

ANALISIS SENTIMEN ULASAN APLIKASI LITA DI PLAY STORE MENGUNAKAN ALGORITMA NAIVE BAYES

Tia Septiani Gumilar¹, Rini Astuti², Yudhistira Arie Wijaya³

¹ Teknik Informatika, Stmik Ikmi Cirebon

² Sistem Infrmasi, Stmik Likmi Bandung

³ Sistem Informasi, Stimik Ikmi Cirebon

Jl. Perjuangan No. 10 B Majasem Kec. Kesambi Kota Cirebon Tlp. 0231) 490480 - 490481
tiagumilar292@gmail.com

ABSTRAK

Aplikasi media sosial yang memberikan Jasa Teman Mabur yang dapat menemani bermain beragam jenis game online, seperti *Mobile Legends*, *PUBG Mobile*, *Free Fire*, dan lainnya. Lita merupakan platform dengan konsep baru dalam dunia game. Kamu dapat bertemu dengan Pro Player ataupun gamer perempuan yang cantik. Ulasan pengguna pada aplikasi mobile merupakan sumber data penting bagi pengembang aplikasi untuk mengetahui respon dari pengguna dan meningkatkan kualitas aplikasi. Ulasan pengguna biasanya mengandung ungkapan subjektif yang mencerminkan sentimen positif atau negatif. Klasifikasi sentimen pada ulasan mobile diperlukan untuk secara otomatis menganalisis sentimen positif dan negatif dari sejumlah besar data ulasan. Penelitian ini bertujuan untuk mengklasifikasikan sentimen ulasan aplikasi Lita di Play Store ke dalam dua kelas yaitu positif dan negatif, dengan menggunakan metode algoritma Naive Bayes. Data yang digunakan dalam penelitian ini adalah 1000 ulasan aplikasi Lita di Play Store dalam bahasa Indonesia, diambil secara acak. Data ulasan mengandung kata tidak penting yang dihapus melalui text preprocessing. Kemudian, pembobotan dilakukan dengan metode TF-IDF untuk mengetahui pentingnya sebuah istilah untuk suatu dokumen. Hasil dari penelitian ini mengenai Klasifikasi Data Sentimen Ulasan Pengguna Lita Google Play Store berjumlah 1000 data, dapat disimpulkan bahwa ulasan pengguna Lita tergolong positif dengan hasil presentase 95% nilai accuracy, 95% nilai precision, dan 95% nilai recall nya.

Kata kunci : Analisis sentimen, Ulasan Pengguna, Aplikasi Lita, Algoritma Naive Bayes.

1. PENDAHULUAN

Perkembangan cepat dibidang informatika telah mengubah cara kita berinteraksi dengan teknologi khususnya melalui aplikasi mobile Lita. Perubahan paradigma ini mencerminkan ketergantungan masyarakat pada aplikasi perangkat lunak untuk memenuhi berbagai kebutuhan. Aplikasi Lita yang diunduh dan diulas melalui platform Play Store menjadi trend yang tidak dapat dihindari. Penelitian ini bertujuan untuk menganalisis sentimen dari ulasan pengguna terhadap aplikasi Lita, dengan menggunakan metode algoritma Naive Bayes. Analisis ini dapat memberikan pemahaman mendalam mengenai sentimen positif, negatif, atau netral dalam ulasan sehingga dapat mengidentifikasi aspek-aspek yang mendukung atau mengecewakan pengguna.

Dalam era modern yang dikuasai oleh aplikasi mobile, pengguna aplikasi mobile seperti aplikasi Lita telah mencapai tingkat popularitas yang signifikan di Play Store. Meskipun aplikasi ini mendapat respons positif dari sejumlah pengguna, tantangan utama yang muncul adalah kompleksitas analisis sentimen pada ulasan yang diberikan oleh para pengguna. Tantangan ini menjadi semakin kritis seiring dengan pertumbuhan jumlah pengguna yang terus meningkat. Relevansi permasalahan ini terletak pada dampaknya terhadap pengembangan dan peningkatan kualitas aplikasi. Terdapat kesenjangan dalam literatur terkait dengan analisis sentimen spesifik pada ulasan aplikasi mobile, khususnya dengan penggunaan metode

