

Analisis Sentimen pada Ulasan Aplikasi My TelU Menggunakan Perbandingan Metode *Naïve Bayes* dan *Support Vector Machine* (Studi Kasus: *Google Play Store*)

1st Zidane Aldani Fitrah Ramadhan
Fakultas Informatika

Universitas Telkom Purwokerto
Banyumas, Indonesia

zidanealdanifr@student.telkomuniversity.ac.id

2nd Nicolaus Euclides Wahyu Nugroho
Fakultas Informatika

Universitas Telkom Purwokerto
Banyumas, Indonesia

nicolausn@telkomuniversity.ac.id

3rd Annisaa Utami
Fakultas Informatika

Universitas Telkom Purwokerto
Banyumas, Indonesia

annisaa@telkomuniversity.ac.id

Abstrak - Implementasi teknologi informasi di lingkungan perguruan tinggi merupakan langkah strategis untuk meningkatkan mutu layanan akademik yang terintegrasi dan responsif. Universitas Telkom memanfaatkan aplikasi My TelU sebagai sarana pendukung berbagai aktivitas akademik mahasiswa. Namun, dalam implementasinya, aplikasi tersebut masih menghadapi sejumlah kendala teknis yang berpotensi memengaruhi kualitas layanan akademik digital serta persepsi pengguna. Ulasan pengguna pada platform Google Play Store menjadi sumber informasi penting untuk mengidentifikasi permasalahan tersebut, tetapi pengolahan ulasan secara manual dan tidak terstruktur berisiko menimbulkan bias dalam evaluasi kualitas layanan.

Oleh karena itu, penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna aplikasi My TelU guna menggambarkan persepsi pengguna secara objektif, sekaligus membandingkan kinerja algoritma *Naïve Bayes* dan *Support Vector Machine* (SVM) dalam mengklasifikasikan sentimen. *Dataset* yang digunakan terdiri dari 139 ulasan pengguna yang telah melalui tahap pra-pemrosesan, meliputi *case folding*, *normalization*, *tokenizing*, dan *stemming*. Evaluasi performa model dilakukan menggunakan metode *K-Fold Cross-Validation* dengan skenario $K = 3, 5$, dan 10 .

Hasil pengujian menunjukkan bahwa algoritma *Naïve Bayes* pada skenario $K = 10$ menghasilkan performa terbaik dengan tingkat akurasi sebesar 91,98%, lebih tinggi dibandingkan performa tertinggi algoritma SVM sebesar 89,84%. Temuan ini menunjukkan bahwa *Naïve Bayes* lebih efektif dalam mengklasifikasikan sentimen ulasan pengguna pada aplikasi My TelU.

Kata kunci – analisis sentimen, *Naïve Bayes*, *Support Vector Machine*, My TelU, Universitas Telkom.

I. PENDAHULUAN

Sebagai lembaga yang fokus pada edukasi, perguruan tinggi turut mengalami hambatan dalam penyediaan layanan akademik berbasis digital [1]. Di tengah era globalisasi yang memacu kecepatan informasi lewat kemajuan teknologi, peran teknologi informasi menjadi sangat vital dalam mendukung penyediaan layanan akademik yang lebih efektif dan responsif [2]. Pemanfaatan teknologi ini memungkinkan sistem informasi menarik data terkini (real-time), yang pada

akhirnya membuat alur pekerjaan menjadi lebih optimal dan tepat guna dalam mendukung layanan akademik kepada mahasiswa [3].

Salah satu contoh penerapan teknologi di perguruan tinggi adalah aplikasi My TelU dari Universitas Telkom. Pemanfaatan aplikasi My TelU menggambarkan peran teknologi sebagai alat penting dalam pengelolaan aktivitas, terutama di era digital yang menuntut layanan pendidikan tinggi semakin terintegrasi dan responsif. Namun, dalam implementasinya, seringkali terjadi masih ditemukan berbagai kendala teknis yang berpotensi memengaruhi kualitas layanan akademik digital dan tingkat kepuasan pengguna. Seiring dengan meningkatnya penggunaan aplikasi My TelU, ulasan pengguna di platform seperti Google Play Store mencerminkan pengalaman dan persepsi pengguna terhadap kualitas layanan akademik yang disediakan oleh aplikasi My TelU.

Sebagai kanal distribusi utama dalam ekosistem Android, platform ini tidak hanya berfungsi sebagai media pengunduhan aplikasi, tetapi juga memfasilitasi interaksi pengguna melalui kolom ulasan [4]. Fitur ini memungkinkan pengguna untuk memberikan opini jujur mengenai performa My TelU, yang secara tidak langsung menjadi data krusial untuk mengevaluasi keunggulan maupun kelemahan layanan yang tersedia. Dengan menganalisis sentimen dari ulasan tersebut, dapat dipahami persepsi pengguna terhadap aplikasi serta langkah perbaikan yang diperlukan. Analisis sentimen merupakan proses yang digunakan untuk mengklasifikasikan dan mengidentifikasi polaritas teks, sehingga memudahkan dalam menentukan kategorinya sebagai positif atau negatif [5].

