

ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI FACEBOOK DI GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA NAÏVE BAYES DAN K-NEAREST NEIGHBOR

Tulus Ramadha Triputra, Albaar Rubhasy

Sistem Informasi, Universitas Nasional

Jl. Sawo Manila No.61, RT.14/RW.7, Pejaten Bar., Ps. Minggu, Jakarta, Indonesia

tulusramadhatriputra.2021@student.unas.ac.id

ABSTRAK

Facebook adalah jejaring sosial yang memudahkan untuk terhubung dan berbagi konten dengan keluarga dan teman secara online. Awalnya facebook ditujukan untuk pelajar, Facebook didirikan oleh Mark Zuckerberg pada tahun 2004 saat kuliah di Universitas Harvard. Pada tahun 2006, siapa pun yang berusia di atas 13 tahun dengan alamat email yang benar dapat bergabung dengan Facebook. Untuk melakukan analisis sentimen terhadap ulasan pengguna aplikasi Facebook di Google Play Store dengan menggunakan algoritma *Naïve Bayes* dan *K-Nearest Neighbor (KNN)* untuk mengetahui persepsi pengguna terhadap aplikasi tersebut. Data ulasan dikumpulkan melalui web scraping dan diproses menggunakan teknik *Natural Language Processing (NLP)*, termasuk pre-processing teks, *TF-IDF vectorization*, dan pembagian dataset dalam beberapa skenario. Evaluasi model dilakukan dengan menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*, serta divisualisasikan dalam bentuk confusion matrix. Hasil penelitian menunjukkan bahwa *Naïve Bayes* memiliki akurasi lebih tinggi (73%), *precision* (51%), *recall* (52%), *F1-Score* (50%), dibandingkan *KNN* (72%), *precision* (55%), *recall* (53%), *F1-Score* (52%), dengan performa lebih baik dalam mengklasifikasikan sentimen positif dan negatif.

Kata kunci : Analisis Sentimen, Naive Bayes, K-Nearest Neighbor, Facebook, Google Play Store

1. PENDAHULUAN

Kemajuan teknologi ini membuka berbagai peluang baru dalam dunia pemasaran. Munculnya berbagai platform online seperti komunitas online, e-commerce, dan media sosial. Media sosial seperti Facebook telah menjadi bagian yang tidak dapat dipisahkan dari kehidupan masyarakat. Facebook adalah jejaring sosial yang memudahkan untuk terhubung dan berbagi konten dengan keluarga dan teman secara online. Awalnya facebook ditujukan untuk pelajar,[1]

Facebook didirikan oleh Mark Zuckerberg pada tahun 2004. Sejauh ini Facebook terus berinovasi untuk memberikan fitur dan layanan yang lebih baik kepada penggunanya [2].

Namun, seiring dengan meningkatnya jumlah pengguna dan fitur yang disediakan, tidak jarang munculnya keluhan dan ketidakpuasan dari pengguna. Hal ini dapat dilihat dari ulasan dan komentar yang ditinggalkan di platform Facebook maupun di Google Play Store. Keluhan yang penulis temukan adalah Gangguan seperti *buffering*, aplikasi yang sering *force-close*, dan fitur yang tidak berfungsi dengan baik dapat menyebabkan pengguna merasa frustrasi dan beralih ke platform lain, Ulasan negatif yang terus bertambah yang menurunkan reputasi Facebook di Google Play Store, sehingga pengguna baru enggan mengunduh aplikasi [3].

Mengacu pada penelitian sebelumnya terkait perbandingan algoritma *Naive Bayes* dan *KNN*. Analisis Sentimen Ulasan Aplikasi Mypertamina.

Penelitian ini menghasilkan bahwa algoritma *Naïve Bayes* mencapai tingkat akurasi yang lebih tinggi yaitu 75% dibandingkan dengan algoritma *K-Nearest Neighbor* yang berkisar antara 56% hingga 73% dalam analisis sentimen ulasan pada aplikasi MyPertamina. Penelitian tersebut menyarankan penelitian lebih lanjut untuk membandingkan *Naïve Bayes* dengan berbagai algoritma klasifikasi lainnya [4].

Berdasarkan latar belakang permasalahan tersebut, penulis tertarik untuk mengidentifikasi dan memahami konsep-konsep penting dalam analisis sentimen, machine learning, dan pemrosesan bahasa alami pada ulasan pengguna aplikasi Facebook, yang diharapkan dapat mengarah pada layanan Facebook yang lebih baik dan sesuai dengan kebutuhan penggunanya. Adapun perbedaan dari penelitian sebelumnya adalah dengan menambahkan satu algoritma yaitu K-Nearest Neighbor dengan tujuan membandingkan dua algoritma yaitu *Naïve Bayes* dan *K-Nearest Neighbor*.

Dari banyaknya ulasan pengguna aplikasi Facebook di Google Play Store menimbulkan tantangan dalam melakukan analisis sentimen, karena data ulasan yang terus bertambah serta beragam. Ulasan-ulasan ini juga menggunakan bahasa alami yang terdiri dari teks, simbol, dan berbagai gaya bahasa yang berbeda, sehingga mempersulit proses klasifikasi otomatis. Sehingga, pengukuran kinerja analisis sentimen juga menjadi sulit jika tidak ada metode evaluasi yang tepat untuk menilai keakuratan hasil klasifikasi.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Bidang ilmiah yang mempelajari bagaimana opini, keluhan, opini, dan penilaian lainnya dianalisis. Analisis sentimen digunakan untuk menentukan suka dan tidak suka suatu produk atau layanan. Tujuan dari analisis sentimen tidak hanya untuk mengetahui sentimen saja, namun juga untuk mengetahui apakah opini terhadap topik atau objek dalam sebuah teks mempunyai sentimen positif atau negatif [5].