algoritma Naive Bayes. Kesenjangan ini menghambat pemahaman mendalam terkait kepuasan pengguna dan aspek-aspek yang perlu ditingkatkan pada aplikasi. Isu-isu aktual seperti ketidakpastian pasar dan persaingan yang semakin ketat dalam industri aplikasi mobile menambah urgensi untuk mengatasi permasalahan analisis sentimen ini. Oleh karena itu, penelitian ini memiliki nilai signifikan dalam mengisi kesenjangan pengetahuan, merespons tren terbaru dalam analisis sentimen aplikasi dan memberikan kontribusi pada pemahaman praktis yang mendalam dalam konteks pengembangan aplikasi mobile yang kompetitif. Analisis sentimen yang akurat dapat memberikan wawasan mendalam mengenai kepuasan pengguna, memungkinkan pengembang untuk mengambil langkah-langkah yang sesuai untuk meningkatkan kualitas dan daya saing aplikasi, serta menanggapi kebutuhan pasar yang dinamis. Dengan demikian, penelitian ini memiliki relevansi yang tinggi dalam mengatasi tantangan praktis yang dihadapi dalam pengembangan aplikasi di era digital saat ini.

Menurut penelitian sebelumnya analisis sentimen dapat diklasifikasikan dengan baik menggunakan salah satu metode algoritma naive bayes. Menurut penelitian Penggunaan Algoritma Naive Bayes dalam Menganalisis Sentimen terhadap Aplikasi Jobstreet bertujuan untuk menganalisis sentimen pengguna terhadap aplikasi Jobstreet menggunakan metode Naive Bayes. Dalam konteks ini, algoritma Naive Bayes digunakan untuk mengklasifikasikan sentimen

pengguna menjadi kategori positif, negatif, atau netral berdasarkan ulasan dan feedback yang diberikan terkait pengalaman mereka menggunakan aplikasi Jobstreet[1]. Menurut penelitian Analisis Sentimen Terhadap Pengguna Qris (Quick Respond Code Indonesian Standart) Pada Twitter Menggunakan Metode Naive Bayes bertujuan untuk menganalisis sentimen pengguna terhadap teknologi QRIS (Quick Response Code Indonesian Standard) menggunakan metode Naive Bayes pada platform Twitter. Dalam konteks ini, penelitian ini bertujuan untuk mengidentifikasi dan mengklasifikasikan sentimen pengguna (baik itu positif, negatif, atau netral) terhadap QRIS berdasarkan data dan konten yang diunggah oleh pengguna Twitter[2]. Menurut penelitian Analisis Sentimen Ulasan Aplikasi Transportasi Online Menggunakan Multinomial Naive Bayes dan Query Expansion Ranking memahami bagaimana pengguna merasakan dan merespons aplikasi transportasi online, serta mengelompokkan ulasan-ulasan tersebut berdasarkan sentimennya. Dengan menggunakan metode bertujuan Multinomial Naive Bayes, penelitian ini bertujuan untuk mengklasifikasikan sentimen-sentimen tersebut dengan akurasi tinggi, memungkinkan penyedia layanan untuk mendapatkan wawasan tentang kekuatan dan kelemahan aplikasi mereka berdasarkan umpan balik pengguna[3].

Tujuan utama dari penelitian ini adalah untuk melakukan analisis sentimen pada ulasan aplikasi Lita di Play Store menggunakan metode algoritma Naive Bayes. Penelitian ini bertujuan untuk menyelidiki dan memahami sentimen pengguna terkait aplikasi Lita baik positif, negatif, atau netral. Selain itu, penelitian ini bertujuan untuk mengidentifikasi faktor-faktor yang mendukung atau mengecewakan pengguna dalam penggunaan aplikasi ini. Hasil analisis sentimen yang diperoleh dapat menjadi panduan berharga bagi pengembang aplikasi Lita untuk meningkatkan kualitas layanan mereka, merespons umpan balik pengguna, dan mengoptimalkan fitur-fitur yang paling dihargai oleh pengguna. Dengan demikian, penelitian ini tidak hanya bersifat akademis tetapi juga memberikan dampak nyata dalam pengembangan aplikasi mobile yang memenuhi kebutuhan dan ekspektasi pengguna secara lebih efektif.

Dalam menghadapi kompleksitas analisis sentimen pada ulasan aplikasi Lita di Play Store, penelitian ini akan menggunakan metode algoritma Naive Bayes. Metode ini dipilih karena kemampuannya yang terbukti dalam klasifikasi sentimen dalam teks, khususnya pada ulasan pengguna. Proses analisis sentimen akan melibatkan langkah-langkah yang melibatkan pemodelan probabilistik untuk memprediksi sentimen positif, negatif, atau netral pada setiap ulasan. Pendekatan eksperimental akan melibatkan pengumpulan dataset ulasan pengguna Lita dari Play Store, yang akan digunakan sebagai bahan penelitian. Teknik analisis data yang terintegrasi dengan metode ini akan

memungkinkan identifikasi pola-pola sentimen yang signifikan dalam ulasan-ulasan tersebut.