Dalam penelitian ini, analisis sentimen dilakukan dengan membandingkan dua algoritma populer dalam machine learning. *Naïve Bayes* merupakan pengklasifikasi probabilistik dengan cara menjumlahkan kombinasi nilai frekuensi dari kumpulan data. Algoritma ini bekerja dengan menghitung nilai probabilitas untuk setiap label kelas input berdasarkan Teorema Bayes, dengan asumsi dasar bahwa setiap atribut bersifat independen atau tidak terpengaruh oleh atribut lainnya [6]. Di sisi lain, *Support Vector Machine* (SVM) menawarkan pendekatan berbeda dalam kerangka supervised learning. Mekanisme utama SVM

berfokus pada penentuan fungsi pemisah optimal untuk membedakan kelas data [7].

Penelitian terkait mengenai analisis sentimen pada aplikasi My TelU sebelumnya telah dilakukan oleh Prabaswara [8] dengan menggunakan algoritma *Naïve Bayes*. Penelitian tersebut berfokus pada evaluasi aplikasi menggunakan ulasan Google Play Store dengan menerapkan teknik *oversampling* (SMOTE) untuk menyeimbangkan data, dan menghasilkan performa akurasi yang sangat baik. Meskipun demikian, penelitian tersebut terbatas pada penggunaan satu algoritma klasifikasi tunggal sehingga belum memberikan gambaran perbandingan kinerja algoritma dalam mengklasifikasikan ulasan pengguna. Oleh karena itu, diperlukan penelitian lanjutan yang menerapkan pendekatan komparatif untuk menguji konsistensi performa *Naïve Bayes* jika dibandingkan dengan algoritma lain seperti *Support Vector Machine* (SVM) pada kondisi data ulasan yang organik guna mendapatkan validasi model yang lebih objektif.

Proses akuisisi dataset dilakukan melalui ekosistem Python dengan memanfaatkan modul *google-play-scraper*. Pustaka ini berfungsi sebagai jembatan untuk menarik data publik yang tersedia di laman Google Play Store secara otomatis [9].

Analisis kualitas layanan tidak hanya memberikan informasi mengenai mutu layanan, tetapi juga membantu dalam mengidentifikasi masalah, hambatan, serta membandingkan kinerja pelayanan sebelum dan sesudah perbaikan. Oleh karena itu, analisis kualitas layanan akademik di era yang kompetitif ini sangat penting untuk dilakukan [10]. Berdasarkan urgensi tersebut, penelitian ini bertujuan untuk menggali informasi mendalam tentang persepsi pengguna terhadap layanan My TelU, sekaligus membandingkan kinerja metode *Naïve Bayes* dan *Support Vector Machine* (SVM) guna menentukan model klasifikasi yang paling optimal serta mengimplementasikannya ke dalam dashboard berbasis *web* Streamlit agar hasil analisis dapat diakses secara mudah oleh pengguna.

II. KAJIAN TEORI

A. Analisis Sentimen

Dalam konteks tekstual, sentimen merepresentasikan manifestasi subjektivitas penulis yang mencakup pandangan pribadi, luapan emosi, hingga sikap tertentu. Meskipun data ini pada dasarnya berbentuk informasi mentah yang tidak terstruktur (*unstructured*) dan terpecah (*scattered*), penerapan metode pengolahan yang tepat dapat mengonversi data tersebut menjadi wawasan (*insight*) yang lebih bermakna [21].

Analisis sentimen adalah studi komputasi yang mengumpulkan dan mengevaluasi opini individu yang diekspresikan dalam bentuk teks [22]. Analisis sentimen merupakan bidang yang krusial dalam ilmu data, bertujuan untuk memahami pandangan umum pengguna mengenai berbagai topik. Dengan menerapkan analisis ini, dapat diidentifikasi tingkat kepuasan dan persepsi pengguna terkait produk, layanan, atau isu lainnya yang relevan [23].

B. Ulasan Pengguna

Layanan yang berkualitas tinggi berperan penting dalam meningkatkan kepercayaan serta loyalitas pengguna. Untuk mengukur kepuasan ini, survei pengguna biasanya

menggunakan kuesioner. Namun, pendekatan ini memiliki kelemahan berupa potensi bias dalam jawaban responden. Oleh karena itu, analisis data faktual seperti ulasan di media sosial atau aplikasi menjadi alternatif yang lebih objektif untuk mengetahui tingkat kepuasan pengguna [24].

Ulasan pengguna memiliki dampak signifikan terhadap bagaimana orang memandang aplikasi *mobile* dan memberikan informasi berharga bagi pengembang tentang cara meningkatkan fungsi dan kualitas produk mereka [25].

C. Metode *Naïve Bayes*

Metode klasifikasi *Naïve Bayes* dikenal sebagai metode klasifikasi yang sederhana, namun menawarkan keunggulan berupa kecepatan serta tingkat akurasi yang tinggi. Algoritma ini terbukti mampu beroperasi dengan baik, bahkan pada kondisi *dataset* yang memiliki ketergantungan fitur yang kuat [26].

Persamaan Teorema Bayes menggunakan persamaan (2.2).

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)} \quad (2.2)$$

$P(C|X)$: Probabilitas hipotesis berdasar kondisi
 $P(C)$: Probabilitas hipotesis (*prior probability*)
 $P(X|C)$: Probabilitas berdasarkan hipotesis
 $P(X)$: Probabilitas X [27]

D. Metode *Support Vector Machine*

Sebagai bagian dari algoritma *supervised learning*, *Support Vector Machine* (SVM) bekerja dengan prinsip pencarian *hyperplane* atau fungsi pemisah yang paling optimal. Keunggulan metode ini terletak pada fleksibilitasnya dalam memetakan data, baik secara linear maupun non-linear, di mana tingkat akurasi model sangat bergantung pada ketepatan konfigurasi parameter serta fungsi *kernel* yang diterapkan [28].