2.2. Data Mining

Memungkinkan pengguna mengakses data dalam jumlah besar dengan cepat. Definisi data mining yang lebih spesifik adalah alat dan aplikasi yang menggunakan analisis data statistik. *Data Mining* melibatkan penggalan sejumlah besar data dan informasi yang sebelumnya tidak diketahui dari database besar sehingga dapat dipahami, dimanfaatkan.[6]

2.3. Natural Language Processing (NLP)

Cabang dari ilmu komputer yang menggabungkan linguistik, ilmu data, dan kecerdasan buatan (AI) untuk memungkinkan komputer memahami, menganalisis, menafsirkan, dan memproses bahasa manusia secara alami. *NLP* bertujuan untuk menjembatani komunikasi antara manusia dan mesin menggunakan bahasa yang dipahami oleh manusia, seperti teks atau ucapan.[7]

2.4. Machine Learning

Salah satu bagian dari kecerdasan buatan yang berfokus pada pengembangan algoritma dan model statistik untuk membuat prediksi serta keputusan tanpa diprogram secara eksplisit [8].

Algoritma pembelajaran mesin menggunakan teknik seperti pengenalan pola dan penambangan data untuk membangun model berdasarkan data yang diberikan. Dengan memasukkan data dalam jumlah besar, algoritma belajar mengenali pola dan hubungan dalam data tersebut serta menerapkannya pada data baru untuk membuat prediksi tanpa aturan yang telah ditentukan sebelumnya.

2.5. Facebook

Facebook merupakan sebuah platform jejaring sosial yang didirikan oleh Mark Zuckerberg dengan teman-temannya, Eduardo Saverin, Andrew McCollum, Dustin Moskovitz, dan Chris Hughes di Universitas Harvard. Facebook pertama kali diluncurkan pada tahun 2004 [9].

2.6. Algoritma Naïve Bayes

Algoritma *Naïve Bayes*, yang juga dikenal sebagai klasifikasi *Naïve Bayes*, adalah algoritma klasifikasi probabilistik yang menghitung probabilitas dengan menjumlahkan frekuensi dan kombinasi nilai dari kumpulan data tertentu. *Naïve Bayes* mengasumsikan bahwa setiap fitur dalam data adalah

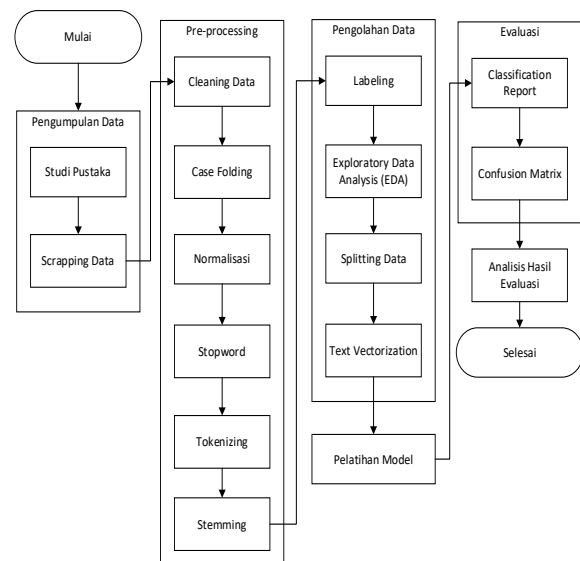
independent (tidak saling bergantung) dan memberikan kontribusi yang sama terhadap probabilitas yang dihitung. Meskipun asumsi independensi fitur ini jarang terpenuhi dalam kasus nyata, algoritma *Naïve Bayes* tetap sering digunakan karena kesederhanaan dan kinerjanya yang baik dalam banyak tugas klasifikasi. [10]

2.7. Algoritma K-Nearest Neighbor

KNN merupakan algoritma supervised learning, di mana algoritma ini menggunakan data pelatihan yang telah berlabel untuk memprediksi kelas atau nilai numerik dari data baru.. Salah satu keunggulan utama algoritma *KNN* adalah kesederhanaan dan kemudahan implementasinya. Algoritma ini tidak memerlukan pelatihan model seperti pada algoritma lain, melainkan hanya menyimpan seluruh data pelatihan dan menghitung jarak saat diperlukan. Selain itu, *KNN* dapat digunakan untuk masalah klasifikasi multi-kelas dan regresi tanpa perlu penyesuaian khusus. [10].

3. METODE PENELITIAN

Pada gambar 1 menunjukkan tahapan dalam penelitian ini.



Gambar 1. Tahapan Penelitian

Pada gambar 1 menunjukkan bahwa penelitian ini dibagi menjadi beberapa tahapan yaitu:

a. Pengumpulan Data

Mengumpulkan ulasan pengguna Facebook di Google Play Store, menggunakan teknik *web scrapping* dengan *library google-play-scraper*.

b. Cleaning Data

Cleaning Data merupakan proses untuk membersihkan data dari kata-kata yang tidak penting seperti *URL*, *emoticon*, dan penghapusan semua karakter selain spasi dan huruf..

c. Case Folding

Pada proses *Case Folding* bertujuan membuat semua huruf dalam teks menjadi huruf kecil

untuk membuat analisis teks menjadi lebih konsisten.

d. Normalisasi

Untuk mengubah teks menjadi lebih rapih, efisien, dan terstruktur.

e. Stopword

Stopword berfungsi untuk menghilangkan kata yang memiliki nilai yang tidak terlalu penting contohnya kata, yang, di, dan pada. Dengan menggunakan *library Sastrawi*.

f. Tokenizing

Tokenizing atau tokenisasi adalah proses membagi teks menjadi unit yang lebih kecil dan lebih mudah untuk diproses.

g. Stemming

Stemming digunakan untuk mengubah beberapa kata menjadi kata dasar dengan menghilangkan imbuhan atau akhir kata.

h. Labeling

Labeling bertujuan untuk mengkategorikan teks atau dokumen ke dalam kategori positif, negatif, dan netral. Pada pelabelan ini ulasan yang memiliki score 5 dan 4 akan diklasifikasikan sebagai positif, sedangkan score 3 diklasifikasikan sebagai netral, dan score 2 dan 1 diklasifikasikan sebagai negatif.

i. Exploratory Data Analysis (EDA)

Proses analisis awal terhadap data untuk memahami karakteristik, pola, dan distribusi data sebelum melakukan pemodelan. EDA dalam penelitian ini dilakukan untuk distribusi kata-kata yang sering muncul, dan mengidentifikasi negatif dalam ulasan pengguna.

