

Perbandingan Kinerja Algoritma Naïve Bayes dan K-Nearest Neighbor dalam Menganalisis Sentimen Pengguna Game Free Fire

Nyoman Dinda Indira Sudiasta Putri¹, I Made Dendi Maysanjaya², I Made Gede Sunarya³

^{1,2,3}Sistem Informasi, Fakultas Teknik dan Kejuruan, Universitas Pendidikan Ganesha
Jalan Udayana No.11, Singaraja

Informasi Naskah:

Diterima: 25 Juli 2025/ Direview: 27 Agustus 2025/ Direvisi: 09 September 2025/ Disetujui Terbit: 09 Oktober 2025

DOI: 10.33369/pseudocode.12.2.53-59

*Korespondensi : dinda.indira@undiksha.ac.id

Abstract

Free Fire is one of the most popular online games in Indonesia, yet it continues to receive a wide range of user reviews regarding gameplay experiences. These reviews reflect diverse user perceptions, including both praise and criticism, making sentiment analysis essential to understanding user satisfaction. This study aims to classify user sentiments toward Free Fire using a combined dataset collected from the Google Play Store and App Store, and to compare the performance of two text classification algorithms: Naive Bayes and K-Nearest Neighbor (KNN). The data were collected using web scraping techniques and manually labeled by expert validators. Text preprocessing involved cleansing, tokenizing, stopword removal, and stemming, followed by term weighting using the Term Frequency-Inverse Document Frequency (TF-IDF) method. The experimental results show that the Naive Bayes algorithm achieved the highest accuracy of 72.78%, while the KNN algorithm recorded a maximum accuracy of 45.91%. Based on these findings, Naive Bayes is proven to be more effective in classifying user sentiments related to Free Fire. The results of this study are expected to provide constructive insights for developers to improve the quality and user experience of the game.

Keywords: Sentiment analysis; Free Fire; Naive Bayes; K-Nearest Neighbor; TF-IDF

1. Pendahuluan

Perkembangan teknologi informasi telah membawa perubahan besar dalam pola hiburan masyarakat modern. Salah satu bentuk hiburan yang kini mendominasi di berbagai kalangan usia adalah game online. Game online telah menjadi hobi yang sangat populer di kalangan muda dan dewasa pada era modern ini. Orang yang memainkan game online atau yang biasa disebut dengan gamers bisa menghabiskan banyak waktu hanya untuk bermain game online tersebut [1]. Aktivitas ini tidak hanya digunakan sebagai sarana melepas penat, tetapi juga telah berkembang menjadi ajang kompetisi dan bahkan profesi bagi sebagian orang.

Salah satu game online yang paling populer saat ini adalah Free Fire. Game ini memiliki basis pemain yang sangat besar dan aktif, baik secara global maupun di Indonesia. Game ini bisa dimainkan di perangkat android dan ios serta bisa dimainkan di semua kalangan, mulai dari anak-anak hingga orang dewasa serta game ini sudah sering ditandingkan pada kanca nasional maupun internasional [2]. Kesuksesan Free Fire ditunjang oleh *gameplay* yang cepat, kemudahan akses lintas perangkat, serta fitur sosial seperti sistem pertemanan dan kompetisi turnamen. Aspek-aspek tersebut menjadikan Free Fire bukan hanya sebagai permainan, tetapi juga sebagai sarana interaksi sosial dan pelatihan kerja sama tim serta

strategi dalam permainan.

Meski demikian, popularitas Free Fire tidak membuatnya lepas dari kritik. Banyak pengguna menyampaikan keluhan terkait kualitas grafis yang dianggap kurang baik, adanya *bug* dan *lag* saat bermain, serta ketidakseimbangan dalam sistem permainan. Ulasan yang diberikan pengguna di Google Play Store dan App Store menunjukkan bahwa meskipun game ini memiliki banyak keunggulan, masih terdapat berbagai aspek yang perlu diperbaiki oleh pengembang. Beberapa ulasan bahkan memperlihatkan kontradiksi antara rating dan komentar, di mana pengguna memberikan rating tinggi namun menyampaikan kritik tajam terhadap game tersebut. Contohnya, ulasan yang memberi bintang lima namun mengeluhkan grafis buruk dan *bug* yang mengganggu. Fenomena ini menunjukkan bahwa terdapat kesenjangan antara penilaian numerik (rating) dan opini sebenarnya dalam komentar pengguna.

