

Analisis Sentimen Ulasan Aplikasi Maxim Untuk Peningkatan Layanan Menggunakan Algoritma *Naïve Bayes*

1st Dwivi Afdillah
Fakultas Rekayasa Industri
Universitas Telkom
Bandung, Indonesia

dwiviafdillah@student.telkomuniversit
y.ac.id

2nd Riska Yanu Fa'rifah
Fakultas Rekayasa Industri
Universitas Telkom
Bandung, Indonesia

riskayanu@telkomuniversity.ac.id

3rd Deden Witarsyah
Fakultas Rekayasa Industri
Universitas Telkom
Bandung, Indonesia

dedenw@telkomuniversity.ac.id

Abstrak — Transportasi *online* kian menarik perhatian dan minat masyarakat untuk dijadikan salah satu kebutuhan yang dapat memudahkan aktivitas sehari-hari. Maxim merupakan salah satu transportasi *online* yang menempati urutan ketiga jasa transportasi *online* yang paling sering digunakan. Dengan adanya hasil pemeringkatan jasa transportasi *online* tersebut maka pihak Maxim dapat meningkatkan pelayanan berdasarkan dari ulasan pengguna aplikasi Maxim melalui metode analisis sentimen yang bertujuan untuk mengolah sejumlah besar data secara selektif dan efisien dengan mengelompokkan ke dalam dua ulasan yaitu positif dan negatif dengan menggunakan algoritma *Naïve Bayes*. Tahap untuk melakukan analisis sentimen dibagi menjadi beberapa bagian diantaranya yaitu pengumpulan data, *preprocessing*, *modelling*, dan evaluasi. Pada penelitian ini, menggunakan tiga skenario rasio pembandingan *training testing* yaitu 60:40, 70:30, dan 80:20 dengan menggunakan model *Multinomial Naïve Bayes* menunjukkan bahwa rasio 70:30 menghasilkan model terbaik dengan nilai akurasi sebesar 87,22% yang dilakukan melalui *confusion matrix* dengan nilai *recall* sebesar 98,49%, nilai *precision* sebesar 86,62%, dan *f1 score* menghasilkan nilai sebesar 92,17%. Berdasarkan dari hasil visualisasi sentimen didapatkan hasil ulasan pengguna aplikasi Maxim cenderung mengarah ke sentimen positif dengan jumlah ulasan positif sebanyak 76% dan sentimen negatif sebanyak 24%. Sehingga, hasil kategori sentimen tersebut dapat dijadikan bahan evaluasi kepada pihak *developer* sebagai bentuk peningkatan layanan aplikasi Maxim.

Kata kunci— Analisis Sentimen, Transportasi *Online*, *Naïve Bayes*

I. PENDAHULUAN

Transportasi *online* di Indonesia kian menarik perhatian dan minat masyarakat untuk dijadikan salah satu kebutuhan yang dapat memudahkan aktivitas sehari-hari. Dengan perkembangan teknologi di era digitalisasi saat ini yang berjalan sangat pesat, sehingga berpengaruh pada aspek ekonomi dikarenakan terciptanya perusahaan yang menyediakan layanan transportasi *online*. Berdasarkan dari data Asosiasi Penyelenggara Jasa Internet Indonesia (APJII), tercatat sebanyak 210.026.769 jiwa yang terkoneksi dengan internet di Indonesia tahun 2021 - 2022 dari total populasi 272.682.600 jiwa penduduk Indonesia tahun 2021 [1].

Dengan berbagai macam aplikasi yang ada, alasan masyarakat menggunakan internet salah satunya karena adanya jasa transportasi *online* yang dapat memudahkan masyarakat untuk melakukan perjalanan tanpa harus ke pangkalan ojek ataupun menunggu dipinggir jalan untuk mendapatkan transportasi baik itu motor maupun mobil. Sehingga berdasarkan dari data Asosiasi Penyelenggara Jasa Internet Indonesia (APJII), transportasi *online* berada pada peringkat ke 8 dari 9 alasan masyarakat menggunakan internet [1]. Berdasarkan dari riset yang telah dilakukan oleh ShopBack pada 1.000 pengguna transportasi *online* di lima kota besar yang ada di Indonesia yaitu Jabodetabek, Makassar, Medan, Bandung, dan Surabaya, diperoleh hasil bahwa alasan penggunaan transportasi *online* yang utama karena murah dengan persentase 55.4 %, nyaman 19.6%, 12.5% karena respon yang cepat, aman sebesar 8.2%, dan 4.3% alasan lainnya. Tidak hanya itu, dapat dilihat juga bahwa mayoritas responden telah menggunakan jasa transportasi *online* sebesar 91.4% dibandingkan dengan yang belum menggunakan jasa transportasi *online* sebesar 8.6% [2].