Hasil penelitian ini memiliki potensi implikasi yang signifikan dalam meningkatkan pemahaman kita tentang sentimen pengguna terhadap aplikasi mobile, khususnya dalam konteks aplikasi Lita di Play Store. Dengan menerapkan metode algoritma Naive Bayes pada analisis sentimen, temuan penelitian ini memiliki potensi untuk memberikan pemahaman yang lebih mendalam tentang kecenderungan dan kepuasan pengguna terhadap fitur-fitur dan kinerja aplikasi tersebut. Implikasi positif dari penelitian ini adalah bahwa temuan analisis sentimen dapat menjadi panduan bagi pengembang aplikasi Lita untuk memperbaiki aspek-aspek yang mendukung dan mengatasi masalah-masalah yang menciptakan sentimen negatif. Selain itu, temuan dari penelitian ini memiliki potensi untuk meningkatkan pemahaman terhadap kecenderungan dan preferensi pengguna dalam memanfaatkan aplikasi mobile, sehingga dapat menjadi dasar untuk pengembangan aplikasi selanjutnya. Para praktisi di industri teknologi dapat memanfaatkan temuan penelitian ini untuk meningkatkan pengalaman pengguna pada aplikasi mobile secara umum. Analisis sentimen yang lebih akurat dan terperinci dapat membantu praktisi dalam mengoptimalkan strategi pengembangan, memahami kebutuhan pasar, dan merespons perubahan dinamis dalam persepsi pengguna terhadap aplikasi. Bagi peneliti di bidang Informatika, penelitian ini memberikan sumbangan berharga pada literatur terkait analisis sentimen pada ulasan aplikasi mobile. Hasil penelitian dapat menjadi sumber referensi yang penting untuk penelitian-penelitian selanjutnya terkait metode analisis sentimen dalam konteks aplikasi mobile.

2. LITERATURE REVIEW

Hasil literature review yang telah dilakukan pada jurnal-jurnal penelitian terkait topik Analisis Sentimen dapat dijabarkan sebagai berikut :

Analisis Sentimen Aplikasi Pedulilindungi Menggunakan Metode Naive Bayes Classifier (NBC) dan Support Vector Machine (SVM). Hasil evaluasi terhadap aplikasi PeduliLindungi menunjukkan bahwa pembaruan yang dilakukan dengan memproses ulasan pengguna di Google Play Store menggunakan analisis sentimen memberikan hasil yang positif. Dua metode yang digunakan, Naive Bayes Classifier (NBC) dan Support Vector Machine (SVM), menunjukkan kinerja yang baik. Algoritma NBC mencapai akurasi sebesar 76%, presisi 76%, recall 82%, dan skor f1 79%. Sementara itu, algoritma SVM mencapai tingkat akurasi sebesar 80%, presisi 83%, recall 80%, dan skor f1 81%[4].

Analisis Sentimen Berbasis Aspek Ulasan Pelanggan Restoran Menggunakan LSTM Dengan Adam Optimizer. Hasil penelitian menunjukkan bahwa analisis sentimen berbasis aspek memiliki akurasi sebesar 78.7%, sementara kategori aspek

memiliki akurasi sebesar 78%. Kedua aspek ini saling berhubungan dan membantu memahami opini pelanggan restoran pada platform TripAdvisor[5].

Analisis Sentimen Calon Presiden 2024 Menggunakan Algoritma SVM Pada Media Sosial Twitter. Hasil penelitian ini menyoroti pentingnya analisis sentimen dalam mengevaluasi opini publik terhadap tokoh politik calon presiden 2024 yang sedang ramai dibahas di media sosial, khususnya di Twitter. Fokus penelitian adalah meningkatkan kinerja algoritma klasifikasi sentimen dengan menggantikan algoritma Naïve Bayes yang memiliki tingkat akurasi rendah dengan algoritma SVM. Melalui pengambilan data Twitter terkait calon presiden tertentu, penelitian berhasil mencapai tingkat akurasi yang mencolok. Algoritma SVM menunjukkan peningkatan signifikan dengan nilai akurasi rata-rata mencapai 98,61%, sedangkan pada penelitian sebelumnya dengan Naïve Bayes hanya mencapai 73,86%. Selain itu, evaluasi sentimen menunjukkan bahwa Ganjar Pranowo mendapat proporsi sentimen positif tertinggi, yaitu 55%, sementara sentimen negatif tertinggi adalah milik Anies dengan 89%[6].