Salah satu cara untuk memahami kepuasan pengguna secara lebih objektif adalah dengan melakukan analisis sentimen terhadap ulasan yang mereka berikan. Melalui pendekatan ini, opini yang tersembunyi di balik rating dapat diungkapkan dengan lebih akurat. Beberapa ulasan bahkan mengandung makna ganda atau sarkasme, yang dapat memengaruhi interpretasi sentimen secara manual. Oleh

karena itu, diperlukan pendekatan berbasis kecerdasan buatan untuk memproses data secara sistematis dan efisien.

Dalam bidang analisis sentimen, dua metode yang sering digunakan adalah *Naïve Bayes* dan *K-Nearest Neighbor* (KNN). Beberapa studi sebelumnya juga menunjukkan hasil yang bervariasi terkait efektivitas *Naïve Bayes* dan *K-Nearest Neighbor* dalam analisis sentimen di berbagai bidang seperti e-commerce, transportasi online, investasi, hingga game. Efektivitas kedua metode ini sangat dipengaruhi oleh karakteristik dataset. Karakteristik dataset yang dimaksud mencakup perbedaan sumber data, seperti komentar dari Twitter yang cenderung bersifat singkat, informal, mengandung singkatan, emotikon, tautan (*link*), hingga sarkasme, serta sering kali menyampaikan opini atau persepsi dari berbagai konteks seperti isu terkini, tren, atau pengalaman singkat. Hal ini menjadikan data dari Twitter memiliki tingkat kebisingan (*noise*) yang tinggi dan interpretasi sentimen yang lebih kompleks.

Berbeda dengan data dari Google Play Store atau App Store yang umumnya lebih panjang, deskriptif, dan fokus pada pengalaman langsung pengguna terhadap fitur atau performa aplikasi tertentu sehingga cenderung lebih mudah diproses dan dianalisis. Perbedaan ini dapat memengaruhi efektivitas algoritma karena masing-masing metode memiliki keunggulan tersendiri dalam menangani karakteristik teks tertentu. Oleh karena itu, penelitian ini berupaya mengisi celah tersebut dengan membandingkan kedua algoritma menggunakan sumber data yang sama, serta mengoptimalkan hasil klasifikasi menggunakan TF-IDF, untuk memastikan perbandingan dilakukan secara adil dan komprehensif.

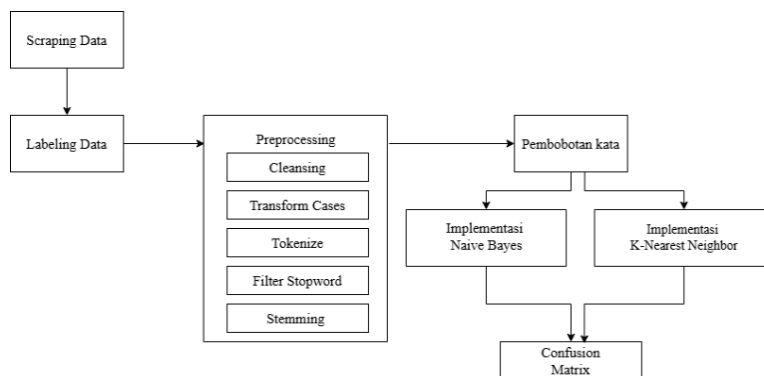
Penelitian ini juga dilakukan karena ditemukan adanya gap dalam dua penelitian sebelumnya mengenai analisis sentimen *Game Free Fire*. Penelitian pertama yang dilakukan oleh Steven [3] menggunakan metode *K-Nearest Neighbor* (KNN) dan memanfaatkan data dari Google Play Store serta App Store. Analisis yang dilakukan bersifat lebih komprehensif karena mempertimbangkan berbagai aspek ulasan pengguna dari kedua platform tersebut. Hasil dari penelitian ini menunjukkan bahwa metode KNN mampu mencapai akurasi sebesar 93,56%. Sementara itu, penelitian kedua yang dilakukan oleh Vera Nauri [4] menggunakan metode *Naïve*

Bayes dan hanya memanfaatkan data dari Twitter. Fokus utama penelitian ini lebih kepada pengukuran sentimen berdasarkan opini pengguna di media sosial. Dari hasil analisis, metode *Naïve Bayes* menunjukkan akurasi sebesar 94,81%, sedikit lebih tinggi dibandingkan dengan metode KNN. Walaupun terdapat selisih sekitar 1%, perbandingan ini semakin menguatkan alasan untuk membandingkan kedua metode tersebut dengan menggunakan sumber data yang sama, yaitu dari Google Play Store dan App Store.

