifik dalam teks ulasan pengguna, sehingga dapat meningkatkan ketepatan dalam mendeteksi sentimen yang terkandung dalam suatu kalimat. Precision yang lebih tinggi pada model ini (0.87) juga menunjukkan bahwa model lebih baik dalam menghindari kesalahan dalam mengklasifikasikan data negatif sebagai positif. Selain itu, nilai F1-score sebesar 0.82 menunjukkan keseimbangan antara precision dan recall yang baik dalam klasifikasi sentimen.

Meskipun GloVe Embedding unggul dalam menangkap hubungan semantik global antar kata, metode ini memiliki kelemahan dalam menangkap urutan kata dan konteks lokal pada dataset yang berukuran kecil hingga menengah. Hal ini terjadi karena GloVe menggunakan representasi vektor rata-rata untuk kata-kata dalam suatu teks, yang dapat menghilangkan informasi penting. Sebagai contoh, frasa seperti "tidak buruk" bisa kehilangan maknanya karena kata "tidak" dirata-ratakan dengan kata lainnya, sehingga sentimen negatifnya dapat tereduksi. Nilai recall pada model ini juga lebih rendah dibandingkan metode lainnya (0.72), yang berarti model ini lebih banyak gagal mendeteksi ulasan yang sebenarnya bersentimen negatif.



Pada SVM tanpa optimasi memiliki recall tertinggi (0.88) tetapi F1-score yang lebih rendah (0.71), yang mengindikasikan adanya trade-off antara recall dan precision. Model ini cenderung lebih sensitif terhadap ulasan dengan sentimen tertentu, tetapi kurang seimbang dalam mengklasifikasikan kedua kategori secara optimal.



Gambar 10. Visualisasi Hasil Perbandingan  
Sumber: Raka Kusuma Wardana

Visualisasi dari hasil perbandingan dapat dilihat pada Gambar 10, yang menunjukkan perbedaan performa antar model secara lebih jelas. Dengan hasil yang diperoleh, penelitian ini menegaskan bahwa N-Gram TF-IDF merupakan pendekatan yang lebih efektif dibandingkan GloVe untuk analisis sentimen ulasan pengguna aplikasi Access by KAI, terutama dalam konteks dataset dengan ukuran kecil hingga menengah. Keunggulan N-Gram TF-IDF dalam menangkap pola spesifik antar kata menjadikannya pilihan yang lebih optimal dalam meningkatkan performa klasifikasi SVM dalam penelitian ini.

## 5. KESIMPULAN DAN SARAN

Penelitian ini mengoptimalkan metode klasifikasi Support Vector Machine (SVM) dalam analisis sentimen aplikasi Access by KAI dengan membandingkan representasi fitur N-Gram TF-IDF dan GloVe Embedding. Hasil menunjukkan bahwa N-Gram TF-IDF dengan kernel linear memberikan performa terbaik dengan akurasi 84,88%, precision negatif 0,83, precision positif 0,87, recall negatif 0,90, recall positif 0,79, serta F1-score negatif 0,86 dan F1-score positif 0,83. Sementara itu, GloVe 300 Dimensi dengan kernel linear hanya mencapai akurasi 79,65%, dan model SVM tanpa optimasi memperoleh 83,43%. Temuan ini menegaskan bahwa pemilihan teknik representasi fitur berpengaruh signifikan terhadap kinerja klasifikasi sentimen.

Berdasarkan hasil penelitian ini, penelitian selanjutnya dapat dilakukan dengan memperluas dataset agar lebih representatif serta menguji model deep learning seperti LSTM atau BERT untuk meningkatkan akurasi. Penggunaan domain-specific embeddings yang lebih sesuai dengan konteks data ulasan aplikasi juga dapat dieksplorasi. Selain itu, metode pelabelan yang lebih efisien, seperti semi-supervised learning atau active learning, dapat

diterapkan untuk mengurangi ketergantungan pada pelabelan manual. Dengan pengembangan lebih lanjut, hasil analisis sentimen yang lebih akurat diharapkan dapat mendukung peningkatan layanan aplikasi Access by KAI serta memperluas wawasan dalam penelitian analisis sentimen.