Adapun fungsi matematis yang diterapkan oleh metode SVM dalam penelitian ini untuk menentukan klasifikasi dapat dilihat pada persamaan (2.3).

$$f(xd) = \sum_{i=1}^{ns} (a_i y_i (x_i \cdot x_d) + b) \quad (2.3)$$

$f(xd)$: Ini adalah skor keputusan untuk data baru (xd)
 $\sum$: Simbol sigma, artinya penjumlahan
 ns : Jumlah total *support vector*.
 a_i : Bobot dari setiap *support vector* ke- i .
 y_i : Label kelas dari *support vector* ke- i
 x_i : Vektor fitur dari *support vector* ke- i
 x_d : Vektor fitur dari data baru.
 b : Bias, nilai ambang batas *hyperlane* [29]

E. Google Play Store

Google Play Store merupakan layanan penyedia konten digital dari Google yang menyediakan berbagai produk toko daring, mulai dari aplikasi, gim, film, musik, hingga buku. Salah satu fitur utama yang tersedia dalam Google Play Store adalah fasilitas pemberian *rating* dan ulasan (*review*) dari pengguna terhadap aplikasi atau layanan yang telah digunakan [30].

F. *Natural Language Processing*

Analisis sentimen merupakan cabang dari *Natural Language Processing* (NLP) yang bertujuan untuk mengembangkan sistem yang dapat mengenali dan mengekstraksi opini dalam bentuk teks [31].

Dalam kerangka *Natural Language Processing* (NLP), terdapat dua pilar fundamental yang menjadi landasan pemahaman bahasa, yakni analisis sintaksis (*syntactic analysis*) serta analisis semantik (*semantic analysis*). Kedua metode ini diterapkan untuk memvalidasi konstruksi bahasa yang diolah. Secara spesifik, analisis sintaksis menitikberatkan pada aspek aturan tata bahasa (gramatikal), sementara analisis semantik berfokus pada penggalian makna atau interpretasi dari kalimat tersebut [32].

G. Text Pre-processing

Sebelum masuk ke tahap inti, data mentah perlu melalui fase *pre-processing* guna memastikan kualitas input yang bersih dan terstruktur. Rangkaian proses ini melibatkan empat tahapan krusial, yakni *case folding*, *normalization*, *tokenizing*, serta *stemming*. Adapun rincian mekanisme dalam penelitian ini adalah sebagai berikut:

1. Case Folding

Tahap ini bertujuan untuk menstandarisasi format huruf dalam dokumen. Seluruh karakter kapital akan dikonversi secara otomatis menjadi huruf kecil (*lower case*) dari rentang 'A' hingga 'Z' [33].

2. Normalization

Normalization berfungsi untuk mengubah teks yang tidak biasa menjadi format yang lazim, untuk memanfaatkan secara efektif sejumlah besar teks informal guna meningkatkan pembangunan model klasifikasi [34].

3. Tokenizing

Mekanisme ini bekerja dengan memecah untaian kalimat menjadi unit-unit kata terpisah yang dikenal sebagai *token*. Proses pemotongan dilakukan dengan mengeliminasi karakter non-alfanumerik, namun tetap mempertahankan karakter khusus seperti "/" dan "-" yang relevan dalam tata bahasa Indonesia [35].

4. Stemming

Stemming merupakan prosedur untuk membuat suatu kata menjadi kata dasar, dengan menghilangkan imbuhan yang ada pada kata tersebut [36].

G. TF-IDF

Dalam tahapan ekstraksi fitur, proses pembobotan kata (*term weighting*) dilakukan dengan menerapkan algoritma *Term Frequency - Inverse Document Frequency* (TF-IDF). Pendekatan statistik ini bekerja dengan cara mengkalkulasi nilai penting suatu kata berdasarkan dua variabel utama, yakni intensitas kemunculan kata dalam dokumen spesifik (*Term Frequency*) dan tingkat keunikan kata tersebut di seluruh koleksi data (*Inverse Document Frequency*). Semakin sering sebuah kata muncul dalam satu dokumen, maka semakin tinggi pula tingkat signifikansi atau relevansi kata tersebut dalam dokumen [37].

Nilai pembobotan untuk setiap dokumen tersebut dihitung berdasarkan persamaan (2.1) berikut:

$$W_{i,j} = tf_{i,j} \times \log\left(\frac{N}{df_i}\right) \quad (2.1)$$

$W_{i,j}$: Bobot kata ke- i pada dokumen ke- j
 $tf_{i,j}$: Frekuensi kata ke- i pada dokumen ke- j
 N : Jumlah total dokumen
 df_i : Jumlah dokumen yang mengandung kata ke- i

H. K-Fold Cross Validation

Guna mengukur reliabilitas model secara objektif, teknik *Cross Validation* diterapkan sebagai instrumen evaluasi dengan memisahkan komposisi data menjadi dua entitas, yakni data latih (*training*) dan data uji (*testing*). Pendekatan ini menjadi strategi krusial, terutama saat menghadapi kendala keterbatasan jumlah *dataset*, agar estimasi performa model tetap valid tanpa mengalami bias yang signifikan [15].