Pendekatan analisis yang berbeda pada kedua penelitian tersebut menunjukkan bahwa ada potensi untuk menggabungkan teknik dari kedua metode guna meningkatkan akurasi dan efektivitas analisis sentimen. Oleh karena itu, penelitian ini tidak hanya bertujuan untuk mengukur kinerja masing-masing metode, tetapi juga untuk mengeksplorasi apakah kombinasi atau adaptasi dari kedua pendekatan tersebut dapat memberikan hasil yang lebih optimal dalam klasifikasi sentimen *Game Free Fire*. Secara khusus, penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna terhadap *Game Free Fire* dari Google Play Store dan App Store guna memahami persepsi pengguna, serta membandingkan efektivitas dan akurasi metode *K-Nearest Neighbor* dan *Naïve Bayes* dalam klasifikasi sentimen, sehingga dapat diketahui metode mana yang lebih unggul dalam mengidentifikasi opini positif dan negatif dari pengguna.

2. Metodologi Penelitian

Metode yang digunakan dalam penelitian ini melibatkan pengumpulan data gabungan berupa ulasan pengguna *Game Free Fire* dari Google Play Store dan App Store. Data yang diperoleh akan diproses melalui tahapan prapemrosesan, meliputi cleansing, transform case, tokenisasi, stopword removal, dan stemming. Selanjutnya, dilakukan pembobotan kata menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Data yang telah dibobot kemudian dianalisis menggunakan dua algoritma klasifikasi, yaitu *K-Nearest Neighbor* dan *Naïve Bayes*, untuk mengklasifikasikan sentimen ke dalam dua kategori: positif dan negatif. Proses analisis ini digambarkan secara rinci pada Gambar 2.1.



Gambar 2.1 Tahapan Analisis Sentimen

2.1. Scraping Data

Data dikumpulkan dari ulasan pengguna *Game Free Fire* di Google Play Store dan App Store menggunakan teknik web scraping melalui library Python seperti BeautifulSoup dan Selenium di platform Google Colab, dengan menerapkan filter tertentu (rentang waktu, bahasa, jumlah data) agar hasil sesuai kebutuhan penelitian dan disimpan dalam format CSV untuk keperluan analisis lebih lanjut. Proses ini dilakukan dengan memanfaatkan keunggulan Google Colab seperti akses fleksibel, integrasi dengan Google Drive, serta dukungan berbagai library, sambil tetap mematuhi etika penelitian dan kebijakan layanan dari kedua platform.

2.2. Labeling Data

Proses pelabelan data dilakukan oleh dua validator yang kompeten, yakni guru Bahasa Indonesia dari SMK Negeri 1 Gerokgak dan SMP Negeri 1 Gerokgak, yang dipilih karena memiliki sertifikasi mengajar sebagai bukti keahlian dalam menganalisis bahasa, dengan tugas mengkategorikan komentar pengguna ke dalam "komentar positif" atau "komentar negatif". Jika terjadi perbedaan pendapat, maka pelabelan akan ditentukan oleh validator ketiga guna memastikan objektivitas, akurasi, dan validitas data untuk analisis selanjutnya.

2.3. Scraping Data

Tahap *preprocessing data* dilakukan menggunakan Python di Google Colab dan mencakup *cleansing* (menghapus karakter selain huruf), *transform case* (mengubah seluruh teks menjadi huruf kecil), *tokenisasi* (memisahkan kalimat menjadi kata-kata), *stopword removal* (menghapus kata umum yang tidak memiliki makna penting seperti "yang", "dan", "di"), serta *stemming* (mengubah kata ke bentuk dasarnya), lalu hasilnya disimpan dalam file `hasil_preprocessing.csv` untuk keperluan pelatihan model klasifikasi.