Maxim merupakan salah satu jasa transportasi *online* yang cukup populer di Indonesia dengan misi meningkatkan interaksi secara terus menerus antara para pengguna dan dapat membantu banyak orang dalam melakukan perjalanan yang ingin dituju. Maxim juga telah merambah ke beberapa kota besar di Indonesia dan bersaing dengan beberapa jasa transportasi *online* yang ada. Berdasarkan dari sumber yang ada Maxim menempati urutan ketiga sebagai aplikasi ojek *online* terbaik dan terpopuler di Indonesia tahun 2022 dan telah bersaing dengan jasa transportasi *online* seperti Grab dan Gojek [3]. Sebagai peningkatan kualitas layanan, Maxim dapat memanfaatkan ulasan pengguna terhadap aplikasi Maxim yang terdapat pada Google Play Store yang bertujuan untuk mengetahui ulasan yang diberikan pengguna baik itu bersifat saran, pujian atau bahkan keluhan dari pengguna. Sehingga dari ulasan tersebut dapat dijadikan acuan untuk memajukan jasa layanan transportasi *online* Maxim. Tetapi akan sulit untuk mengetahui semua informasi ulasan yang ada pada jasa layanan transportasi *online* Maxim jika sekadar membaca satu persatu. Maka dari itu diperlukanlah sebuah metode analisis sentimen yang dapat membantu dalam

mengolah data agar terbilang hemat dari segi waktu serta langsung mengelompokkan ulasan-ulasan tersebut ke dalam dua ulasan yaitu positif dan negatif.

Dalam melakukan analisis sentimen kerap kali menggunakan beberapa jenis algoritma diantaranya yaitu *Naïve Bayes*, *Decision Tree*, *Support Vector Machine*, *C4.5*, dan *K-Nearest Neighbor*. Sehingga untuk menghasilkan analisis sentimen berdasarkan ulasan pengguna aplikasi Maxim, maka dibutuhkan metode dan algoritma yang baik dan akurat. Bahkan masing-masing jenis algoritma memiliki tingkat akurasi yang berbeda pada setiap kasus yang ada. Seperti, pada penelitian sebelumnya dengan melakukan perbandingan algoritma *Naïve Bayes* dan *C4.5* pada analisis sentimen presiden 3 periode di Twitter dan didapatkan hasil akurasi pada *C4.5* memiliki akurasi sebesar 78% dibandingkan dengan akurasi *Naïve Bayes* sebesar 85% lebih tinggi dari akurasi algoritma *C4.5* [4]. Penelitian sebelumnya juga pernah melakukan perbandingan algoritma *Naïve Bayes*, *Support Vector Machine*, *C4.5*, dan *K-Nearest Neighbor* pada analisis sentimen ulasan aplikasi Bibit di Play Store dan didapatkan hasil akurasi yaitu pada algoritma *Naïve Bayes* sebesar 84,91%, 71,77% merupakan akurasi pada algoritma *Support Vector Machine*, akurasi pada algoritma *C4.5* sebesar 76,94%, dan 66,81% merupakan akurasi pada algoritma *K-Nearest Neighbor* [5]. Dari hasil penelitian terdahulu, dapat dilihat bahwa algoritma *Naïve Bayes* memiliki kelebihan dibandingkan dengan algoritma lainnya karena algoritma *Naïve Bayes* memiliki tingkat akurasi dan kecepatan yang tinggi dalam melakukan proses *sentiment analysis* serta merupakan metode yang baik dan akurat untuk digunakan. Sehingga pada penelitian ini akan melakukan Analisis Sentimen ulasan pengguna aplikasi Maxim untuk peningkatan layanan dengan menggunakan Algoritma *Naïve Bayes* dengan tujuan untuk melakukan analisis berdasarkan ulasan pengguna aplikasi Maxim di Google Play Store dan untuk menilai tingkat kepuasan pengguna yang didapatkan melalui ulasan aplikasi Maxim yang bersifat teks serta dapat dikelompokkan menjadi asumsi positif dan negatif berdasarkan dari *user experience*.