Analisis Sentimen Kepuasan Pengguna Aplikasi Whatsapp Menggunakan Algoritma Naïve Bayes Dan Support Vector Machine. Hasil pengujian pada 1500 komentar pengguna di Google Play Store menunjukkan bahwa tingkat akurasi dari algoritma Support Vector Machine lebih tinggi, mencapai 77,00%, sementara Naïve Bayes mencapai 70,40%. Selain itu, nilai AUC (Area Under the Curve) yang tinggi, khususnya pada Support Vector Machine dengan nilai 0,876, menunjukkan bahwa algoritma ini efektif dalam mengkategorikan komentar sebagai positif atau negatif terkait kepuasan pengguna WhatsApp[7].

Analisis Sentimen Pada Situs Google Review dengan Naïve Bayes dan Support Vector Machine. Penelitian ini menggunakan studi kasus Summarecon Mal Bekasi yang menghasilkan volume besar komentar, mencapai 60.000 ulasan. Untuk mengidentifikasi sentimen dari komentar-komentar ini, diterapkan alat bantu komputasi berupa sentiment analysis dengan metode Naïve Bayes dan Support Vector Machine. Pengambilan data dilakukan melalui teknik web scraping, diikuti dengan proses pre-processing dan pelabelan menggunakan kamus Lexicon. Analisis sentimen dilakukan untuk mengkategorikan komentar sebagai positif atau negatif. Hasil penelitian menunjukkan bahwa metode Naïve Bayes dan Support Vector Machine memiliki akurasi yang cukup tinggi dalam melakukan analisis sentimen terhadap 2.143 komentar Summarecon Mal Bekasi, dengan tingkat akurasi berturut-turut 80,95% dan 100%[8].

Analisis Sentimen Pengguna Aplikasi Bank Syariah Indonesia Dengan Menggunakan Algoritma Support Vector Machine (Svm). Dalam penelitian ini, 5 ahli bahasa memberikan label manual pada 55.416 ulasan dan peringkat pengguna. Setelah

menghilangkan data duplikat, 13.568 data digunakan dalam penelitian. Analisis sentimen dilakukan dengan menggunakan metode Support Vector Machine (SVM). Hasil penelitian menunjukkan presisi dan daya ingat sekitar 87,309%, dengan akurasi pelatihan sebesar 85,87%. Proyeksi analisis sentimen menunjukkan tingkat akurasi 85,87%, sementara presisi pelatihan dan pengujian adalah 86,958%[9].

Analisis Sentimen Pengguna Media Sosial Terhadap Aplikasi M-Health Peduli Lindungi Dengan Metode Lexicon Based Dan Naïve Bayes. Analisis sentimen dilakukan terhadap aplikasi Peduli Lindungi pada media sosial YouTube, TikTok, dan Twitter menggunakan kombinasi Lexicon Based dan Naïve Bayes. Hasil analisis menunjukkan akurasi total sebesar 91%, presisi 94%, recall 82%, dan f1_scores 86%. Namun, ketika hasil dianalisis berdasarkan platform media sosial, ditemukan perbedaan. Pada YouTube, akurasi sebesar 90%, presisi 93%, recall 81%, dan f1_scores 84%. Di sisi lain, untuk Twitter, model tidak memberikan hasil yang baik dengan akurasi 70%, presisi 23%, recall 33%, dan f1-scores 28%. Ini menunjukkan bahwa gabungan algoritma Lexicon Based dan Naïve Bayes kurang efektif dalam mengklasifikasikan data dari Twitter yang diambil menggunakan tweet-harvest. Hasil klasifikasi sentimen menunjukkan bahwa TikTok dan YouTube cenderung memiliki opini netral dari masyarakat, sementara Twitter didominasi oleh opini positif[10].