2.4. Term Frequency-Inverse Document Frequency (TF-IDF)

TF-IDF adalah suatu metode yang bisa digunakan untuk pembobotan kata. *Term Weighting* atau pembobotan kata bertujuan untuk memberikan bobot nilai pada setiap kata. Perhitungan bobot ini memerlukan dua hal yaitu *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF). *Term Frequency* merupakan banyaknya jumlah kata atau term tertentu yang ada dalam suatu dokumen. Sementara *Inverse Document Frequency* adalah frekuensi kemunculan kata atau term pada seluruh dokumen [5].

Perhitungan metode TF dapat dilihat pada persamaan 1, yang menentukan seberapa sering suatu kata muncul dalam dokumen tertentu. Sementara itu, perhitungan metode IDF disajikan pada persamaan 2, yang mengukur seberapa penting suatu kata dalam keseluruhan kumpulan dokumen dengan

mempertimbangkan kelangkaannya.

Perhitungan metode TF

$$tf_{i,j} = \frac{n_{i,j}}{\sum kn_{k,j}} \quad (1)$$

dengan penjelasan:

$n_{i,j}$ adalah jumlah kemunculan kata i dalam dokumen j

$\sum kn_{k,j}$ adalah total kata dalam dokumen j

Perhitungan metode IDF

$$idf_i = \log \frac{|D|}{|\{D_i: t_i \in d_i\}|} \quad (2)$$

dengan penjelasan:

$|D|$ adalah jumlah total dokumen dalam *corpus*

$|\{D_i: t_i \in d_i\}|$ adalah jumlah dokumen yang mengandung kata t_i

Dengan tujuan mendapatkan pembobotan yang sesuai untuk tiap *term* dalam tiap dokumen, maka dilakukan kombinasi metode TF dan IDF. Pembobotan kata menggunakan metode TF-IDF diperoleh dengan mengalikan nilai TF dan IDF, sebagaimana ditunjukkan dalam persamaan 3.

$$tf \times idf_{i,j} = tf_{i,j} idf_i \quad (3)$$

2.5. K-Nearest Neighbor (KNN)

K-Nearest Neighbor (KNN) adalah metode klasifikasi terhadap objek baru berdasarkan data training yang memiliki jarak tetangga terdekat (*nearest neighbor*) dengan objek baru tersebut. Dekat atau jauhnya *neighbor* biasanya dihitung berdasarkan jarak *Euclidean* [6]. Berikut ini adalah langkah-langkah algoritma KNN:

1. Menentukan nilai K , yang dapat dihitung menggunakan persamaan 4.

$$K = \sqrt{N} \quad (4)$$

N menentukan banyaknya sampel pada data training

2. Melakukan perhitungan nilai jarak (*euclidean distance*) terhadap masing-masing objek data yang diberikan. Rumus untuk menghitung *euclidean distance* dapat dilihat pada persamaan 5.

$$d_i = \sqrt{(x_{ki} - x_{kj})^2 + (x_{ki} - x_{kj})^2 + \dots + (x_{ki} - x_{kj})^2} \quad (5)$$

Keterangan:

d_i = jarak *euclidean*

x_{ki} = data training ke-1

x_{kj} = data testing ke-1

3. Melakukan pengelompokkan data sesuai dengan perhitungan jarak (*Euclidean distance*)
4. Melakukan pengelompokkan data sesuai dengan nilai tetangga terdekat (*nearest neighbor*) atau berdasarkan data yang mempunyai jarak *Euclidean* terkecil.

5. Memilih nilai mayoritas dari tetangga terdekat sebagai hasil klasifikasi.

2.6. Naïve Bayes

Naive Bayes merupakan algoritma *machine learning* yang masuk pada kategori *superised classification*. Naive Bayes adalah bentuk paling sederhana dari pengklasifikasi jaringan Bayesian. Pengklasifikasi Naive Bayes berakar dari teorema Bayes, yang mengasumsikan bahwa data tidak berhubungan secara statistik [7].

Rumus hitung data uji dapat dilihat pada persamaan 6 dan persamaan 7:

$$P(v_j) = \frac{|Dok_i|}{|training|} \quad (6)$$

$$P(a_i|v_j) = \frac{n_i+1}{n+|kosakata|} \quad (7)$$

Dengan penjelasan:

$P(v_j)$ adalah probabilitas setiap dokumen pada sekumpulan dokumen

$P(a_i|v_j)$ adalah probabilitas kemunculan kata a_i pada suatu dokumen dengan kategori kelas v_j

$|Dok_i|$ adalah frekuensi dokumen pada setiap kategori

$|training|$ adalah jumlah dokumen training yang ada

n_i adalah frekuensi kata ke- k pada setiap kategori

$|kosakata|$ adalah jumlah kosakata yang ada pada dokumen uji

2.7. Confusion Matrix

Confusion Matrix adalah tabel yang menyatakan klasifikasi jumlah data uji yang benar dan jumlah data uji yang salah [8]. Contoh confusion matrix untuk klasifikasi biner ditunjukkan pada Tabel 2.1.