## DAFTAR PUSTAKA

- [1] G. Radiana and A. Nugroho, "Analisis Sentimen Berbasis Aspek Pada Ulasan Aplikasi Kai Access Menggunakan Metode Support Vector Machine," *J. Pendidik. Teknol. Inf.*, vol. 6, no. 1, pp. 1–10, 2023, doi: 10.37792/jukanti.v6i1.836.
- [2] N. B. Sidauruk and N. Riza, "SENTIMEN ANALISIS DATA PENGGUNA TERHADAP KAI ACCESS Systematic Literature Review," *J. Mhs. Tek. Inform.*, vol. 7, no. 2, pp. 1297–1303, 2023.
- [3] H. -, A. Y. Kuntoro, and T. Asra, "Klasifikasi Keluhan Pengguna Kai Access Untuk Pemesanan Tiket Dengan Algoritma Svm Dan Naïve Bayes," *JIKA (Jurnal Inform.)*, vol. 6, no. 2, p. 161, 2022, doi: 10.31000/jika.v6i2.6187.
- [4] E. Suryati, Styawati, and A. A. Aldino, "Analisis Sentimen Transportasi Online Menggunakan Ekstraksi Fitur Model Word2vec Text Embedding Dan Algoritma Support Vector Machine (SVM)," *J. Teknol. Dan Sist. Inf.*, vol. 4, no. 1, pp. 96–106, 2023, [Online]. Available: <https://doi.org/10.33365/jtsi.v4i1.2445>
- [5] N. Sidauruk, N. Riza, and R. N. Siti Fatonah, "Penggunaan Metode Svm Dan Random Forest Untuk Analisis Sentimen Ulasan Pengguna Terhadap Kai Access Di Google Playstore," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 3, pp. 1901–1906, 2023, doi: 10.36040/jati.v7i3.6899.
- [6] E. Mulyani, F. P. B. Muhamad, and K. A. Cahyanto, "Pengaruh N-Gram terhadap Klasifikasi Buku menggunakan Ekstraksi dan Seleksi Fitur pada Multinomial Naïve Bayes," *J. Media Inform. Budidarma*, vol. 5, no. 1, p. 264, 2021, doi: 10.30865/mib.v5i1.2672.
- [7] D. R. Putri, E. Y. Puspaningrum, H. Maulana, U. Pembangunan, N. Veteran, and J. Timur, "Indonesia Neural Network," vol. 12, no. 3, pp. 2759–2769, 2024.
- [8] R. Tineges, A. Triayudi, and I. D. Sholihati, "Analisis Sentimen Terhadap Layanan Indihome Berdasarkan Twitter Dengan Metode Klasifikasi Support Vector Machine (SVM)," *J. Media Inform. Budidarma*, vol. 4, no. 3, p. 650, 2020, doi: 10.30865/mib.v4i3.2181.
- [9] E. Rizqi Mar'atus Sholihah, I. G. Susrama Mas Diyasa, and E. Yulia Puspaningrum, "Perbandingan Kinerja Kernel Linear Dan Rbf Support Vector Machine Untuk Analisis Sentimen Ulasan Pengguna Kai Access Pada Google Play Store," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 728–733, 2024, doi: 10.36040/jati.v8i1.8800.