II. KAJIAN TEORI

A. Google Play Store

Google Play Store merupakan layanan dari google yang biasa disebut dengan istilah Google Play. Google Play Store juga merupakan kumpulan aplikasi yang berbasis digital dan dapat diakses secara *online*. Awalnya Google Play dikenal dengan sebutan Android Market yang berfungsi untuk mendownload aplikasi khusus bagi pengguna android. Namun setelah namanya diganti, semua layanan aplikasi dapat dilihat melalui Google Play. Tidak hanya itu Google Play juga memiliki fitur ulasan pada tiap aplikasi yang ada [6].

B. Maxim

Maxim merupakan layanan transportasi *online* yang mulai beroperasi di Indonesia pada tahun 2018. Dengan misi ingin meningkatkan hubungan interaksi yang baik kepada para pengguna agar menciptakan kenyamanan. Sehingga para pengguna dapat mempercayakan perjalanan dengan menggunakan Maxim. Tidak hanya itu, Maxim juga telah melakukan perluasan ke beberapa kota besar di Indonesia

dengan menyediakan berbagai macam layanan mulai dari layanan transportasi, layanan *delivery* makanan, dan layanan pengiriman barang [7]

C. Text Mining

Text mining merupakan metode yang dimanfaatkan untuk mengetahui jenis *text* berdasarkan tipe atau kategori. Tujuan diciptakannya *text mining* yaitu untuk meningkatkan pencarian *database* yang sifatnya sekuens dan bertujuan untuk mengetahui ulasan atau *review* pengguna pada suatu *system* atau aplikasi. [8].

Text mining adalah sumber data yang digunakan untuk memastikan informasi yang belum tahu akan kepastiannya. *Text mining* dilakukan dengan cara ekstraksi informasi, mengelompokkan teks, dan melakukan ekstraksi data yang bertujuan untuk mengetahui informasi secara pasti atau nyata sehingga dapat dituliskan ke dalam *text* [9].

D. Sentiment Analysis

Sentiment Analysis atau biasa disebut dengan *opinion mining* merupakan bagi dari penelitian *text mining*. *Sentiment analysis* digunakan untuk mengetahui ulasan pengguna yang berkaitan dengan nilai emosional dan perasaan yang dituangkan dalam kata atau teks yang sifatnya saran ataupun keluhan. Berdasarkan dari ulasan yang ada maka dapat dikategorikan ulasan sentimennya yaitu dari segi suka, tidak suka, dan netral [10].

Klasifikasi juga merupakan bagian dari *sentiment analysis* karena pada klasifikasi dilakukan pengelompokan berdasarkan dari sikap atau pendapat pengguna yang sifatnya terbilang problematis. Tidak hanya itu, *sentiment analysis* juga memiliki manfaat yaitu untuk mengetahui tingkat kepuasan masyarakat terkait aplikasi yang digunakan apakah bersifat positif atau negatif. *Sentiment analysis* juga banyak digunakan dalam bidang usaha atau layanan karena berfungsi untuk mengetahui ulasan pengguna terkait sebuah produk. Sehingga dengan adanya *sentiment analysis* maka dapat dilakukan peningkatan pada rencana pemasaran dengan mengetahui kemauan pengguna berdasarkan dari data ulasan yang telah ada sebelumnya. Selain itu, *sentiment analysis* juga memiliki keunggulan yaitu lebih efisien dari segi waktu karena dapat melakukan analisis dengan menyeleksi banyak data [11].

E. Naïve Bayes

Naïve Bayes merupakan algoritma yang digunakan pada pengklasifikasian data dan termasuk algoritma dengan perhitungan yang cepat, algoritma yang sederhana, dan termasuk algoritma yang mempunyai tingkat ketelitian yang tinggi. *Naïve Bayes* sering digunakan pada proses *sentiment analysis*. Sifat dari algoritma *Naïve Bayes* lebih cocok digunakan pada data yang berukuran besar karena dapat mengatasi *missing value* [12].

$$P(A|B) = \frac{P(B|A) \times P(A)}{P(B)} \quad (1)$$

Berdasarkan dari rumus (1), adapun penjelasan parameternya sebagai berikut:

- a) *B* : Data training
- b) *A* : Kelas klasifikasi berupa positif dan negatif
- c) *P(A|B)* : Peluang hipotesis *A* berdasarkan kondisi *B*

- d) $P(B|A)$: Peluang kejadian data B , berdasarkan kondisi dari hipotesis A
- e) $P(A)$: Peluang dari hipotesis A
- f) $P(B)$: Peluang kemunculan data B yang diamati