Tabel 2.1 Kriteria Pelabelan

		Kelas prediksi	
		1	2
Kelas sebenarnya	1	TP	FN
	0	FP	TN

Keterangan:

TP (*True Positive*) = jumlah dokumen dari kelas 1 yang benar diklasifikasikan sebagai kelas 1

TN (*True Negative*) = jumlah dokumen dari kelas 0 yang benar diklasifikasikan sebagai kelas 0

FP (*False Positive*) = jumlah dokumen dari kelas 0 yang salah diklasifikasikan sebagai kelas 1

FN (*False Negative*) = jumlah dokumen dari kelas 1 yang salah diklasifikasikan sebagai kelas 0

Evaluasi yang dilakukan pada *Text Classification* di antaranya dapat menggunakan *accuracy*, *precision*, dan *recall*.

Accuracy adalah tingkat kesesuaian antara nilai prediksi dan nilai aktual yang dapat dilihat pada persamaan 8.

$$accuracy = \frac{TP+TN}{Total} \quad (8)$$

Precision adalah ukuran ketepatan antara respons sistem dengan informasi yang diinginkan oleh pengguna yang dapat dilihat pada persamaan 9.

$$precision = \frac{TP}{TP+FP} \quad (9)$$

Recall adalah ukuran keberhasilan sistem dalam menemukan kembali informasi yang relevan yang dapat dilihat pada persamaan 10.

$$recall = \frac{TP}{TP+FN} \quad (10)$$

3. Hasil dan Pembahasan

Data ulasan pengguna *Game Free Fire* berhasil dikumpulkan dari dua platform aplikasi, yaitu Google Play Store sebanyak 6.000 ulasan dan App Store sebanyak 3.920 ulasan, sehingga total data yang diperoleh sebanyak 9.920 ulasan. Data yang dikumpulkan mencakup isi ulasan, skor, dan tanggal, serta telah difilter berdasarkan bahasa, wilayah, dan jumlah ulasan. Seluruh data kemudian digabungkan dan disimpan dalam format CSV untuk memudahkan proses analisis. Proses pelabelan dilakukan untuk mengklasifikasikan komentar pengguna ke dalam dua kategori, yaitu komentar positif dan komentar negatif, dengan melibatkan dua orang validator yang kompeten di bidang Bahasa Indonesia, masing-masing berasal dari SMK N 1 Gerokgak dan SMP N 1 Gerokgak. Dari hasil pelabelan terhadap seluruh data yang digabungkan, diperoleh 4.026 komentar positif dan 5.894 komentar negatif. Selama proses pelabelan tidak ditemukan perbedaan pendapat antara kedua validator, sehingga tidak diperlukan adanya validator tambahan.

3.1. Klasifikasi

Dalam penelitian ini, dilakukan eksperimen klasifikasi terhadap data ulasan *Game Free Fire* menggunakan dua metode, yaitu K-Nearest Neighbor (KNN) dan Naive Bayes dengan menggunakan satu jenis dataset yang terdiri dari data ulasan pengguna yang telah dikumpulkan dan diproses sebelumnya. Tujuan dari eksperimen ini adalah untuk membandingkan performa kedua metode dalam mengklasifikasikan sentimen dari data ulasan, serta mengevaluasi pengaruh variasi parameter terhadap akurasi klasifikasi. Penelitian ini tidak melakukan balancing data agar distribusi asli tetap terjaga sehingga hasil analisis sentimen lebih relevan dan mencerminkan persepsi pengguna secara autentik. Pada metode KNN, dilakukan delapan skenario eksperimen dengan memvariasikan dua parameter utama,