j. Stemming

Stemming digunakan untuk mengubah beberapa kata menjadi kata dasar dengan menghilangkan imbuhan atau akhir kata.

k. Splitting Data

Merupakan proses membagi dataset menjadi dua bagian yaitu *subset training* dan *subset testing*. Dengan menggunakan metode *Training-Test Split*.

l. Text Vectorization

Proses mengubah teks menjadi angka numerik sehingga dapat digunakan untuk *model machine learning*. Metode yang dipakai adalah *TF-IDF vectorizer*.

m. Pemodelan

Setelah data siap, langkah selanjutnya adalah melatih model menggunakan algoritma *Naive Bayes* dan *K-Nearest Neighbors (KNN)*. Kedua algoritma ini dapat diterapkan menggunakan *MultinomialNB* (untuk *Naive Bayes*) dan *KNeighborsClassifier* (untuk *KNN*) yang tersedia di *scikit-learn*.

n. Evaluasi

Classification report merupakan proses yang menghasilkan sebuah tabel yang berisi beberapa metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan *f1-score* yang dirumuskan sebagai berikut.

a. Akurasi

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

b. Precision

$$Precision = \frac{TP}{TP + FP}$$

c. Recall

$$Recall = \frac{TP}{TP + FN}$$

d. F1-Score

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

Confusion Matrix Berupa tabel yang menunjukkan bagaimana model klasifikasi memprediksi data ke dalam kategori yang berbeda. Metode *confusion matrix* terdiri dari *True Positive (TP)*, *True Negative (TN)*, *False Positive (FP)*, dan *False Negative (FN)*.

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan Data

Memperoleh data sekunder menggunakan scraping data dengan package name yaitu *com.facebook.katana* untuk mengambil ulasan pengguna aplikasi facebook dengan Bahasa Indonesia dan Region Indonesia. Ulasan yang di dapat merupakan ulasan yang relevan dengan semua rating, lalu ditargetkan mengambil 5000 data ulasan. Setelah berhasil mendapatkan ulasan yang berjumlah 5000 ulasan, kemudian mensortir kolom dengan hanya ulasan dan *rating/score* nya saja

4.2. Cleaning Data

Langkah Pre-Processing yang pertama adalah cleaning data proses ini untuk membersihkan ulasan agar lebih siap digunakan untuk analisis sentimen proses ini untuk menghilangkan *URL*, emoji, angka, tanda baca, spasi berlebih, dan spasi di awal/akhir. Pada Gambar 2 berikut.

	content	score	cleaning
0	Tolong di fix, ngedit video pake text di reels...	3	Tolong di fix ngedit video pake text di reels ...
1	Facebook, Please kembalikan fitur edit; untuk ...	5	Facebook Please kembalikan fitur edit untuk me...
2	Aplikasi populer dan masih mampu bertahan. Mas...	4	Aplikasi populer dan masih mampu bertahan Masi...
3	Ini kenapa hobi banget error ya? Saya ga terken...	4	Ini kenapa hobi banget error ya Saya ga terkena...
4	Kenapa ini Facebook setelah update malah makin...	2	Kenapa ini Facebook setelah update malah makin...
5	mohon kembalikan lagi durasi ketika mengunggah...	3	mohon kembalikan lagi durasi ketika mengunggah...
6	terlalu banyak pembatasan padahal gaada pelang...	3	terlalu banyak pembatasan padahal gaada pelang...
7	setelah update terbaru, saat liat story video ...	5	setelah update terbaru saat liat story video a...
8	Terbaik memang fitur2 Facebook yang baru. Teru...	5	Terbaik memang fitur Facebook yang baru Teruta...
9	Upload Video di Reel di FB Thumbnail / sampul ...	2	Upload Video di Reel di FB Thumbnail sampul vl...

Gambar 2. Hasil Cleaning Data

4.3. Case Folding

Pada proses ini bertujuan untuk mengubah ulasan pengguna menjadi huruf kecil (lowercase), hal ini dilakukan untuk menghindari perbedaan akibat huruf besar dan kecil dalam teks. Pada Gambar 3 berikut.

cleaning	case_folding
Tolong di fix ngedit video pake text di reels ...	tolong di fix ngedit video pake text di reels ...
Facebook Please kembalikan fitur edit untuk me...	facebook please kembalikan fitur edit untuk me...
Aplikasi populer dan masih mampu bertahan Masi...	aplikasi populer dan masih mampu bertahan masi...
Ini kenapa hobi banget error ya Saya ga terkena...	ini kenapa hobi banget error ya saya ga terkena...
Kenapa ini Facebook setelah update malah makin...	kenapa ini facebook setelah update malah makin...
mohon kembalikan lagi durasi ketika mengunggah...	mohon kembalikan lagi durasi ketika mengunggah...
terlalu banyak pembatasan padahal gaada pelang...	terlalu banyak pembatasan padahal gaada pelang...
setelah update terbaru saat liat story video a...	setelah update terbaru saat liat story video a...
Terbaik memang fitur Facebook yang baru Teruta...	terbaik memang fitur facebook yang baru teruta...
Upload Video di Reel di FB Thumbnail sampul vi...	upload video di reel di fb thumbnail sampul vi...

Gambar 3. Hasil Case Folding

4.4. Normalisasi

Dilakukannya normalisasi teks bertujuan untuk mengganti kata-kata yang tidak baku atau *slang* menjadi kata yang baku agar lebih mudah diproses oleh model *machine learning*. Pada Gambar 4 berikut.

case_folding	normalisasi
tolong di fix ngedit video pake text di reels ...	tolong di fix ngedit video pake text di reels ...
facebook please kembalikan fitur edit untuk me...	facebook please kembalikan fitur edit untuk me...
aplikasi populer dan masih mampu bertahan masi...	aplikasi populer dan masih mampu bertahan masi...
Ini kenapa hobi banget error ya saya ga terkena...	ini kenapa hobi banget error ya saya tidak terk...
kenapa ini facebook setelah update malah makin...	kenapa ini facebook setelah update malah makin...
mohon kembalikan lagi durasi ketika mengunggah...	mohon kembalikan lagi durasi ketika mengunggah...
terlalu banyak pembatasan padahal gaada pelang...	terlalu banyak pembatasan padahal gaada pelang...
setelah update terbaru saat liat story video a...	setelah update terbaru saat liat story video a...
terbaik memang fitur facebook yang baru Teruta...	terbaik memang fitur facebook yang baru teruta...
upload video di reel di fb thumbnail sampul vi...	upload video di reel di facebook thumbnail sam...