Mekanisme evaluasi berjalan dalam siklus iteratif, di mana setiap segmen secara bergantian difungsikan sebagai data uji, sedangkan segmen lainnya dikonsolidasikan untuk melatih model. Prosedur ini terus berulang hingga seluruh segmen mendapatkan giliran evaluasi sebanyak K kali [35].

I. Confusion Matrix

Untuk mengevaluasi ketepatan model secara mendetail, digunakan instrumen *Confusion Matrix*. Matriks ini bekerja dengan cara memetakan perbandingan antara label data yang sebenarnya (*actual class*) dengan hasil klasifikasi yang diprediksi oleh sistem (*predicted class*). Adapun representasi struktur matriks tersebut dapat dilihat pada Tabel 1.

TABEL 1
(Confusion Matrix)

Confusion Matrix		Actual Class	
		Positive	Negative
Predicted Class	Positive	True Positive	False Positive
	Negative	False Negative	True Negative

Pengukuran mencakup tingkat akurasi, presisi, *recall*, serta nilai *F1-score*, yang bertujuan untuk menganalisis dampak setiap skenario *pre-processing* pada performa model analisis sentimen. Akurasi mengindikasikan seberapa sering model berhasil mengklasifikasikan sentimen dengan benar. Rumus perhitungan akurasi dapat dilihat pada persamaan (2.4).

$$\text{Akurasi} = \frac{TP + TN}{TP + FP + TN + FN} \quad (2.4)$$

Presisi mengukur ketepatan dari *classifier*, di mana presisi tinggi berarti terdapat lebih sedikit *false positive* (FP). Rumus presisi dapat dilihat pada persamaan (2.5).

$$\text{Presisi} = \frac{TP}{TP + FP} \quad (2.5)$$

Metrik *Recall*, atau yang kerap diistilahkan sebagai sensitivitas, berfungsi untuk menakar kapabilitas *classifier*

dalam menjangking kembali data yang relevan secara utuh. Semakin tinggi skor yang diperoleh, mengindikasikan bahwa model semakin andal dalam mendeteksi informasi positif yang tersedia di dalam *dataset*. Adapun formulasi matematis untuk menghitung nilai ini merujuk pada persamaan (2.6).

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2.6)$$

Metrik presisi dan *recall* dapat dikombinasikan untuk membentuk *F1-Score*, yang merepresentasikan rata-rata harmonik dari kedua variabel tersebut. Adapun formulasi *F1-Score* ditunjukkan pada persamaan (2.7).

$$F1\text{-Score} = 2 \times \frac{\text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}} \quad (2.7)$$

Metrik *F1-score* ini berguna sebagai ukuran keseimbangan antara presisi dan *recall*, yang lebih tepat dalam evaluasi model khususnya saat data tidak seimbang [36].

J. Word Cloud

Dalam penyajian data tekstual, *Word Cloud* digunakan sebagai metode representasi visual yang memetakan kata-kata ke dalam bidang dua dimensi. Karakteristik utama dari teknik ini adalah variasi ukuran huruf yang ditampilkan, yang secara langsung mengindikasikan frekuensi kemunculan kata tersebut [39].



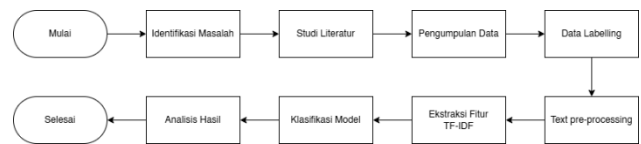
GAMBAR 1
(Contoh Word Cloud)

K. Streamlit

Streamlit merupakan *framework* berbasis Python yang memudahkan pembuatan antarmuka *web* interaktif untuk menampilkan hasil prediksi model *machine learning* secara *real-time* tanpa memerlukan keahlian pengembangan *web* secara mendalam [40].

III. METODE

Penelitian ini dilakukan dengan alur sistematis mulai dari identifikasi masalah, pengumpulan data, hingga evaluasi model, seperti yang ditunjukkan secara komprehensif pada Gambar 1.

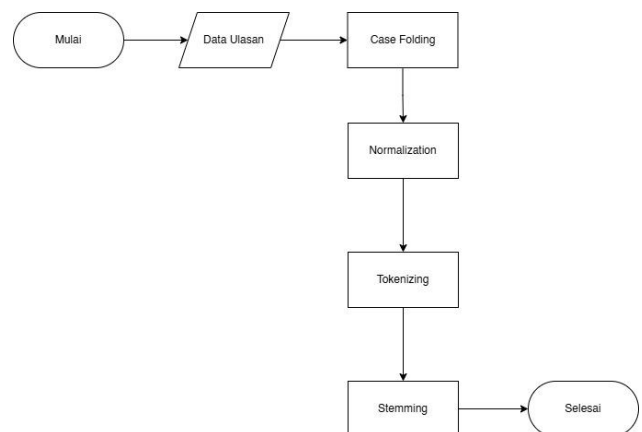


GAMBAR 2
(Diagram Alir Penelitian)

Tahap awal penelitian dimulai dengan Pengumpulan Data. Data ulasan aplikasi My TelU diambil dari platform Google Play Store menggunakan pustaka *google-play-scraper*. Pengambilan data dibatasi pada ulasan yang muncul hingga tanggal 30 November 2025 untuk memastikan relevansi periode penelitian.