yaitu jumlah lipatan pada *k-fold cross-validation* (dengan nilai 5 dan 10) serta jumlah tetangga terdekat (*k*) yang diuji dengan nilai 2,3,9, dan 15. Nilai *k* = 2, 3, 9, dan 15 dipilih karena dianggap mewakili variasi jumlah tetangga dari kecil hingga cukup besar sehingga dapat digunakan untuk mengamati sensitivitas model terhadap perubahan parameter serta menemukan konfigurasi terbaik pada masing-masing dataset. Jika menggunakan *k* = 2 berpotensi terjadi tie (jumlah tetangga positif dan negatif sama), mekanisme penanganannya biasanya dilakukan dengan cara memilih kelas secara acak atau berdasarkan prioritas tertentu. Namun, untuk menghindari ambiguitas tersebut, sebaiknya digunakan nilai *k* ganjil sehingga probabilitas tie dapat diminimalisasi. Setiap nilai *k-fold* dikombinasikan dengan empat nilai *k*, sehingga total terdapat delapan skenario pengujian pada metode KNN. Sementara itu, pada metode Naive Bayes hanya dilakukan dua skenario eksperimen karena algoritma ini tidak memiliki parameter jumlah tetangga seperti KNN. Eksperimen ini berfokus pada pengaruh jumlah lipatan terhadap performa

klasifikasi. Untuk lebih jelasnya, skenario eksperimen ditampilkan pada tabel 3.1.

Tabel 3.1 Skenario Pengujian Klasifikasi

Metode	k-fold	Nilai k (KNN)	Total skenario
<i>K-Nearest Neighbor</i>	5	2, 3, 9, 15	4
<i>K-Nearest Neighbor</i>	10	2, 3, 9, 15	4
<i>Naive Bayes</i>	5	-	4
<i>Naive Bayes</i>	10	-	4

Seluruh kombinasi pengujian yang dilakukan menggunakan metode *Naive Bayes* dijabarkan secara rinci dalam Tabel 3.2. Sementara itu seluruh kombinasi pengujian menggunakan metode *K-Nearest Neighbor* (KNN) dijabarkan secara rinci pada Tabel 3.3.

K fold	Accuracy	Precision Positif	Precision negative	Recall Positif	Recal negative	F1 positive	F1 negative
5	72.42%	83%	70%	40%	95%	54%	80%
10	72.78%	83%	70%	41%	94%	55%	80%

Tabel 3.2 Performa Naive Bayes berdasarkan Nilai k-Fold

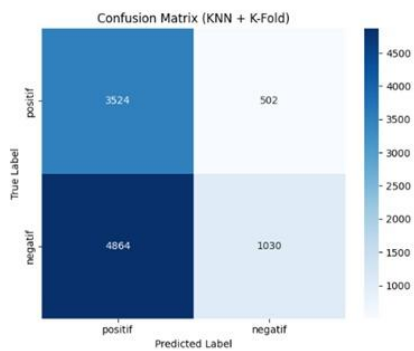
K fold dan k	Accuracy	Precision Positif	Precision negative	Recall Positif	Recal negative	F1 positive	F1 negative
5 (2)	45.82%	42%	68%	88%	17%	57%	27%
10 (2)	45.91%	42%	67%	88%	17%	57%	28%
5 (3)	44.21%	42%	76%	96%	9%	58%	16%
10 (3)	44.55%	42%	76%	96%	10%	58%	17%
5 (9)	43.95%	42%	77%	96%	8%	58%	15%
10 (9)	43.31%	41%	76%	97%	7%	58%	12%
5 (15)	42.55%	41%	81%	98%	4%	58%	8%
10 (15)	42.81%	42%	77%	96%	8%	58%	15%

Tabel 3.3 Performa K-Nearest Neighbor berdasarkan Nilai k-Fold dan k Tetangga

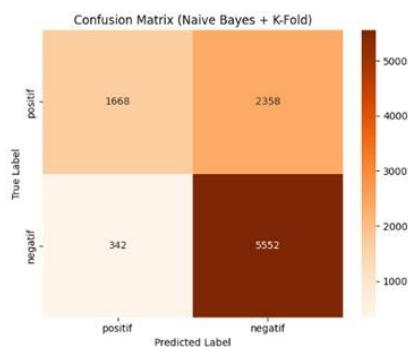
Tabel 3.3 menunjukkan bahwa performa terbaik KNN tercapai pada *k*=2 dan *k-fold*=10, dengan akurasi sebesar 45.91%. *Precision* untuk kelas positif sebesar 42%, dengan *recall* yang sangat tinggi, yakni 88%, yang menghasilkan F1-score positif sebesar 57%. Pada Tabel 3.2, menunjukkan akurasi terbaik sebesar 72.78% saat menggunakan *k-fold*=10. Meskipun *precision positive* berada di angka 83%, *recall positive* hanya mencapai 41%, sehingga F1-score untuk kelas positif berada pada angka 55%.