Terdapat tiga jenis tipe atau model *Naïve Bayes* berdasarkan dari penelitian [13] di antaranya yaitu :

1. Multinomial Naïve Bayes

Multinomial Naïve Bayes merupakan salah satu tipe metode klasifikasi yang berdistribusi *multinomial* dan digunakan dalam klasifikasi teks dari suatu dokumen dengan menghitung frekuensi kemunculan setiap kata yang terdapat dalam dokumen

2. Bernoulli Naïve Bayes

Bernoulli Naïve Bayes merupakan tipe metode klasifikasi yang cocok dengan fitur biner dengan menggunakan simbol angka 0 dan 1. Dimana 0 menunjukkan kata yang tidak muncul dalam dokumen, sedangkan 1 menunjukkan angka yang muncul dalam dokumen

3. Gaussian Naïve Bayes

Gaussian Naïve Bayes merupakan tipe klasifikasi yang menerapkan distribusi Gaussian pada nilai yang sifatnya kontinu atau tidak terbatas.

F. TF-IDF

TF-IDF (*Term Frequency – Inverse Document Frequency*) merupakan sebuah metode pengolahan teks yang bertujuan untuk menghitung nilai bobot setiap kata yang ada dalam dokumen [14].

G. Matrix Evaluasi

Confusion Matrix merupakan metode yang digunakan dalam melakukan evaluasi algoritma dan bertujuan untuk mengetahui jumlah prediksi benar dan salah dari suatu model klasifikasi data *testing* [15]. *Confusion Matrix* juga bertujuan untuk menghitung nilai akurasi, *precision*, *recall*, dan *f1 score* [16]. Berikut ini merupakan tabel confusion matrix.

TABEL 1
CONFUSION MATRIX

Confusion Matrix		Predicted	
		Negative	Positive
Actual	Negative	TN	FP
	Positive	FN	TP

Adapun penjelasan parameternya sebagai berikut:

- TP (*True Positive*) : Jumlah kelas positif yang benar-benar diprediksi benar oleh model
- FP (*False Positive*) : Jumlah kelas negatif yang salah diprediksi sebagai kelas positif oleh model
- TN (*True Negative*) : Jumlah kelas negatif yang benar-benar diprediksi benar oleh model
- FN (*False Negative*) : Jumlah kelas positif yang salah diprediksi sebagai kelas negatif oleh model

Berdasarkan keempat parameter *confusion matrix* di atas dapat digunakan untuk menghitung nilai *accuracy*, *precision*,

recall, dan *f1score*. Berikut ini merupakan rumus *confusion matrix*.

$$Accuracy = \frac{TP + TN}{(TP + TN + FP + FN)} \quad (2)$$

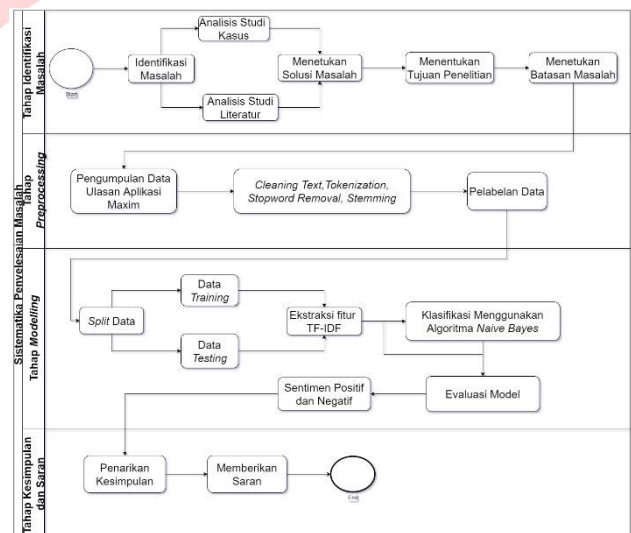
$$Precision = \frac{TP}{(TP + FP)} \quad (3)$$

$$Recall = \frac{TP}{(TP + FN)} \quad (4)$$

$$F1 \text{ Score} = \frac{2 \times Recall \times Precision}{Recall + Precision} \quad (5)$$

III. METODE

Sistematika penyelesaian masalah pada penelitian ini dapat dikelompokkan menjadi 4 tahap yaitu tahap identifikasi masalah, tahap *preprocessing*, tahap *modelling*, dan yang terakhir yaitu tahap kesimpulan dan saran. Di bawah ini merupakan gambar sistematika penyelesaian masalah yang digunakan pada penelitian.