Gambar 4. Hasil Normalisasi

4.5. Stopword Removal

Pada proses ini bertujuan untuk menghapus kata-kata umum seperti "dan", "yang", dan "dengan", tidak memiliki makna penting dalam analisis sentimen sehingga dapat dihapus untuk mengurangi kompleksitas data. Proses ini dilakukan menggunakan *library sastrawi*, yang dirancang khusus untuk Bahasa Indonesia. Pada Gambar 5 berikut.

normalisasi	stopword
tolong di fix ngedit video pake text di reels ...	fix ngedit video pake text reels setiap text t...
facebook please kembalikan fitur edit untuk me...	facebook please kembalikan fitur edit menambah...
aplikasi populer dan masih mampu bertahan masi...	aplikasi populer mampu bertahan masih beberapa...
Ini kenapa hobi banget error ya saya tidak terk...	kenapa hobi banget error terkena pl apapun tiba...
kenapa ini facebook setelah update malah makin...	ini facebook update malah makin jelek suka kel...
mohon kembalikan lagi durasi ketika mengunggah...	mohon kembalikan durasi mengunggah reels faceb...
terlalu banyak pembatasan padahal gaada pelang...	terlalu banyak pembatasan padahal gaada pelang...
setelah update terbaru saat liat story video a...	update terbaru liat story video ataupun postin...
terbaik memang fitur facebook yang baru Teruta...	terbaik memang fitur facebook baru terutama ka...
upload video di reel di facebook thumbnail sam...	upload video reel facebook thumbnail sampul vi...

Gambar 5. Hasil Stopword Removal

4.6. Tokenizing

Dalam proses ini bertujuan untuk memecah teks ulasan pengguna menjadi lebih kecil, seperti kata, yang disebut sebagai "*token*". Pada Gambar 6 berikut.

stopword	tokens
fix ngedit video pake text reels setiap text t...	[fix, ngedit, video, pake, text, reels, setiap...
facebook please kembalikan fitur edit menambah...	[facebook, please, kembalikan, fitur, edit, me...
aplikasi populer mampu bertahan masih beberapa...	[aplikasi, populer, mampu, bertahan, masih, be...
kenapa hobi banget error terkena pl apapun tiba...	[kenapa, hobi, banget, error, terkena, pl, apap...
ini facebook update malah makin jelek suka kel...	[ini, facebook, update, malah, makin, jelek, s...
mohon kembalikan durasi mengunggah reels faceb...	[mohon, kembalikan, durasi, mengunggah, reels,...
terlalu banyak pembatasan padahal gaada pelang...	[terlalu, banyak, pembatasan, padahal, gaada, ...
update terbaru liat story video ataupun postin...	[update, terbaru, liat, story, video, ataupun,...
terbaik memang fitur facebook baru terutama ka...	[terbaik, memang, fitur, facebook, baru, terut...
upload video reel facebook thumbnail sampul vi...	[upload, video, reel, facebook, thumbnail, sam...

Gambar 6. Hasil Tokenizing

4.7. Stemming

Pada proses ini mengubah kata-kata ke bentuk dasarnya maupun akarnya, bertujuan untuk mengurangi variasi kata yang tidak perlu, sehingga model *machine learning* dapat memfokuskan pada kata dasar dari kata tersebut, contoh nya pada kata "kesalahan" menjadi "salah". Pada Gambar 7 berikut.

tokens	stemmed
['fix', 'ngedit', 'video', 'pake', 'text', 're...	fix ngedit video pake text reels tiap text tex...
['facebook', 'please', 'kembalikan', 'fitur', ...	facebook please kembali fitur edit tambah foto...
['aplikasi', 'populer', 'mampu', 'bertahan', '...	aplikasi populer mampu tahan masih beberapa ke...
['kenapa', 'hobi', 'banget', 'error', 'terkena'...	kenapa hobi banget error kena pl apa tibatiba l...
['ini', 'facebook', 'update', 'malah', 'makin'...	ini facebook update malah makin jelek suka kel...
['mohon', 'kembalikan', 'durasi', 'mengunggah'...	mohon kembali durasi unggah reels facebook jad...
['terlalu', 'banyak', 'pembatasan', 'padahal', ...	terlalu banyak batas padahal gaada langgar buk...
['update', 'terbaru', 'liat', 'story', 'video'...	update baru liat story video atau postingan vi...
['terbaik', 'memang', 'fitur', 'facebook', 'ba...	baik memang fitur facebook baru utama kami cre...
['upload', 'video', 'reel', 'facebook', 'thumb...	upload video reel facebook thumbnail sampul vi...

Gambar 7. Hasil Stemming

Sebelum proses *labeling* dilakukan adalah memeriksa informasi dataset, beberapa kolom memiliki nilai yang hilang (*null*). Untuk memastikan dataset bersih dan dapat digunakan untuk analisis lebih lanjut, perlu menghapus baris yang memiliki nilai yang hilang dengan *data.dropna()*. Pada Gambar 8 berikut.

```
Data columns (total 8 columns):
```

#	Column	Non-Null Count	Dtype
0	content	4991 non-null	object
1	score	4991 non-null	int64
2	cleaning	4991 non-null	object
3	case_folding	4991 non-null	object
4	normalisasi	4991 non-null	object
5	stopword	4991 non-null	object
6	tokens	4991 non-null	object
7	stemmed	4991 non-null	object

```
dtypes: int64(1), object(7)
```