Sebagai pelengkap dari hasil pengujian yang telah dijabarkan sebelumnya, *confusion matrix* ditampilkan hanya untuk kombinasi parameter terbaik pada metode *K-Nearest Neighbor* (KNN) dan *Naive Bayes* dengan menggunakan

dataset yang telah disiapkan. Untuk metode KNN, kombinasi parameter terbaik diperoleh pada nilai *k* = 2 dengan *k-fold* = 10, sedangkan untuk metode *Naive Bayes*, kombinasi terbaik diperoleh pada *k-fold* = 10. Hasil *confusion matrix* pada kedua metode bisa dilihat pada Gambar 3.1 *K-Nearest Neighbor* dan Gambar 3.2 *Naive Bayes*



Gambar 3.1. Confusion Matrix metode KNN



Gambar 3.2 Confusion Matrix metode Naive Bayes

3.2. Perbandingan metode *K-Nearest Neighbor* dan *Naïve Bayes*

Berdasarkan pengujian terhadap dataset yang digunakan dalam penelitian ini, perbandingan antara metode *K-Nearest Neighbor* (KNN) dan *Naïve Bayes* menunjukkan perbedaan yang cukup signifikan dalam hal akurasi dan efektivitas. Secara umum, metode *Naïve Bayes* menunjukkan akurasi yang lebih tinggi dibandingkan KNN, dengan akurasi tertinggi sebesar 72.78% pada dataset ini, sementara KNN hanya mencapai akurasi maksimal sebesar 45.91%. Dari segi precision terhadap kelas positif, *Naïve Bayes* menunjukkan keunggulan yang jelas, dengan nilai precision mencapai 83%. Namun, kelemahan utama metode ini terletak pada nilai recall yang rendah, yaitu hanya sebesar 41%. Artinya, model *Naïve Bayes* sangat selektif dalam mengenali ulasan positif dan cenderung melewatkan banyak data positif. Berdasarkan confusion matrix, hal ini tercermin dari tingginya jumlah false negative, yakni sebanyak 2.358.

Di sisi lain, metode KNN memiliki kekuatan pada aspek recall yang sangat tinggi, mencapai 88%, yang menunjukkan kemampuannya dalam menjangkau sebagian besar ulasan positif. Hal ini dibuktikan dengan jumlah true positive yang tinggi, yaitu sebanyak 3.524. Namun, tingginya recall ini diikuti dengan kelemahan pada nilai precision yang rendah, yakni hanya 42%, akibat jumlah false positive yang sangat tinggi, yaitu sebesar 4.864. Confusion matrix KNN menunjukkan kecenderungan model untuk mengklasifikasikan

data secara berlebihan ke kelas positif, sehingga berdampak negatif terhadap keakuratan prediksi ulasan yang benar-benar positif. Tingginya recall ini dapat disebabkan oleh dominasi kata-kata umum yang sering muncul pada ulasan positif, seperti “bagus”, “seru”, atau “mantap”, yang juga kadang digunakan pada ulasan negatif dengan konteks berbeda, misalnya “*grafiknya bagus, tapi sering lag*”. Selain itu, distribusi data yang tidak seimbang antara ulasan positif dan negatif membuat tetangga terdekat dalam KNN lebih sering berasal dari kelas positif, sehingga model cenderung memasukkan lebih banyak data ke kelas positif. Hal inilah yang menjelaskan mengapa recall KNN sangat tinggi, tetapi di sisi lain precision menjadi rendah..