Setelah data terkumpul, dilakukan proses Pelabelan Data secara otomatis berdasarkan skor rating yang diberikan pengguna. Ulasan dengan rating 4 dan 5 dikategorikan sebagai sentimen positif (Label 0), sedangkan ulasan dengan rating 1 dan 2 dikategorikan sebagai sentimen negatif (Label 1). Ulasan dengan rating 3 dianggap netral dan tidak digunakan dalam proses pelatihan model untuk menjaga ketegasan polaritas sentimen.

Data yang telah berlabel kemudian masuk ke tahap *Text Pre-processing* untuk membersihkan dan menstandarisasi teks agar siap diolah oleh algoritma *machine learning*. Alur proses ini dapat dilihat pada Gambar 3:



GAMBAR 3
(Diagram Alir Text Pre-processing)

Langkah selanjutnya adalah Klasifikasi dan Evaluasi. Fitur-fitur teks diekstraksi menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*). Klasifikasi kemudian dilakukan dengan membandingkan dua algoritma, yaitu *Naïve Bayes* dan *Support Vector Machine* (SVM). Untuk mengukur kinerja dan validitas model, digunakan metode evaluasi *Stratified K-Fold Cross Validation* dengan variasi nilai $K = 3, 5, 10$, yang menjamin proporsi seimbang antar kelas pada setiap lipatan pengujian.

IV. HASIL DAN PEMBAHASAN

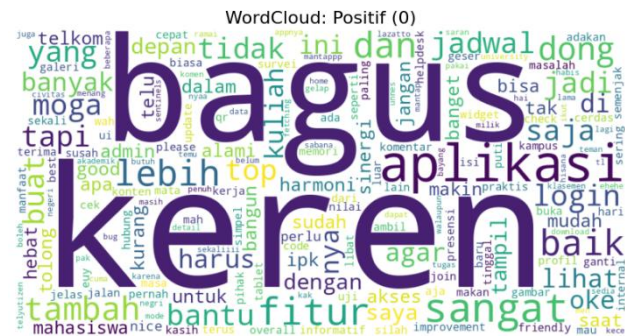
Bab ini menguraikan hasil dari setiap tahapan penelitian yang telah dilaksanakan, mulai dari karakteristik data yang dikumpulkan, hasil visualisasi kata dominan, hingga performa komparatif model klasifikasi. Berikut adalah pembahasan mendalam mengenai temuan-temuan tersebut..

A. Hasil Pengumpulan Data

Dari total data yang dikumpulkan, setelah penyaringan data netral, didapatkan 139 ulasan bersih yang terdiri dari 82 ulasan positif dan 57 ulasan negatif.

B. Visualisasi *Word Cloud*

Untuk memetakan distribusi istilah yang paling sering muncul dalam dataset yang telah bersih, dilakukan visualisasi menggunakan *Word Cloud*. Visualisasi ini merepresentasikan frekuensi kemunculan kata, di mana ukuran teks yang semakin besar menandakan bahwa kata tersebut semakin sering dibahas oleh pengguna dalam ulasannya.



Akurasi: *Naïve Bayes* mencatatkan akurasi tertingginya di angka 91.98%, sedangkan SVM memperoleh akurasi sebesar 89.84%.

Performa Kelas Negatif: *Naïve Bayes* kembali mengungguli SVM secara menyeluruh pada kelas negatif dengan *Precision* 92.50% dan *F1-Score* 89.70%, dibandingkan SVM yang memiliki *Precision* 90.12% dan *F1-Score* 87.21%.

Performa Kelas Positif: Pada sentimen positif, *Naïve Bayes* sangat dominan dengan kemampuan mendeteksi data positif (*Recall*) sebesar 95.14% dan *F1-Score* 93.38%, jauh lebih baik dibandingkan SVM yang mencatat *Recall* 92.78% dan *F1-Score* 91.45%.

D. Analisis Confusion Matrix

Selain analisis metrik di atas, evaluasi diperdalam menggunakan *Confusion Matrix* untuk melihat rincian prediksi benar dan salah pada setiap skenario. Hasilnya disajikan pada Tabel 3 berikut:

TABEL 3
(Confusion Matrix Gabungan)

Skenario K	Model	True Positive (TP)	True Negative (TN)	False Negative (FN)	False Positive (FP)
K = 3	<i>Naïve Bayes</i>	77	48	5	9
	SVM	75	48	7	9
K = 5	<i>Naïve Bayes</i>	77	50	5	7
	SVM	75	49	7	8
K = 10	<i>Naïve Bayes</i>	78	50	4	7
	SVM	76	49	6	8

Berdasarkan Tabel 3, rincian distribusi prediksi benar dan salah untuk setiap skenario pengujian adalah sebagai berikut:

1. Analisis Confusion Matrix Skenario K = 3

Pada skenario ini, *Naïve Bayes* menunjukkan keunggulan dalam mendeteksi sentimen positif, sementara Support Vector Machine (SVM) memiliki tingkat kesalahan yang sedikit lebih tinggi pada data positif yang terlewat.

True Positive (TP): *Naïve Bayes* lebih unggul dengan memprediksi 77 data positif secara benar, sedangkan SVM memprediksi jumlah yang lebih sedikit, yaitu 75 data.

True Negative (TN): Kedua model menunjukkan performa yang setara dalam mengenali sentimen negatif, di mana baik *Naïve Bayes* maupun SVM sama-sama berhasil memprediksi 48 data negatif dengan tepat.