Gambar 8. Membersihkan nilai yang hilang

4.8. Labeling

Pada proses ini untuk mengkategorikan ulasan pengguna aplikasi Facebook berdasarkan skor *rating* yang telah diberikan di Google Play Store. Proses ini bertujuan untuk mengkategorikan teks ke dalam salah satu kategori sentimen positif, netral dan negatif. Jika *rating* adalah 5 dan 3 maka dikategorikan sebagai positif dengan jumlah 2067 ulasan, jika *rating* adalah 3 maka dikategorikan netral dengan jumlah 472 ulasan, dan jika *rating* adalah 2 dan 1 maka dikategorikan sebagai negatif dengan jumlah 2452 negatif. Pada Gambar 9 berikut.

	stemmed	label
0	fix ngedit video pake text reels tiap text tex...	netral
1	facebook please kembali fitur edit tambah foto...	positif
2	aplikasi populer mampu tahan masih beberapa ke...	positif
3	kenapa hobi banget error kena pl apa tibatiba l...	positif
4	ini facebook update malah makin jelek suka kel...	negatif
5	mohon kembali durasi unggah reels facebook jad...	netral
6	terlalu banyak batas padahal gaada langgar buk...	netral
7	update baru liat story video atau postingan vi...	positif
8	baik memang fitur facebook baru utama kami cre...	positif
9	upload video reel facebook thumbnail sampul vi...	negatif

Gambar 9. Hasil Labeling

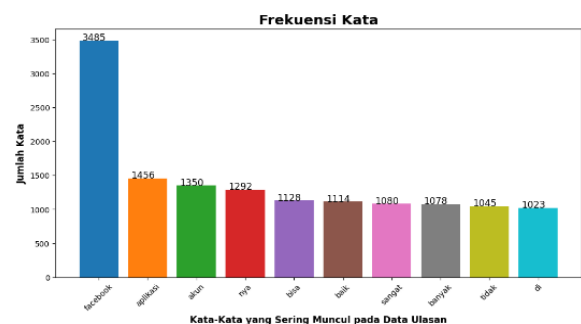
4.9. Exploratory Data Analysis (EDA)

Pada proses ini terdapat beberapa visualisasi yaitu terdapat awan kata dan frekuensi 10 kata yang sering muncul, visualisasi tersebut bertujuan untuk mengetahui kata-kata yang sering muncul pada keseluruhan data. Visualisasi tersebut berupa WordCloud dan Bar Chart.



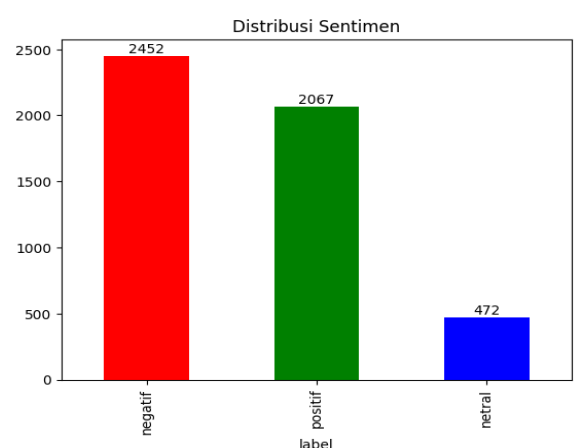
Gambar 10. WordCloud

Pada Gambar 10. *WordCloud* menunjukkan distribusi kata yang paling sering muncul dalam ulasan pengguna. Kata-kata yang tampil dengan ukuran lebih besar menandakan bahwa kata tersebut lebih sering digunakan dalam ulasan yang diberikan.



Gambar 11. Bar Chart 10 Kata teratas

Pada Gambar 11. *Bar Chart* tersebut memberikan informasi mengenai kata-kata yang paling dominan dalam ulasan pengguna. Dari grafik ini, dapat terlihat bahwa terdapat beberapa kata yang sering muncul di dalam data ulasan. Hal ini memberikan gambaran tentang kata-kata kunci yang sering digunakan pengguna dalam mengekspresikan pendapatnya.

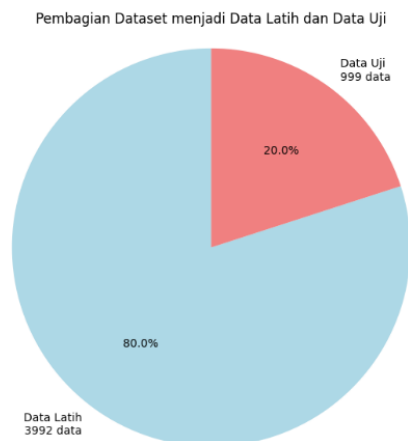


Gambar 12. Bar Chart Distribusi Sentimen

Pada Gambar 12. *Bar Chart* tersebut menampilkan jumlah ulasan berdasarkan kategori sentimen yang dihasilkan dari analisis, memberikan wawasan tentang proporsi sentimen positif, negatif, dan netral dalam dataset yang digunakan.

4.10. Splitting Data

Pada proses ini melakukan pembagian dataset ulasan menjadi dua bagian utama yaitu data latih (*training data*), dan data uji (*testing data*) dengan pembagian 80:20 berjumlah 3992 Data latih dan 999 Data Uji. Proses ini sangat penting dalam *machine learning* untuk melatih model pada sebagian data dan menguji kinerjanya pada data baru yang tidak terdapat pada datasetnya jumlah pembagian datanya divisualisasikan menggunakan *pie chart*. Pada Gambar 13 berikut.