Sentimen Analisis Aplikasi Belajar Online Menggunakan Klasifikasi SVM. Data diperoleh melalui teknik scraping data dengan bantuan library google-play-scraper, mengumpulkan 30.000 data yang dibagi menjadi 10.000 data untuk masing-masing Ruang Guru, Zenius, dan Quipper. Proses web scraping terdiri dari tahap Fetching, Extraction, dan Transformation. Setelah itu, data mengalami tahap preprocessing, termasuk normalisasi, case folding, cleaning, tokenizing, dan stopword. Data dibagi menjadi 90% data latih dan 10% data uji, dengan memberikan label nilai 1 (positif), 0 (netral), dan -1 (negatif). Hasil klasifikasi sentimen menggunakan SVM menunjukkan bahwa Ruang Guru memiliki nilai positif tertinggi, meskipun respon pengguna memberikan nilai positif untuk semua aplikasi. Akurasi klasifikasi sentimen dengan SVM adalah 99% untuk Ruang Guru, 96% untuk Zenius, dan 82% untuk Quipper[11].

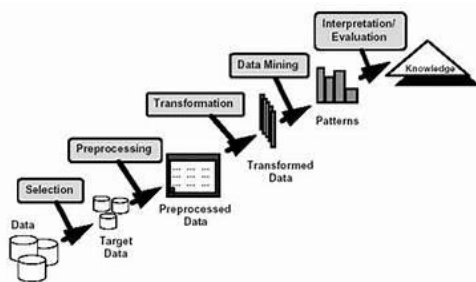
Analisis Sentimen Terhadap Pengguna Gojek Menggunakan Metode K-Nearest Neighbors. Dalam penelitian ini muncul masalah dalam mengklasifikasikan respon pengguna Twitter ke dalam kategori positif, negatif, dan netral, yang jika dilakukan secara manual akan memakan waktu lama. Oleh karena itu, diperlukan sistem klasifikasi menggunakan metode K-Nearest Neighbor untuk memudahkan pengklasifikasian respon pengguna Gojek di Twitter. Hasil klasifikasi menggunakan metode K-Nearest Neighbor menunjukkan tingkat akurasi sebesar 79,43%, dengan nilai k=15, dan dapat

digunakan oleh perusahaan Gojek sebagai bahan evaluasi dan penilaian terhadap layanannya[12].

Perbandingan Metode KNN, Decision Tree, dan Naïve Bayes Terhadap Analisis Sentimen Pengguna Layanan BPJS. Hasil penelitian menunjukkan bahwa analisis sentimen terhadap layanan BPJS di Twitter menggunakan metode KNN mencapai tingkat akurasi 95.58%, dengan class precision untuk prediksi negatif sebesar 45.00%, prediksi positif 0.00%, dan prediksi netral 96.83%. Pada metode Decision Tree, tingkat akurasi mencapai 96.13%, dengan class precision untuk prediksi negatif sebesar 55.00%, prediksi positif 0.00%, dan prediksi netral 97.28%. Terakhir, metode Naïve Bayes mencapai akurasi 89.14%, dengan class precision untuk prediksi negatif sebesar 16.67%, prediksi positif 1.64%, dan prediksi netral 98.40%[13].

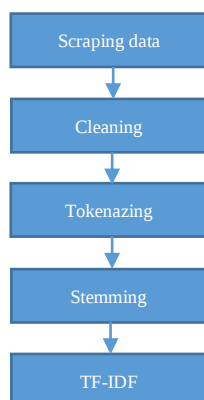
3. METODE PENELITIAN

Metode penelitian ini menggunakan metode klasifikasi dengan alur dibawah ini. Penggunaan Knowledge Discovery on Database (KDD) dianggap sebagai metode yang sesuai dalam konteks penelitian ini. KDD merupakan suatu proses pengekplorasian data untuk meraih informasi baru dari database yang besar. Proses ini terdiri dari empat tahap utama, yakni Pengumpulan Data, Pra-pemrosesan Data, Pengekplorasian Data, dan evaluasi[14]. Alur Proses KDD Modelling dapat dilihat pada gambar 1 dibawah ini :



Gambar 1 Proses Metode KDD

Berdasarkan gambar 1 diatas, tahapan metode Knowledge Discovery on Database (KDD) dapat dijelaskan :



3.1. Scraping Data

Teknik pengumpulan data menggunakan Scraping Data dengan cara mengambil ulasan/review

aplikasi lita melalui Google Play Store, mengambil 1000 ulasan dari aplikasi lita secara acak. Informasi review Google Play Store di scraping menggunakan program Python dan hasilnya disimpan dalam format .csv. Untuk mengkategorikan data termasuk ke dalam klasifikasi positif dan negatif data yang diperoleh dilabeli perkata dengan manual.

3.2. Cleaning

Data cleaning untuk menghapus karakter khusus, tanda baca, dan simbol yang tidak diperlukan. Konversi teks menjadi huruf kecil (lowercasing) untuk konsistensi. Hapus kata-kata penghubung (stop words) yang umumnya tidak memberikan informasi signifikan.