False Negative (FN): *Naïve Bayes* memiliki tingkat kesalahan yang lebih rendah dengan hanya 5 data positif yang dianggap negatif, berbanding terbalik dengan SVM yang melakukan kesalahan pada 7 data.

False Positive (FP): Kedua model mencatatkan jumlah kesalahan yang sama dalam menganggap data negatif sebagai positif, yaitu sebanyak 9 data.

2. Analisis Confusion Matrix Skenario K = 5

Pada skenario ini, *Naïve Bayes* memperlebar keunggulannya dengan mencatatkan performa yang lebih baik daripada *Support Vector Machine* (SVM) di hampir semua aspek matriks kebingungan.

True Positive (TP): *Naïve Bayes* mempertahankan konsistensinya dengan memprediksi 77 data positif secara benar, mengungguli SVM yang mencatat 75 prediksi benar.

True Negative (TN): *Naïve Bayes* lebih baik dalam mengenali sentimen negatif dengan 50 prediksi benar, sedangkan SVM tertinggal tipis dengan 49 prediksi benar.

False Negative (FN): Tingkat kesalahan *Naïve Bayes* lebih minim dengan hanya 5 data *False Negative*, lebih baik dibandingkan SVM yang memiliki 7 data kesalahan.

False Positive (FP): *Naïve Bayes* juga lebih unggul dalam meminimalisir kesalahan prediksi positif palsu dengan 7 data, dibandingkan SVM yang mencatat 8 data kesalahan.

3. Analisis Confusion Matrix Skenario K = 10

Skenario ini membuktikan stabilitas tertinggi dari *Naïve Bayes*, yang secara konsisten mengungguli *Support Vector Machine* (SVM) dalam memprediksi kelas yang benar maupun meminimalisir kesalahan.

True Positive (TP): *Naïve Bayes* mencapai angka tertinggi dalam penelitian ini dengan 78 data benar, lebih tinggi dibandingkan SVM yang memprediksi 76 data.

True Negative (TN): *Naïve Bayes* menjaga performa terbaiknya pada kelas negatif dengan 50 data benar, mengungguli SVM yang memprediksi 49 data.

False Negative (FN): *Naïve Bayes* mencatatkan kesalahan terendah dengan hanya 4 data, menunjukkan sensitivitas yang lebih tinggi dibandingkan SVM yang memiliki kesalahan pada 6 data.

False Positive (FP): *Naïve Bayes* kembali menunjukkan presisi yang lebih baik dengan menekan kesalahan menjadi 7 data, sementara SVM mencatat kesalahan yang lebih tinggi sebanyak 8 data.

Berdasarkan analisis menyeluruh terhadap metrik performa dan rincian *Confusion Matrix* di atas, *Naïve Bayes* dengan skenario K = 10 ditetapkan sebagai model dengan performa terbaik secara keseluruhan. Model ini menghasilkan akurasi tertinggi sebesar 91.98%, dengan kombinasi *True Positive* (78) dan *True Negative* (50) yang maksimal, serta tingkat kesalahan *False Negative* yang paling minimal (4 data) dibandingkan seluruh skenario pengujian lainnya.

E. Implementasi Streamlit

Aplikasi web interaktif Streamlit diimplementasikan sebagai alat pendukung dan dashboard visualisasi untuk menyajikan, memverifikasi, dan mendemonstrasikan hasil dari penelitian ini. Peran streamlit adalah menjembatani logika pemrosesan data dan evaluasi kinerja model agar dapat disajikan secara interaktif.

1. Peran Streamlit dalam Validasi Model

Halaman utama berfungsi sebagai dashboard ringkasan yang menampilkan statistik penting dari data ulasan sebelum proses *pre-processing*. Bagian ini menyajikan metrik penting seperti Total Ulasan, Jumlah Sentimen Positif, dan Jumlah Sentimen Negatif.

2. Integrasi Sistem dan Model Persisten

Implementasi Streamlit menjamin konsistensi antara tahapan pelatihan dan pengujian serta memvalidasi hasil kinerja yang diperoleh:

Pemuatan Model: Model klasifikasi *Naïve Bayes* dan SVM, beserta objek TF-IDF yang telah dilatih, dimuat secara persisten dari file .pkl.

Akses Hasil Evaluasi: Aplikasi memuat hasil evaluasi kinerja model untuk ditampilkan di modul evaluasi model, memungkinkan perbandingan langsung.

3. Fungsionalitas Modul Aplikasi

Fungsionalitas aplikasi dibagi menjadi empat modul utama: Modul *Pre-processing*: Alat visual untuk verifikasi transformasi teks secara bertahap (*Case Folding, Normalization, Tokenizing, Stemming*).

Modul Evaluasi Model: Menyajikan visualisasi interaktif hasil pengujian model (*Akurasi, Presisi, Recall, F1-Score*) dan *Confusion Matrix*.

Modul *Word Cloud*: Memvalidasi fitur-fitur yang tersisa setelah *pre-processing* dan mendukung identifikasi kata dominan.

Modul Klasifikasi *Live*: Mendemonstrasikan kapabilitas model secara *real-time* untuk memprediksi sentimen dari teks baru yang dimasukkan pengguna.