Gambar 13. Pie Chart jumlah data latih dan data uji

4.11. Pemodelan Algoritma

Pada tahap pemodelan, algoritma *Naïve Bayes* (*MultinomialNB*) dan *K-Nearest Neighbor* (*KNeighborsClassifier*) digunakan untuk membangun model klasifikasi sentimen berdasarkan data ulasan pengguna. Proses dimulai dengan melakukan text vectorization menggunakan metode *TF-IDF* (*Term Frequency - Inverse Document Frequency*), yang mengubah teks menjadi representasi numerik berdasarkan frekuensi relatif kata dalam dokumen, sehingga dapat digunakan sebagai input untuk algoritma *machine learning*.

a. Klasifikasi *Naïve Bayes*

	precision	recall	f1-score	support
negatif	0.67	0.95	0.78	491
netral	0.00	0.00	0.00	94
positif	0.87	0.62	0.73	414
accuracy			0.73	999
macro avg	0.51	0.52	0.50	999
weighted avg	0.69	0.73	0.69	999

Gambar 14. Classification Report Naive Bayes

Pada Gambar 14. Hasil evaluasi menunjukkan bahwa model *Naive Bayes* memiliki akurasi sebesar 73%, yang menunjukkan tingkat keberhasilan dalam mengklasifikasikan sentimen ulasan. Pada kelas negatif, *recall* cukup tinggi (0.95), yang berarti model mampu mengidentifikasi hampir semua ulasan negatif, tetapi *precision*-nya hanya 0.67, menunjukkan adanya prediksi yang salah. Sebaliknya, pada kelas positif,

precision lebih tinggi (0.87), tetapi *recall* lebih rendah (0.62), yang mengindikasikan bahwa beberapa ulasan positif diklasifikasikan sebagai kategori lain. Sementara itu, kelas netral memiliki *precision*, *recall*, dan *F1-score* 0.00, yang menunjukkan bahwa model tidak mampu mengenali atau memprediksi ulasan dengan sentimen netral dengan baik, kemungkinan karena jumlah data yang lebih sedikit atau karena fitur yang kurang representatif. Secara keseluruhan, nilai *macro avg* (0.50 *F1-score*) yang lebih rendah dibandingkan *weighted avg* (0.69 *F1-score*) menunjukkan ketidakseimbangan performa antar kelas, dengan model lebih baik dalam mengenali ulasan positif dan negatif dibandingkan netral.

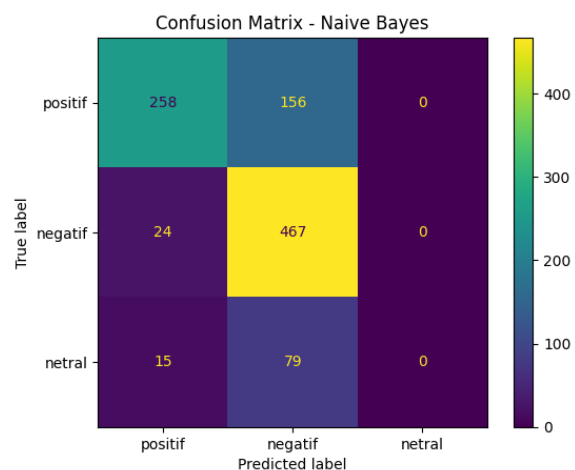
b. Klasifikasi *K-Nearest Neighbor*

	precision	recall	f1-score	support
negatif	0.71	0.86	0.78	491
netral	0.18	0.02	0.04	94
positif	0.76	0.72	0.74	414
accuracy			0.72	999
macro avg	0.55	0.53	0.52	999
weighted avg	0.68	0.72	0.69	999

Gambar 15. Classification Report KNN

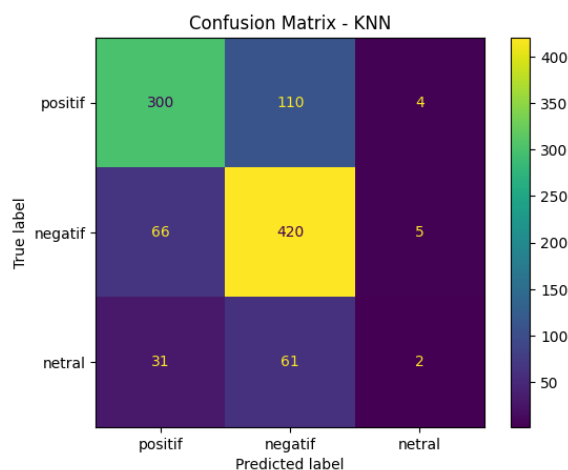
Pada Gambar 15. Hasil evaluasi model menggunakan algoritma *K-Nearest Neighbors* (KNN) menunjukkan bahwa akurasi model sebesar 72%. Pada kelas negatif, *recall* cukup tinggi (0.86), yang berarti model mampu menangkap sebagian besar ulasan negatif, dengan *precision* 0.71. Untuk kelas positif, *precision* dan *recall* masing-masing 0.76 dan 0.72, yang berarti model cukup baik dalam mengenali ulasan positif tetapi masih memiliki beberapa kesalahan klasifikasi. Namun, kelas netral memiliki performa sangat rendah, dengan *precision* 0.18 dan *recall* 0.02, yang menunjukkan model mengalami kesulitan dalam mengenali ulasan netral. Nilai *macro avg* (0.52 *F1-score*) lebih rendah dibandingkan *weighted avg* (0.69 *F1-score*), mengindikasikan ketidakseimbangan dalam performa antar kelas.

4.12. Confusion Matrix



Gambar 16. Confusion Matrix Naive Bayes

Pada Gambar 16, *Confusion matrix* untuk *Naïve Bayes* menunjukkan bahwa model berhasil mengklasifikasikan 258 ulasan positif dengan benar, namun 156 ulasan positif salah diklasifikasikan sebagai negatif. Untuk kelas negatif, model mampu mengenali 467 ulasan negatif dengan benar, tetapi masih terdapat 24 ulasan negatif yang salah dikategorikan sebagai positif. Sedangkan untuk kelas netral, model tidak dapat mengenali satupun ulasan dengan benar, dengan semua ulasan netral diklasifikasikan sebagai negatif (79 ulasan) atau positif (15 ulasan), yang menunjukkan kesulitan model dalam mendeteksi kategori netral.