GAMBAR 1
SISTEMATIKA PENYELESAIAN MASALAH

A. Tahap Identifikasi Masalah

Identifikasi masalah merupakan tahap pertama yang dilakukan dalam proses penyelesaian masalah. Pada tahap ini, peneliti mencari tahu dengan melakukan identifikasi terkait ulasan pengguna berdasarkan layanan yang diberikan oleh aplikasi Maxim. Setelah melakukan identifikasi masalah, selanjutnya peneliti melakukan analisis studi kasus dan analisis studi literatur. Kemudian peneliti menentukan solusi dengan menggunakan algoritma yang tepat berdasarkan referensi yang ada. Lalu, dilanjutkan dengan menentukan tujuan penelitian. Sehingga dari tujuan penelitian tersebut dilanjutkan dengan menentukan batasan masalah yang cocok agar terbilang lebih spesifik.

B. Tahap Preprocessing

Preprocessing merupakan tahap kedua yang dilakukan dalam proses penyelesaian masalah. Data *preprocessing* merupakan bagian dari data *mining* yang bertujuan untuk mengubah data mentah menjadi data yang mudah dipahami.

Pada tahap *preprocessing* terdapat beberapa proses yang dilakukan diantaranya yaitu melakukan pengumpulan data ulasan pengguna aplikasi Maxim pada Google Play Store dengan menggunakan metode *web scraping* yang bertujuan untuk mengekstrak data secara otomatis. Tahap selanjutnya yaitu melakukan *cleaning text* yang bertujuan untuk membersihkan dan menghilangkan elemen yang tidak sesuai dari data teks, terdapat beberapa langkah umum dalam melakukan *cleaning text* yaitu diantaranya *case folding* yang bertujuan untuk mengubah semua huruf yang ada dalam dokumen menjadi huruf kecil dan bertujuan untuk menghapus simbol yang ada pada data ulasan. Kemudian, dilanjutkan dengan melakukan proses *punctuation removal* yang bertujuan untuk menghapus semua tanda baca ataupun karakter yang ada dalam teks. *Tokenization* merupakan proses yang bertujuan untuk memisahkan karakter teks menjadi bentuk token. Lalu, dilanjutkan dengan proses *stopword removal* yang bertujuan untuk menghapus kata yang tidak cocok atau aneh, kemudian dilanjutkan dengan melakukan proses *stemming* yang bertujuan untuk mengubah kata menjadi bentuk kata dasar yang sifatnya baku dengan menghilangkan imbuhan dari suatu kata. Tahap terakhir yang dilakukan pada proses *preprocessing* yaitu tahap pelabelan data dengan membagi data menjadi dua label yaitu label positif dan label negatif.

C. Tahap Modelling

Modelling merupakan tahap ketiga yang dilakukan dalam proses penyelesaian masalah yang bertujuan untuk mengolah data yang telah diproses sebelumnya pada tahap *preprocessing*. Tahap *modelling* terdiri dari beberapa bagian diantaranya melakukan *split* data yang bertujuan membagi data menjadi dua bagian yaitu data *training* dan data *testing*. Kemudian dilakukan ekstraksi fitur TF-IDF yang bertujuan untuk mengetahui nilai bobot kata dari data *training* dan data *testing*. Setelah itu, nilai bobot kata pada data *training* digunakan untuk melatih model klasifikasi menggunakan algoritma *Naïve Bayes*. Sedangkan, untuk nilai bobot kata pada data *testing* digunakan untuk melakukan evaluasi model yang telah dilatih dengan menggunakan tiga rasio pembandingan untuk data *training* dan data *testing* yaitu rasio 60:40, 70:30, dan 80:20.

D. Tahap Kesimpulan dan Saran

Kesimpulan dan saran merupakan tahap akhir yang dilakukan pada proses penyelesaian masalah yang terdiri dari hasil analisis yang diperoleh dari tahap sebelumnya. Kemudian, dilanjutkan pada kesimpulan dan saran.