4. Kesimpulan

Berdasarkan hasil pengujian terhadap dataset dalam penelitian ini, metode *Naïve Bayes* secara umum menunjukkan performa yang lebih baik dibandingkan *K-Nearest Neighbor* (KNN) dalam hal akurasi dan precision. *Naïve Bayes* mencapai akurasi tertinggi sebesar 72.78% dan precision terhadap kelas positif sebesar 83%, menunjukkan kemampuannya dalam menghasilkan prediksi yang tepat terhadap ulasan positif. Hal ini dapat dijelaskan karena *Naïve Bayes* bekerja dengan pendekatan probabilistik yang mampu menangani data teks berdimensi tinggi serta memanfaatkan distribusi kata-kata yang khas pada ulasan game, sehingga model lebih mudah mengenali kata-kata yang merepresentasikan sentimen positif maupun negatif. Namun, kelemahan metode *Naïve Bayes* ini terletak pada nilai recall yang rendah (41%), yang menandakan bahwa banyak ulasan positif tidak berhasil dikenali (false negative tinggi).

Sebaliknya, metode KNN menunjukkan keunggulan dalam hal recall yang sangat tinggi (88%), dengan jumlah true positive yang besar. Akan tetapi, hal ini tidak diimbangi dengan precision yang baik (hanya 42%), karena model cenderung mengklasifikasikan terlalu banyak data sebagai positif, menghasilkan false positive yang tinggi. Kondisi ini dapat terjadi karena KNN sangat bergantung pada perhitungan jarak antar data, sementara data teks ulasan memiliki dimensi yang tinggi dan distribusi yang tidak seimbang, sehingga model lebih rentan menghasilkan prediksi yang kurang spesifik dan cenderung bias terhadap kelas mayoritas.

Untuk pengembangan selanjutnya, *Naïve Bayes* direkomendasikan karena memberikan akurasi dan precision yang tinggi. Namun, untuk meningkatkan sensitivitas terhadap kelas minoritas, peneliti selanjutnya disarankan mengeksplorasi teknik seperti *threshold tuning*, seleksi fitur, atau *resampling*. Di sisi lain, performa KNN dapat ditingkatkan dengan mencoba variasi metode pencarian tetangga seperti *KD-Tree*, penggunaan fungsi jarak yang berbeda, atau pendekatan ensemble seperti *Bagging-KNN*. Selain itu, penelitian berikutnya juga dapat mempertimbangkan pendekatan hybrid, seperti kombinasi antara *Naïve Bayes* dan KNN, atau menguji algoritma lain

seperti SVM, Random Forest, dan LSTM untuk mendapatkan hasil klasifikasi sentimen yang lebih seimbang dan akurat.

Referensi

- [1] M. O. Saputra, *Communication Patterns on Addicted Elementary School-Age Children in Playing Free Fire Games*, 2023. [Online]. Available: <http://journals.telkomuniversity.ac.id/liski124JurnalIlmiahLISKI>
- [2] S. H. Harahap and Z. H. Ramadan, "Dampak Game Online Free Fire terhadap Hasil Belajar Siswa Sekolah Dasar," *Jurnal Basicedu*, vol. 5, no. 3, pp. 1304–1311, Apr. 2021, doi: 10.31004/basicedu.v5i3.895.
- [3] Steven, *Optimasi Algoritma Klasifikasi K-Nearest Neighbor pada Perbandingan Analisis Sentimen Game Free Fire dan PUBG Mobile*, 2024.
- [4] V. Nuari, *Analisis Sentimen pada X mengenai Game Online PUBG Mobile dan Free Fire di Indonesia Menggunakan Metode Naive Bayes Classifier*, 2024.
- [5] M. K. Maulidina and E. I. Sela, "Analisis Sentimen Komentar Warganet terhadap Postingan Instagram Menggunakan Metode Naive Bayes Classifier dan TF-IDF," 2024.
- [6] H. A. D. Fasnuari, H. Yuana, and M. T. Chulkamdi, "Penerapan Algoritma K-Nearest Neighbor untuk Klasifikasi Penyakit Diabetes Melitus," *Antivirus: Jurnal Ilmiah Teknik Informatika*, vol. 16, no. 2, pp. 133–142, Oct. 2022, doi: 10.35457/antivirus.v16i2.2445.
- [7] R. A. Husen, R. Astuti, L. Marlia, R. Rahmaddeni, and L. Efrizoni, "Analisis Sentimen Opini Publik pada Twitter terhadap Bank BSI Menggunakan Algoritma Machine Learning," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 3, no. 2, pp. 211–218, Oct. 2023, doi: 10.57152/malcom.v3i2.901.
- [8] D. Normawati and S. A. Prayogi, "Implementasi Naïve Bayes Classifier dan Confusion Matrix pada Analisis Sentimen Berbasis Teks pada Twitter," 2021.