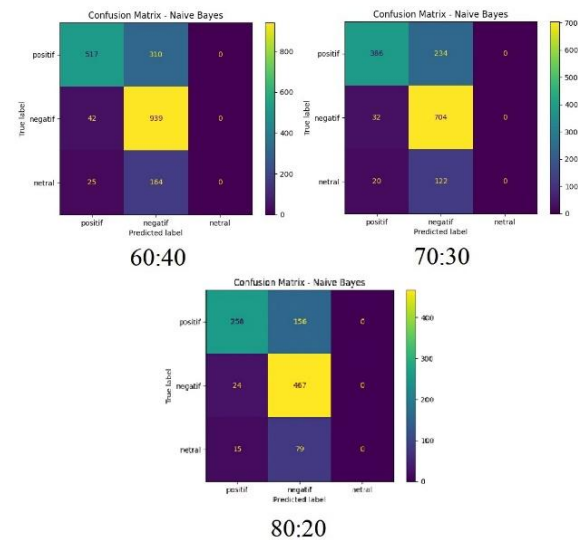
Gambar 17. Confusion Matrix KNN

Pada Gambar 17, *Confusion matrix* untuk *K-Nearest Neighbors (KNN)* menunjukkan bahwa model mengklasifikasikan 300 ulasan positif dengan benar, namun terdapat 110 ulasan positif yang salah dikategorikan sebagai negatif dan 4 ulasan positif yang diklasifikasikan sebagai netral. Untuk kelas negatif, model mampu mengenali 420 ulasan negatif dengan benar, tetapi masih terdapat 66 ulasan negatif yang dikategorikan sebagai positif dan 5 ulasan negatif yang salah diklasifikasikan sebagai netral. Pada kelas netral, model hanya mampu mengenali 2 ulasan netral dengan benar, sedangkan 31 ulasan netral dikategorikan sebagai positif dan 61 ulasan netral diklasifikasikan sebagai negatif, yang menunjukkan bahwa model juga mengalami kesulitan dalam mengenali kategori netral.

4.13. Perbandingan 3 Skenario Pembagian Data Naive Bayes dan KNN

Pada Gambar 18 dan Tabel 1. Menunjukkan jumlah prediksi yang benar dan salah untuk setiap kelas (positif, negatif, dan netral) berdasarkan tiga skenario pembagian data: 60:40, 70:30, dan 80:20. Dari hasil tersebut, model cenderung lebih akurat dalam mengklasifikasikan sentimen negatif dibandingkan sentimen positif atau netral. Hal ini terlihat dari nilai diagonal yang lebih tinggi pada kategori negatif. Dari *classification report*, model memiliki akurasi sebesar 72% untuk dataset 60:40, serta meningkat sedikit menjadi 73% pada dataset

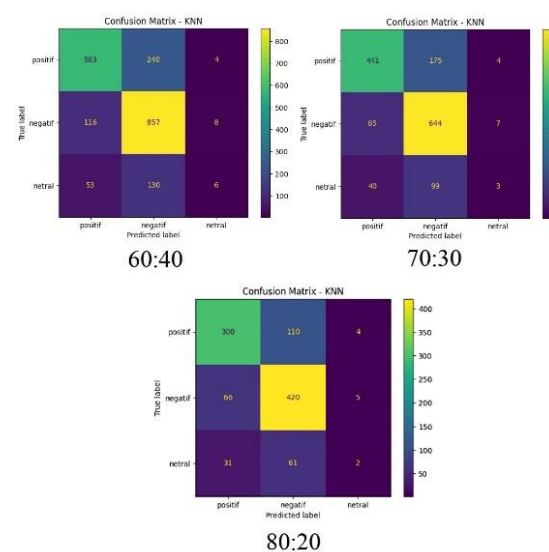
70:30 dan 80:20. *Precision* berkisar antara 50% hingga 51%, menunjukkan bahwa dari seluruh prediksi positif yang dibuat model, sekitar setengahnya benar. *Recall* juga berada di kisaran 51% hingga 52%, menandakan bahwa model berhasil mengidentifikasi lebih dari setengah data positif yang sebenarnya. *F1-Score*, sebagai gabungan *precision* dan *recall*, berada di kisaran 49% hingga 50%.



Gambar 18. 3 Confusion Matrix Naive Bayes

Tabel 1. Perbandingan Classification Matrix Naive Bayes

Matrix%	Dataset 60:40	Dataset 70:30	Dataset 80:20
Accuracy	72%	73%	73%
Precision	50%	51%	51%
Recall	51%	51%	52%
F1-Score	49%	50%	50%



Gambar 19. 3 Confusion Matrix KNN

Tabel 2. Perbandingan Classification Matrix KNN

Matrix%	Dataset 60:40	Dataset 70:30	Dataset 80:20
Accuracy	71%	71%	72%
Precision	53%	54%	55%
Recall	52%	53%	53%
F1-Score	51%	51%	52%

Pada Gambar 19 dan Tabel 2. Menunjukkan performa klasifikasi model berdasarkan tiga skenario pembagian dataset: 60:40, 70:30, dan 80:20. Dari hasil tersebut, model lebih banyak mengklasifikasikan sentimen negatif dengan benar dibandingkan sentimen positif atau netral. Hal ini terlihat dari jumlah prediksi yang benar pada kelas negatif yang lebih tinggi dibandingkan kelas lainnya. Dari *classification report*, model memiliki akurasi sebesar 71% pada dataset 60:40 dan 70:30, serta meningkat menjadi 72% pada dataset 80:20. *Precision* berkisar antara 53% hingga 55%, menunjukkan bahwa dari seluruh prediksi positif yang dibuat model, lebih dari setengahnya benar. *Recall* berada di kisaran 52% hingga 53%, menandakan bahwa model berhasil mengidentifikasi lebih dari setengah data positif yang sebenarnya. *F1-Score*, sebagai gabungan *precision* dan *recall*, berada di kisaran 51% hingga 52%.