IV. HASIL DAN PEMBAHASAN

A. Pengumpulan Data

Data yang digunakan merupakan data ulasan pengguna aplikasi Maxim pada Google Play Store dengan melakukan teknik *web scraping* data menggunakan bahasa pemrograman *python*. Jumlah ulasan pengguna yang digunakan untuk analisis sentimen pada penelitian ini sebanyak 20000 ulasan yang sifatnya terbaru dan data yang digunakan telah disortir terlebih dahulu berdasarkan dari waktu dan tanggal terbaru sampai dengan tanggal pengambilan data yaitu 11 Maret 2023. Data ulasan yang telah disortir kemudian disimpan ke dalam format file CSV

yang bertujuan untuk mempermudah dalam mengolah data pada tahap *preprocessing*. Berikut merupakan contoh data ulasan pengguna aplikasi Maxim yang belum diolah:

TABEL 2
CONTOH DATA ULASAN

No	Data Ulasan
1	terimakasih sudah mengantarkan kami dengan selamat sampai tujuan 🙏🙏
2	Maxim memang murah cuman terlalu lambat + map nya kurang
3	Mohon diperbaiki lagi yaa, tiap kali saya mau order kok dimaps app maxim gaada yaa. Sampe pusing bgt nyari alamat ga terdaftar di maps

B. Preprocessing Data

Preprocessing data dilakukan untuk mengubah data mentah menjadi data yang mudah dipahami agar proses klasifikasi sentimen dapat berjalan dengan optimal. Adapun tahap *preprocessing* terdiri dari pengumpulan data, *cleaning text*, *tokenization*, *stopword removal*, *stemming*, dan pelabelan data. Berikut ini merupakan contoh *preprocessing* data:

TABEL 3
HASIL PREPROCESSING

Data Ulasan	mantap dan yang terbagus di maxim, sangat cepat dan aman
Cleaning Text	mantap dan yang terbagus di maxim sangat cepat dan aman
Tokenization	['mantap', 'dan', 'yang', 'terbagus', 'di', 'maxim', 'sangat', 'cepat', 'dan', 'aman']
Stopword Removal	['mantap', 'terbagus', 'maxim', 'cepat', 'aman']
Stemming	['mantap', 'bagus', 'maxim', 'cepat', 'aman']

C. Modelling

Pada tahap ini, dilakukan pemodelan pada data ulasan yang telah melalui proses *preprocessing*. Data ulasan tersebut akan dibagi menjadi dua bagian yaitu data *training* dan data *testing*. Sehingga pada tahap pemodelan terdiri dari tiga sub-tahapan diantaranya yaitu *split* data, TF-IDF, dan pembentukan model.

1. Split Data Training dan Data Testing

Pada tahap ini akan dilakukan *split* data dengan membagi data menjadi dua bagian yaitu data *training* dan data *testing*. Tujuan dilakukan *split* data yaitu untuk mengetahui jumlah data *training* dan data *testing* yang optimal untuk algoritma machine learning[17]. Pada tahap ini data akan dibagi menjadi tiga rasio pembandingan yaitu rasio 60:40, rasio 70:30, dan rasio 80:20. Berikut ini merupakan tabel *split* data.

TABEL 4
SPLIT DATA

Rasio	Pembagian Data		Total Data
	Training	Testing	
60:40	10.945	7.297	18.242
70:30	12.769	5.473	
80:20	14.593	3.649	

2. Ekstraksi Fitur TF-IDF

TF-IDF merupakan gabungan dari TF (*Term Frequency*) yang dan IDF (*Inverse Document Frequency*) yang berfungsi untuk menghitung bobot pada setiap kata (*term*) dalam dokumen ulasan pengguna aplikasi Maxim. Berikut ini merupakan contoh hasil TF-IDF.

TABEL 5
HASIL TF-IDF

Term	TF-IDF		
	d1	d2	d3
bagus	0,4472136	0	0
cepat	0	0	0,5
driver	0	0,5	0
fast	0	0,5	0
jalan	0,4472136	0	0
jemput	0	0	0,5
mantap	0	0,5	0
motor	0,4472136	0	0
nyaman	0,4472136	0	0
ramah	0,4472136	0	0
respon	0	0,5	0
sesuai	0	0	0,5
titik	0	0	0,5
Total bobot tiap dokumen	2,236068	2	2

3. Pembentukan Model

Pada penelitian ini, dilakukan klasifikasi pembagian data dengan masing-masing rasio berbeda yang bertujuan untuk mengetahui rasio yang menghasilkan nilai akurasi model terbaik. Berikut ini merupakan klasifikasi menggunakan *Multinomial Naïve Bayes*.