3.3. Tokenizing

Tokenisasi/tokenazing adalah langkah untuk menetapkan batas kata pada suatu kalimat dengan menggunakan tanda baca koma (.). Langkah ini dilakukan setelah proses case folding untuk menghilangkan tanda baca yang tidak diperlukan.

3.4. Stemming

Proses stemming ini memanfaatkan basis data kata dari Sastrawi untuk mengubah setiap kata ke bentuk dasarnya dengan maksud mengelompokkan kata-kata yang memiliki makna serupa.

3.5. Tf-Idf

Proses pengklasifikasian suatu baris data dalam proses evaluasi ada 4 kemungkinan yang akan terjadi True positive (TP), False Negative (FN), False positive (FP) dan True negative (TN)[15].

4. HASIL DAN PEMBAHASAN

4.1. Scraping data

Tahapan dimulai dengan proses pengumpulan data. Proses pengumpulan data yang diperoleh dari hasil scraping sebanyak 1000 data, terdiri dari 1 atribut yaitu content seperti pada gambar 2 dibawah.

	content
238	Saya senang dengan apk ini, menemani waktu lua...
108	aplikasi yang bagus. bisa cari teman mabar dis...
916	Lita seru bngtt!!!! Selain bisa mabar dgn pro ...
189	Jual aplikasi cuma mau cari keuntungan pribad...
227	Aplikasi sangat membantu untuk mencari teman m...

Gambar 2 Read Excel.

Berdasarkan pelabelan dalam penelitian ini data yang sudah terkumpul diberikan label secara manual dengan menambahkan kolom label di sisi kanan content. Pelabelan data dilakukan mengacu kepada kamus berkonotasi negatif dan positif menjadi 2 kelas sentimen yang akan digunakan, yaitu positif dan negatif. Kamus dapat dilihat pada link <https://github.com/SadamMahendra/ID-NegPos>[16].

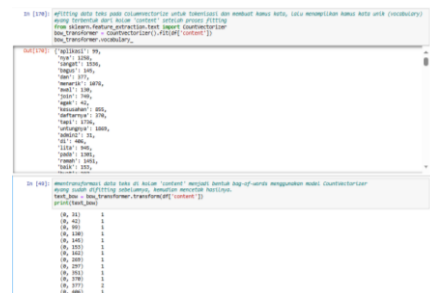
Tabel 1 pelabelan

4.2. Cleaning

[illegible]

Gambar 5 Stemming

Fitting data teks pada ColumnVectorize untuk tokenisasi dan membuat kamus kata, lalu menampilkan kamus kata unik (vocabulary) yang terbentuk dari kolom 'content' setelah proses fitting terdapat pada gambar 6 dibawah.



Gambar 6 Proses dan Hasil Fitting

Proses mentransformasikan representasi bag-of-words menjadi bentuk pembobotan tfidf, kemudian mencetak model beserta matriks tf-idf yang terbentuk, serta shape-nya. Berguna sebagai pembuatan fitur tfidf dari data teks untuk proses selanjutnya terdapat pada gambar 6 dibawah.

	content	text	text_clean	text_wordcount
1	Apakah yang sangat banyak dan bermacam...	apakah yang sangat banyak dan bermacam...	apakah ya banyak memang banyak dan bermacam...	10
2	Apakah ya banyak untuk dan bermacam...	apakah ya banyak untuk dan bermacam...	apakah ya banyak dan bermacam untuk dan bermacam...	10
3	Apakah banyak dan bermacam sekali. Untuk...	apakah banyak dan bermacam sekali untuk...	apakah banyak dan bermacam sekali untuk...	10
4	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
5	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
6	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
7	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
8	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
9	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10
10	Apakah yang banyak sangat dan bermacam...	apakah yang banyak sangat dan bermacam...	apakah banyak sangat dan bermacam...	10

Gambar 3 Cleaning

Pada gambar 4 merupakan hasil dari text tokens. Tokenisasi/tokenazing adalah langkah untuk menetapkan batas kata pada suatu kalimat dengan menggunakan tanda baca koma (.). Langkah ini dilakukan setelah proses case folding untuk menghilangkan tanda baca yang tidak diperlukan.