V. KESIMPULAN

Berdasarkan hasil pengujian dan pembahasan yang telah dilakukan pada bab sebelumnya mengenai analisis sentimen aplikasi "My TelU" menggunakan algoritma *Naïve Bayes* dan *Support Vector Machine* (SVM), dapat ditarik beberapa kesimpulan sebagai berikut:

Persepsi pengguna terhadap kualitas layanan akademik digital pada aplikasi My TelU tercermin secara jelas melalui pola sentimen dan visualisasi kata dominan dalam ulasan Google Play Store. Berdasarkan hasil visualisasi *Word Cloud*, sentimen positif didominasi oleh kata-kata seperti "keren", "bagus", "aplikasi", "bantu", "fitur", dan "jadwal", yang menunjukkan bahwa pengguna memersepsikan aplikasi My TelU sebagai alat yang membantu aktivitas akademik, khususnya dalam pengelolaan informasi jadwal dan fitur pendukung akademik lainnya. Sebaliknya, sentimen negatif secara konsisten didominasi oleh kata "tidak", "bisa", dan "login", yang dikelilingi oleh istilah teknis seperti "error", "server", "loading", dan "absen". Pola ini mengindikasikan bahwa ketidakpuasan pengguna terutama berkaitan dengan kendala teknis pada stabilitas server, proses autentikasi (login), serta performa sistem yang lambat. Dengan demikian, persepsi pengguna terhadap My TelU bersifat dualistik, yaitu positif dari sisi fungsionalitas fitur akademik, namun negatif dari sisi keandalan infrastruktur sistem.

Perbandingan kinerja metode *Naïve Bayes* dan *Support Vector Machine* (SVM) menunjukkan bahwa *Naïve Bayes* memiliki performa yang lebih optimal dan konsisten dalam mengklasifikasikan sentimen ulasan pengguna aplikasi My TelU. Berdasarkan hasil pengujian menggunakan *K-Fold Cross-Validation* pada seluruh skenario, algoritma *Naïve Bayes* secara konsisten memperoleh nilai *Accuracy, Precision, Recall*, dan *F1-Score* yang lebih tinggi dibandingkan SVM pada kedua kelas sentimen (positif dan negatif). Performa terbaik diperoleh pada skenario $K = 10$, di mana *Naïve Bayes* mencapai akurasi sebesar 91,98%, dengan nilai *Recall* sentimen positif tertinggi sebesar 95,14%, serta tingkat kesalahan *False Negative* yang paling rendah, yaitu 4 data. Hasil evaluasi melalui *Confusion Matrix* juga menunjukkan bahwa *Naïve Bayes* lebih efektif

dalam mendeteksi ulasan positif maupun negatif secara seimbang dan meminimalkan kesalahan klasifikasi, sehingga lebih sesuai digunakan untuk pemetaan persepsi pengguna pada dataset ulasan aplikasi My TelU dalam penelitian ini.

REFERENSI

- [1] W. Agustiono, M. C. Fajrin, and F. H. Rachman, "Rencana strategi teknologi informasi pada perguruan tinggi di Indonesia: Sebuah tinjauan pustaka," *Sistemasi: Jurnal Sistem Informasi*, vol. 10, no. 1, pp. 197–211, Jan. 2021.
- [2] I. Yuniarto, H. D. Purnomo, and S. Y. J. Prasetyo, "Analisa sistem informasi akademik menggunakan WebQual dan PIECES frameworks pada Universitas XYZ," *J. Media Inform. Budidarma*, vol. 5, no. 3, pp. 987–1007, Jul. 2021, doi: 10.30865/mib.v5i3.3046.
- [3] U. U. Sufandi, "Analisis kebutuhan dan dokumentasi sistem informasi tiras dan transaksi bahan ajar Universitas Terbuka," *Jurnal Nasional Pendidikan Teknik Informatika (JANAPATI)*, vol. 11, no. 2, pp. 112–122, Jul. 2022.
- [4] N. Andriani, B. Warsito, and R. Santoso, "Analisis sentimen aplikasi Microsoft Teams berdasarkan ulasan Google Play Store menggunakan model neural network dengan optimasi Adaptive Moment Estimation (ADAM)," *J. Gaussian*, vol. 13, no. 1, pp. 168–179, 2024, doi: 10.14710/j-gauss.13.1.168-179.
- [5] F. Putra, P. Subandi, F. Romadlon, I. Nurisusilawati, and A. Chindyana, "Sentiment analysis of Indonesian interest in Korean food based on Naïve Bayes algorithm," *J. Inf. Syst. Explor. Res.*, vol. 21, no. 3, pp. 337–346, 2022.
- [6] M. A. Saddam and E. K. D., "Analisis sentimen fenomena PHK massal menggunakan Naive Bayes dan Support Vector Machine," *Jurnal Informatika: Jurnal Pengembangan IT (JPIT)*, vol. 8, no. 3, pp. 226–233, 2023.
- [7] J. Friadi and D. Ely, "Analisis sentimen ulasan wisatawan terhadap Alun-Alun Kota Batam: Perbandingan kinerja metode Naive Bayes dan Support Vector Machine," *Jurnal Infotel*, vol. 14, no. 4, pp. 403–407, 2024, doi: 10.21456/vol14iss4pp403-407.
- [8] R. Prabaswara, "Analisis sentimen publik pada aplikasi My Tel-U menggunakan algoritma Naive Bayes," Tugas Akhir, Universitas Telkom, Bandung, 2024.
- [9] H. Utami, "Analisis sentimen dari aplikasi Shopee Indonesia menggunakan metode Recurrent Neural Network," *Indones. J. Appl. Stat.*, vol. 5, no. 1, pp. 31–38, Mei 2022, doi: 10.13057/ijas.v5i1.56825.
- [10] M. Asfi, "Analisis kepuasan alumni terhadap pelayanan akademik dengan metode Importance Performance Analysis berbasis web (Studi Kasus: Universitas CIC)," *J. Media Inform. Budidarma*, vol. 4, no. 4, pp. 1007–1015, 2020, doi: 10.30865/mib.v4i4.2343.
- [11] M. R. Pratama, F. A. G. S., R. R. Zhafari, Rendy, and H. N. Irmanda, "Sentiment analysis of beauty