4.14. Analisis Perbandingan Hasil Evaluasi

	precision	recall	f1-score	support
negatif	0.67	0.95	0.78	491
netral	0.00	0.00	0.00	94
positif	0.87	0.62	0.73	414
accuracy			0.73	999
macro avg	0.51	0.52	0.50	999
weighted avg	0.69	0.73	0.69	999

Gambar 20. Classification Report Naive Bayes

	precision	recall	f1-score	support
negatif	0.71	0.86	0.78	491
netral	0.18	0.02	0.04	94
positif	0.76	0.72	0.74	414
accuracy			0.72	999
macro avg	0.55	0.53	0.52	999
weighted avg	0.68	0.72	0.69	999

Gambar 21. Classification Report KNN

Berdasarkan pembagian data (*Splitting Data*) 60:40, 70:30, dan 80:20 sebelumnya, pembagian data 80:20 merupakan yang terbaik, oleh karena itu dilakukan analisis perbandingan pada pembagian data tersebut. Berdasarkan hasil *classification report* model *Naive Bayes* dan *K-Nearest Neighbor* Pada Gambar 20 dan Gambar 21. memiliki performa yang berbeda dalam beberapa aspek. Dari segi akurasi, *Naive Bayes* memiliki 73%, sedikit lebih tinggi dibandingkan *KNN* yang hanya 72%. Namun, jika dilihat dari *precision*, *recall*, dan *F1-score*, *KNN* menunjukkan nilai lebih seimbang. *Precision Naive Bayes* untuk kelas negatif (0.67) lebih rendah dibandingkan *KNN* (0.71), tetapi *precision* kelas positif *Naive Bayes* lebih tinggi (0.87)

dibandingkan *KNN* (0.76). *Recall Naive Bayes* untuk kelas negatif lebih tinggi (0.95) dibandingkan *KNN* (0.86), namun *recall* untuk kelas positif lebih rendah (0.62 *Naive Bayes*, 0.72 *KNN*). *F1-score* keseluruhan pada *Naive Bayes* adalah 0.73, sedangkan *KNN* sedikit lebih rendah di 0.72. Namun, jika melihat *macro average* dan *weighted average*, *KNN* lebih unggul dalam distribusi performa antar kelas, terutama karena netral memiliki nilai *recall* dan *precision* yang lebih baik dibandingkan *Naive Bayes*.

Berdasarkan analisis, *Naive Bayes* memiliki akurasi lebih tinggi dan *F1-score* keseluruhan yang lebih baik. Meskipun *KNN* memiliki *precision* yang lebih baik untuk beberapa kelas, *Naive Bayes* lebih unggul dalam menangkap kelas negatif dan positif dengan lebih seimbang, serta memiliki hasil lebih stabil dalam kondisi dataset yang lebih kecil. Selain itu, *Naive Bayes* lebih ringan dalam komputasi dibandingkan *KNN*, yang memerlukan lebih banyak sumber daya saat melakukan prediksi karena berbasis perhitungan jarak dengan data lain.

5. KESIMPULAN DAN SARAN

Hasil penelitian ini menunjukkan bahwa *Naive Bayes* memiliki performa lebih baik dibandingkan *KNN*. Model dievaluasi menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*, serta divisualisasikan melalui *confusion matrix*. *Naive Bayes* menghasilkan akurasi sebesar 73%, sedangkan *KNN* memiliki akurasi 72%. Selain itu, *Naive Bayes* lebih unggul dalam mengklasifikasikan sentimen negatif dengan *recall* 0.95, meskipun *precision*-nya lebih rendah dibandingkan *KNN*. Sementara itu, *KNN* memiliki distribusi performa yang lebih seimbang, terutama pada kelas netral, meskipun masih mengalami kesulitan dalam mendeteksi kategori tersebut. Adapun beberapa saran untuk penelitian selanjutnya antara lain yaitu: meningkatkan jumlah data yang digunakan, salah satu keterbatasan utama adalah rendahnya performa model dalam mengenali sentimen netral, penelitian selanjutnya dapat membandingkan algoritma lain, seperti *Support Vector Machine (SVM)* atau *Deep Learning*, guna memperoleh hasil yang lebih optimal. Untuk meningkatkan akurasi model, bisa dilakukan eksperimen dengan parameter tuning lebih lanjut menggunakan teknik seperti *Grid Search* atau *Random Search*.

DAFTAR PUSTAKA

- [1] I. Maili Ningrum and J. Ervina Sari, "Penelitian tentang Facebook," 2020.
- [2] I. Kurniasari and H. AlFatta, "ANALISIS SENTIMEN KOMENTAR FACEBOOK BERBASIS LEXICON DAN SUPPORT VECTOR MACHINE" 2020.
- [3] S. R. Salam and A. Nilogiri, "ANALISIS SENTIMEN PADA MEDIA SOSIAL FACEBOOK TERHADAP MARKETPLACE ONLINE DI INDONESIA MENGGUNAKAN

- METODE SUPPORT VECTOR MACHINE.” 2020.
- [4] H. Taufiqurrahman, F. Tri Anggraeny, and M. Muharrom Al Haromainy, “PERBANDINGAN ALGORITMA NAÏVE BAYES DAN K-NEAREST NEIGHBOR PADA ANALISIS SENTIMEN ULASAN APLIKASI MYPERTAMINA,” 2023.
- [5] V. Fitriyana *et al.*, “Analisis Sentimen Ulasan Aplikasi Jamsostek Mobile Menggunakan Metode Support Vector Machine,” 2023.
- [6] C. Zai and T. Komputer, “IMPLEMENTASI DATA MINING SEBAGAI PENGOLAHAN DATA.” 2022.
- [7] M. Febima, L. Magdalena, M. Asfi, M. Hatta, R. Fahrudin, and S. Informasi, “Implementasi Optimasi NLP dan KNN untuk User Review Aplikasi SAMPEAN Cirebon.” 2024.
- [8] H. Jurnal and D. Saputra, “JURNAL RISET SISTEM INFORMASI UPAYA PENDIDIKAN MENGGUNAKAN MACHINE LEARNING,” *Januari*, vol. 1, no. 1, pp. 57–62, 2024.
- [9] R. M. Kamal and E. Rainarli, “ANALISIS SENTIMEN CYBERBULLYING PADA KOMENTAR FACEBOOK DENGAN METODE KLASIFIKASI SUPPORT VECTOR MACHINE.” 2020.
- [10] Moh. A. Miftakurahmat, N. Safitri, P. A. Kusnadi, and C. Rozikin, “Klasifikasi Pengguna Hashtag pada Aplikasi Tiktok Menggunakan Perbandingan Metode K-Nearest Neighbors dan Naive Bayes Classifier,” *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 11, no. 3, Aug. 2023, doi: 10.23960/jitet.v11i3.3150.