(No)	content	label	test_cian	test_Silwerid	test_Jukuna
1	Aplahira ma sangat bagas dan Aurid jo	protest	Aplahira ma sangat bagas dan Aurid jo	Aplahira ma sangat bagas dan kesamban dadi	Aplahira, ma, sangat, bagas, kesamban-1
2	Aplahira jo bagas Mata gema	protest	Aplahira jo bagas Mata gema	Aplahira jo bagas dan kesamban dadi	Aplahira, jo, bagas, dan, kesamban, kesamban-1
3	Aplahira bagas dan remasan Tikisa mudi	protest	Aplahira bagas dan remasan Tikisa mudi	Aplahira bagas remasan Tikisa mudi	Aplahira, bagas, remasan, Tikisa, mudi
4	Aplahira jo remasan bagas dan huul gema	protest	Aplahira jo remasan bagas dan huul gema	Aplahira remasan bagas dan huul gema	Aplahira, remasan, bagas, dan, huul, gema
5	Aplahira bagas bagas bagas dan gema	protest	Aplahira bagas bagas bagas dan gema	Aplahira bagas bagas bagas dan gema	Aplahira, bagas, bagas, bagas, dan, gema
6	Aplahira ma hgl adremasa samah bagas	protest	Aplahira ma hgl adremasa samah bagas	Aplahira ma hgl adremasa samah bagas	Aplahira, ma, hgl, adremasa, samah, bagas
7	Aplahira ma sangat bagas	protest	Aplahira ma sangat bagas	Aplahira bagas	Aplahira, bagas
8	Aplahira ma sangat bagas dan remasan gema	protest	Aplahira ma sangat bagas dan remasan gema	Aplahira bagas bagas dan remasan gema	Aplahira, bagas, remasan, gema

Gambar 4 Tokenezing

Pada gambar 5 merupakan hasil dari text stemming. Proses stemming ini memanfaatkan basis data kata dari Sastrawi untuk mengubah setiap kata ke bentuk dasarnya dengan maksud mengelompokkan kata-kata yang memiliki makna serupa.

Berdasarkan nilai akurasi, digunakan sebagai indikator untuk mengukur sejauh mana data sentimen berhasil diklasifikasikan dengan benar. Evaluasi merupakan salah satu tahap krusial dalam pengembangan program; melalui tahap evaluasi, pengguna dapat menilai setiap aspek akurasi, presisi, dan recall dari program yang mereka bangun. Salah satu cara untuk mengevaluasi hasil analisis sentimen adalah dengan menggunakan Confusion Matrix. Confusion Matrix adalah suatu metode yang memudahkan program untuk menentukan nilai penting dalam mengevaluasi akurasi, presisi, dan recall [11].

```
In [114]: membagi data menjadi data training dan testing dengan test_size = 0.20 dan random state nya 0
from sklearn.model_selection import train_test_split
X_train, X_test, y_train, y_test = train_test_split(data_clean['content'], data_clean['label'],
                                                test_size = 0.20,
                                                random_state = 0)

In [117]: melakukan vektorisasi data teks menjadi bentuk TF-IDF untuk data latih (X_train) dan data uji (X_test).
melakukan ekstraksi fitur TF-IDF dari data teks sebelum digunakan untuk pemodelan.
shape dicetak untuk memastikan data latih dan data uji sesuai
from sklearn.feature_extraction.text import TfidfVectorizer

tfidf_vectorizer = TfidfVectorizer()
tfidf_train = tfidf_vectorizer.fit_transform(X_train)
tfidf_test = tfidf_vectorizer.transform(X_test)

print(X_train.shape)
print(y_train.shape)
print(X_test.shape)
print(y_test.shape)

(200,)
(200,)
(200,)
(200,)
```

Gambar 8 Split data

```
import pandas as pd
from sklearn.feature_extraction.text import TfidfVectorizer
from sklearn.naive_bayes import MultinomialNB
from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report
from sklearn.metrics import confusion_matrix

clf = MultinomialNB()
clf.fit(X_train, y_train)
predicted = clf.predict(X_test)

print("MultinomialNB Accuracy:", accuracy_score(y_test, predicted))
print("MultinomialNB Precision:", precision_score(y_test, predicted, average="micro", pos_label="Negatif"))
print("MultinomialNB Recall:", recall_score(y_test, predicted, average="micro", pos_label="Negatif"))
print("MultinomialNB f1_score:", f1_score(y_test, predicted, average="micro", pos_label="Negatif"))

print(f"Confusion Matrix:\n {confusion_matrix(y_test, predicted)}")
print("=====\n")
print(classification_report(y_test, predicted, zero_division=0))

# load dataset
data_clean = pd.read_csv('hasil_TextPreProcessing_lita.csv')
```

Gambar 9 Script Hasil

Pada gambar 8 diatas merupakan splitting data yaitu proses membagi dataset menjadi dua atau lebih subset yang berbeda untuk tujuan tertentu. Proses splitting data yang dilakukan adalah data training 80 data testing 20.