- [10] N. A. Jiana and B. Hartono, "Sentimen Analisa Ulasan Aplikasi Access by KAI pada Google Play Store menggunakan Algoritma K-NN," *J. Media Inform. Budidarma*, vol. 8, no. 3, p. 1388, 2024, doi: 10.30865/mib.v8i3.7730.
- [11] Rima Nurmalah, Muhibudin Wijaya Laksana, and Hilma Mutiara Winata, "Pengaruh Inovasi Pada Aplikasi Access by KAI Terhadap Kualitas Pelayanan," *PUBLIKA J. Ilmu Adm. Publik*, vol. 10, no. 1, pp. 137–149, 2024, doi: 10.25299/jiap.2024.16859.
- [12] S. Rabbani, D. Safitri, N. Rahmadhani, A. A. F. Sani, and M. K. Anam, "Perbandingan Evaluasi Kernel SVM untuk Klasifikasi Sentimen dalam Analisis Kenaikan Harga BBM," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 3, no. 2, pp. 153–160, 2023, doi: 10.57152/malcom.v3i2.897.
- [13] A. M. Priyatno and F. I. Firmanand, "Fitur n-gram untuk perbandingan metode machine learning pada sentimen judul berita keuangan," vol. 1, no. 1, pp. 1–6, 2022.
- [14] M. S. Ibrahim, S. Gauch, T. Gerth, and B. Cox, "WOVe: Incorporating Word Order in GloVe Word Embeddings," *Int. J. Eng. Sci. Technol.*, vol. 4, no. 2, pp. 124–129, 2022, doi: 10.46328/ijonest.83.
- [15] A. R. Fitriansyah and Y. Sibaroni, "Analisis Sentimen Terhadap Pembangunan Kereta Cepat Jakarta-Bandung Pada Media Sosial Twitter Menggunakan Metode SVM dan GloVe Word Embedding," *e-Proceeding Eng.*, vol. 10, no. 2, p. 1713, 2023.
- [16] R. Maulana, A. Voutama, and T. Ridwan, "Analisis Sentimen Ulasan Aplikasi MyPertamina pada Google Play Store menggunakan Algoritma NBC," *J. Teknol. Terpadu*, vol. 9, no. 1, pp. 42–48, 2023, doi: 10.54914/jtt.v9i1.609.
- [17] A. S. Pamungkas and N. Cahyono, "Analisis Sentimen Review ChatGPT di Play Store menggunakan Support Vector Machine dan K-Nearest Neighbor," *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 1–10, 2024, doi: 10.29408/edumatic.v8i1.24114.
- [18] R. Rakarahayu Putri and N. Cahyono, "Analisis Sentimen Komentar Masyarakat Terhadap Pelayanan Publik Pemerintah Dki Jakarta Dengan Algoritma Super Vector Machine Dan Naive Bayes," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 2, pp. 2363–2371, 2024, doi: 10.36040/jati.v8i2.9472.
- [19] G. A. Joni, N. Cahyono, A. Baita, and N. Aini, "RANGKA ESAF HONDA MENGGUNAKAN ALGORITMA NAIVE BAYES DAN SUPER VECTOR MACHINE," vol. 8, no. 5, pp. 10879–10885, 2024.
- [20] D. Setiyawati and N. Cahyono, "Analisa Sentimen Pengguna Sosial Media Twitter Terhadap Perokok di Indonesia," *Indonesian Journal of Computer Science*, vol. 12, no. 1, pp. 262–272, 2023. DOI: <https://doi.org/10.33022/ijcs.v12i1.3154>.
- [21] F. Baehaqi and N. Cahyono, "Analisis Sentimen Terhadap Cyberbullying Pada Komentar Di Instagram Menggunakan Algoritma Naive Bayes," *Indonesia. J. Comput. Sci.*, vol. 13, no. 1, pp. 1051–1063, 2024, doi: 10.33022/ijcs.v13i1.3301.
- [22] N. Cahyono and A. O. Praneswara, "Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes," *Indonesian Journal of Computer Science*, vol. 12, no. 6, Dec. 2023, doi: 10.33022/ijcs.v12i6.3473.
- [23] B. Wicaksono and N. Cahyono, "Analisis Sentimen Komentar Instagram Pada Program Kampus Merdeka Dengan Algoritma Naive Bayes Dan Decision Tree," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 2, pp. 2372–2381, 2024, doi: 10.36040/jati.v8i2.9473.
- [24] R. G. A. Pratama and N. Cahyono, "Comparison of Deep Learning Methods on Sentiment Analysis Using Word Embedding," *JITK*, vol. 10, no. 1, pp. 1–8, Jul. 2024. DOI: <https://doi.org/10.33480/jitk.v10i1.5280>
- [25] A. N. Safira, E. Pujastuti, H. Hanafi, B. Setiaji, D. Prabowo, and N. Cahyono, "Sentiment Analysis to Find Out Positive or Negative Opinions on Ride Hailing Application," in 2023 International Conference on Informatics, Multimedia, Cyber and Informations System (ICIMCIS), Jakarta Selatan, Indonesia, 2023, pp. 341–345, doi: 10.1109/ICIMCIS60089.2023.10349083.