TABEL 6
PSEUDOCODE NAÏVE BAYES

Input:	- Dataset Pelatihan (tf_train, y_train), dataset pengujian (tf_test)
Output:	- Array hasil prediksi model <i>Multinomial Naïve Bayes</i>
Step:	- Melakukan inisialisasi variabel untuk menyimpan probabilitas kelas ($P(B)$) - Melakukan inisialisasi matriks untuk menyimpan probabilitas fitur untuk setiap kelas ($P(F_i B)$) - Setiap sampel (x) dalam dataset pelatihan - Mengambil label kelas (y) dari sampel - Pada jumlah sampel dalam kelas y ($N(B)$) ditambahkan angka 1 - Untuk setiap fitur F_i dalam x ditambahkan 1 pada jumlah sampel dalam kelas y melalui fitur F_i ($N(B, F_i)$) - Untuk setiap kelas B dilakukan perhitungan probabilitas kelas B sebagai $P(B) = N(B) / \text{total sampel}$ - Untuk setiap fitur F_i dilakukan perhitungan probabilitas fitur F_i untuk kelas B sebagai $P(F_i B) = (N(C, F_i) + 1) / (N(C) + \text{total_features})$
Melakukan inisialisasi array untuk menyimpan prediksi kelas setiap sampel pengujian menggunakan <i>Multinomial Naïve Bayes</i>	

D. Evaluasi Performansi Model Klasifikasi

Pada tahap ini, dilakukan proses pengujian dan penilaian untuk mengetahui seberapa baik model klasifikasi yang digunakan dengan masing-masing skenario rasio data yang berbeda. Berikut ini merupakan tabel hasil *accuracy*, *precision*, *recall*, dan *f1 score* yang didapatkan dari proses klasifikasi menggunakan *confusion matrix* dengan model *Multinomial Naïve Bayes*.

TABEL 7
EVALUASI PERFORMANSI MODEL

Rasio	Accuracy	Precision	Recall	F1 Score
-------	----------	-----------	--------	----------

60:40	86,95%	86,25%	98,53%	91,98%
70:30	87,22%	86,62%	98,49%	92,17%
80:20	87,14%	86,59%	98,38%	92,11%

Berdasarkan tabel di atas, hasil evaluasi menunjukkan bahwa model terbaik yaitu pada rasio 70:30 dengan nilai akurasi sebesar 87,22%, nilai *recall* sebesar 98,49%, nilai *precision* sebesar 86,62%, dan *f1 score* menghasilkan nilai sebesar 92,17%.

Berikut ini tabel *confusion matrix* untuk model klasifikasi yang memiliki nilai akurasi terbaik yaitu pada rasio 70:30.

TABEL 8
CONFUSION MATRIX RASIO 70:30

Confusion Matrix		Predicted	
		Negatif	Positif
Actual	Negatif	656	636
	Positif	63	4118

V. KESIMPULAN

Berdasarkan implementasi *sentiment analysis* yang dilakukan pada data ulasan aplikasi Maxim di Google Play Store, model yang optimal didapatkan dari implementasi pada rasio 70:30 menggunakan algoritma *Naïve Bayes* dengan distribusi *Multinomial Naïve Bayes* didapatkan hasil akurasi sebesar 87,22% yang menunjukkan bahwa prediksi yang dilakukan oleh model adalah benar dan terbilang bagus dalam klasifikasi data. Sedangkan hasil analisis sentimen, aplikasi Maxim cenderung mengarah ke sentimen positif dengan jumlah ulasan positif sebesar 76% dan didominasi dengan kata bagus, mantap, cepat, ramah, puas, murah, terimakasih, dan beberapa kata lainnya yang menunjukkan kepuasan pengguna pada saat menggunakan layanan transportasi Maxim. Dalam hal ini, pengguna menyampaikan kepuasan terhadap layanan aplikasi Maxim, seperti adanya *driver* yang bersikap ramah kepada pengguna dengan memberikan pelayanan yang bagus dan cepat. Tetapi, terdapat juga sentimen negatif sebesar 24%, sehingga perlu dijadikan sebagai bahan evaluasi bagi pengembangan untuk meningkatkan kinerja dan fungsionalitas aplikasi.