- product e-commerce using Support Vector Machine method," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 6, no. 2, pp. 269–274, 2022, doi: 10.29207/resti.v6i2.3920.
- [12] Y. Astuti, Y. Ruldeviyani, F. Salbari, and A. Prayogi, "Sentiment analysis of electricity company service quality using Naïve Bayes," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 7, no. 2, pp. 389–396, 2023, doi: 10.29207/resti.v7i2.4851.
- [13] Kamrozi, A. N. Hidayanto, K. Y. P. M. Yudhakusuma, M. A. Virgananda, and R. R. Suryono, "Sentiment analysis of cryptocurrency trading platform service quality on Playstore data: A case of Indodax," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 7, no. 3, pp. 445–456, 2023, doi: 10.29207/resti.v7i3.4912.
- [14] D. Hardiansyah, R. Z. A. Aziz, and M. S. Hasibuan, "The classification method is used for sentiment analysis in My Telkomsel," *Int. J. Artif. Intell. Res.*, vol. 8, no. 2, pp. 169–179, Des. 2024, doi: 10.29099/ijair.v8i2.1229.
- [15] L. Budi and A. Mude, "Perbandingan metode klasifikasi Support Vector Machine dan Naïve Bayes untuk analisis sentimen pada ulasan tekstual di Google Play Store," *Informatika Mulawarman: Jurnal Ilmiah Ilmu Komputer*, vol. 12, no. 2, pp. 154–161, 2020.
- [16] K. Wabang, O. D. Nurhayati, and Farikhin, "Application of the Naïve Bayes classifier algorithm to classify community complaints," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 6, no. 5, pp. 872–876, 2022, doi: 10.29207/resti.v6i5.4418.
- [17] E. B. Satriawan, S. H. Wijoyo, and D. E. Ratnawati, "Analisis sentimen terhadap pendapat masyarakat mengenai Pilkada 2024 menggunakan metode Support Vector Machine (SVM)," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 1, no. 1, pp. 2548–964, 2024.
- [18] N. S. Ramadan and D. Darwis, "Perbandingan metode Naïve Bayes dan SVM untuk sentimen analisis masyarakat terhadap serangan ransomware pada data KIP-K," *J. Sist. Inf. dan Inform.*, vol. 8, no. 1, pp. 12–23, 2024, doi: 10.47080/simika.v8i1.3621.
- [19] M. Umair and E. R. Sutanto, "Analisis sentimen ulasan pengguna pada aplikasi BRImo BRI menggunakan metode klasifikasi algoritma Naive Bayes," *J. Media Inform. Budidarma*, vol. 8, no. 2, pp. 1149–1159, Apr. 2024, doi: 10.30865/mib.v8i2.7381.
- [20] P. Arsi, R. Wahyudi, and R. Waluyo, "Optimasi SVM berbasis PSO pada analisis sentimen wacana pindah ibu kota Indonesia," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 5, no. 2, pp. 231–237, 2021, doi: 10.29207/resti.v5i2.2843.
- [21] V. R. Prasetyo, N. Benarkah, and V. J. Chrisintha, "Implementasi Natural Language Processing dalam pembuatan chatbot pada program Information Technology Universitas Surabaya," *Teknika*, vol. 10, no. 2, pp. 114–121, 2021, doi: 10.34148/teknika.v10i2.370.
- [22] A. N. Kasanah, Muladi, and U. Pujianto, "Penerapan teknik SMOTE untuk mengatasi imbalance class dalam klasifikasi objektivitas berita online menggunakan algoritma KNN," *Jurnal RESTI (Rekayasa Sist. dan Teknol. Inf.)*, vol. 3, no. 2, pp. 196–201, 2019, doi: 10.29207/resti.v3i2.945.
- [23] F. Muttakin and N. Andrika, "Sentiment analysis of shoe product reviews on Indonesian e-commerce platform using Lexicon Based and Support Vector Machine," *J. Teknol. Inf. dan Komun.*, vol. 6, no. 2, pp. 839–854, 2025.
- [23] A. E. Budiman, "Analisis pengaruh teks preprocessing terhadap deteksi plagiarisme pada dokumen Tugas Akhir," *Sistemasi: Jurnal Sistem Informasi*, vol. 6, pp. 475–488, 2020.
- [24] E. H. Muktafin and E. T. Luthfi, "Analisis sentimen pada ulasan pembelian produk di marketplace Shopee menggunakan pendekatan Natural Language Processing," *Eksplora Informatika*, pp. 32–42, 2020, doi: 10.30864/eksplora.v10i1.390.
- [25] V. Westley, D. Thomas, and F. Rumaisa, "Analisis sentimen ulasan hotel bahasa Indonesia menggunakan Support Vector Machine dan TF-IDF," *J. Media Inform. Budidarma*, vol. 6, no. 3, pp. 1767–1774, 2022, doi: 10.30865/mib.v6i3.4218.