REFERENSI

- [1] APJII, "Profil Internet Indonesia 2022," *Apji.or.Od*, no. June, p. 10, 2022, [Online]. Available: apji.or.id
- [2] ShopBack Indonesia, "Sering Membandingkan Harga Transportasi Online? Aplikasi Ini Akan Memudahkan Penggunanya," 2018. <https://www.shopback.co.id/katashopback/transportasi-online-makin-digemari>
- [3] Listiorini, "15 Aplikasi Ojek Online Terbaik dan Terpopuler di Indonesia," *Carisinyal*, 2022. <https://carisinyal.com/aplikasi-ojek-online/>
- [4] F. Albasithu and A. Wibowo, "Perbandingan Algoritma Naïve Bayes Dan C4.5 pada Analisis Sentimen Presiden 3 Periode di Twitter," *Semin. Nas. Mhs. Fak. Teknol. Inf. Jakarta-Indonesia*, no. September, pp. 510–516, 2022, [Online]. Available: <https://senafiti.budiluhur.ac.id/index.php/>
- [5] A. Z. Kamalia, A. Al Zaroni, and M.

- Wangsadanureja, "Analisis Sentimen Pada Ulasan Aplikasi Bibit Di Play Store Dengan Metode Naive Bayes, Support Vector Machine, C4.5 Dan K-Nearest NEIGHBOR," vol. 13, no. 1, 2022.
- [6] Kompasiana, "Apa Bedanya Google Play Store dengan Google Store?," *Kompasiana*, 2021. https://www.kompasiana.com/ruangmuda1780/6025ff5ed541df70da59f302/apa-bedanya-google-play-store-dengan-google-store?page=2&page_images=1
- [7] Layanan Maxim, "Tentang Perusahaan Maxim," *Maxim*, 2022. <https://id.taximaxim.com/id/2093-jakarta/about/>
- [8] K. B. Cohen and L. Hunter, "Getting started in text mining," *PLoS Comput. Biol.*, vol. 4, no. 1, pp. 0001–0003, 2008, doi: 10.1371/journal.pcbi.0040020.
- [9] T. Kwartler, "What is Text Mining?," *Text Min. Pract. with R*, pp. 1–15, 2017, doi: 10.1002/9781119282105.ch1.
- [10] S. Redhu, "Sentiment Analysis Using Text Mining: A Review," *Int. J. Data Sci. Technol.*, vol. 4, no. 2, p. 49, 2018, doi: 10.11648/j.ijdst.20180402.12.
- [11] P. Nomleni, M. Hariadi, and I. K. E. Purnama, "Sentiment Analysis Berbasis Big Data," *Semin. Nas. Rekayasa Teknol. Ind. dan Inf.*, vol. 9, pp. 142–149, 2014.
- [12] I. Romli, T. Pardamean, S. Butsianto, T. N. Wiyatno, and E. Bin Mohamad, "Naive Bayes Algorithm Implementation Based on Particle Swarm Optimization in Analyzing the Defect Product," *J. Phys. Conf. Ser.*, vol. 1845, no. 1, 2021, doi: 10.1088/1742-6596/1845/1/012020.
- [13] S. Xu, "Bayesian Naïve Bayes classifiers to text classification," *J. Inf. Sci.*, vol. 44, no. 1, pp. 48–59, 2018, doi: 10.1177/0165551516677946.
- [14] Scikit-learn:Machine Learning in python, "Feature extraction," 2011. https://scikit-learn.org/stable/modules/feature_extraction.html#text-feature-extraction
- [15] P. S. Saragih, D. Witarsyah, F. Hamami, and J. M. MacHado, "Sentiment Analysis of Social Media Twitter with Case of Large Scale Social Restriction in Jakarta using Support Vector Machine Algorithm," *2021 Int. Conf. Adv. Data Sci. E-Learning Inf. Syst. ICADEIS 2021*, vol. 19, no. January 2020, pp. 1–6, 2021, doi: 10.1109/ICADEIS52521.2021.9701961.
- [16] R. Andika and Suhartjito, "Effects of using wordnet and spelling checker on classification methods in sentiment analysis for datasets using Bahasa," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 25, no. 3, pp. 1662–1671, 2022, doi: 10.11591/ijeecs.v25.i3.pp1662-1671.
- [17] S. W. Iriananda, R. P. Putra, and K. S. Nugroho, "Analisis Sentimen Dan Analisis Data Eksploratif Ulasan Aplikasi Marketplace Google Playstore," *4th Conf. Innov. Appl. Sci. Technol. (CIASTECH 2021)*, no. Ciastech, pp. 473–482, 2021.
- [18] Hubert, P. Phoenix, R. Sudaryono, and D. Suhartono, "Classifying Promotion Images Using Optical Character Recognition and Naïve Bayes Classifier," *Procedia Comput. Sci.*, vol. 179, no. 2020, pp. 498–506, 2021, doi: 10.1016/j.procs.2021.01.033.