Pada gambar 9 diatas merupakan script untuk hasil proses evaluasi dan pada tabel 2 merupakan uraian script untuk memperoleh hasil dari processing lita.

Tabel 2 Uraian Script

Script	Uraian
import pandas as pd	Mengimpor Library Pandas
from sklearn.feature_extraction.text import TfidfVectorizer	Membuat Vektorisasi Teks Dengan TF-IDF
from sklearn.naive_bayes import MultinomialNB	Mengimpor Model Naive Bayes Multinomial Dari Sklearn
from sklearn.metrics import accuracy_score, precision_score,	Membagi Dataset Menjadi Data Latih Dan Data Uji
from sklearn.model_selection import train_test_split	Membuat Laporan Klasifikasi Yang Berisi Beberapa Metrik
from sklearn.metrics import classification_report	Membuat Confusion Matrix
from sklearn.metrics import confusion_matrix	Membuat Instance Model Multinomialnb
clf = MultinomialNB()	Melatih Model Dengan Data Latih
clf.fit(X_train, y_train)	Membuat Prediksi Pada Data Uji
predicted = clf.predict(X_test)	Mencetak Nilai Akurasi
print("MultinomialNB Accuracy:", accuracy_score(y_test, predicted))	Mencetak Nilai Precision
print("MultinomialNB Precision:", precision_score(y_test, predicted, average="micro", pos_label="Negatif"))	Mencetak Nilai Recall
print("MultinomialNB Recall:", recall_score(y_test, predicted, average="micro", pos_label="Negatif"))	Mencetak Nilai F1-Score
print("MultinomialNB f1_score:", f1_score(y_test, predicted, average="micro", pos_label="Negatif"))	Mencetak Confusion Matrix
print(f"confusion_matrix:\n {confusion_matrix(y_test, predicted)}")	Mencetak Classification Report
print(classification_report(y_test, predicted, zero_division=0))	Menyimpan Dataset Csv
data_clean = pd.read_csv('hasil_TextPreProcessing_lita.csv')	

Dari hasil training 80 dan testing 20 pada gambar 8 hasil dari nilai akurasi, precision, recall dan f1_score spliting data training 80 data testing 20 menghasilkan Akurasi 95%, Precision 95%, Recall 95%.

```
MultinomialNB Accuracy: 0.95
MultinomialNB Precision: 0.95
MultinomialNB Recall: 0.95
MultinomialNB f1_score: 0.9500000000000001
confusion_matrix:
[[ 6  1  1]
 [ 0  1  8]
 [ 0  0 183]]
=====
              precision    recall  f1-score   support

negatif         1.00        0.75        0.86         8
netral           0.50        0.11        0.18         9
positif          0.95        1.00        0.98        183

accuracy                   0.95         200
macro avg                  0.82        0.62        0.67         200
weighted avg               0.93        0.95        0.94         200
```

Gambar 7 Hasil Akurasi 80, 20

Tabel 3 Hasil Evaluasi

Training	Testing	Akurasi
80	20	95%
70	30	96%
60	40	96%
50	50	95%
90	10	95%

Pada tabel 3 merupakan hasil akurasi dari spliting data 70 dan 20 terdapat hasil 96%, 60 40 terdapat hasil 96%, 50 50 terdapat hasil 95% dan 90 10 terdapat hasil 95%. Pada gambar 10 dan 11 terdapat hasil dari visualisasi sentimen positif dan negatif.



Gambar 10 Hasil sentimen positif



Gambar 10 Hasil sentimen negatif

5. KESIMPULAN & SARAN

Tingkat akurasi yang didapatkan adalah 95% yang artinya aplikasi lita menunjukkan tingkat akurasi yang sangat tinggi, sebagian besar pengguna aplikasi Lita adalah positif. Hanya sedikit data negatif yang terdeteksi model. Penelitian ini menghasilkan nilai precision 95% dan recall 95% dengan rasio data training 80% dan data testing 20%.

Karena aplikasi Lita tergolong tanggapan positif dari penggunaanya, disarankan untuk terus mempertahankan dan meningkatkan kualitas serta manfaat aplikasi ini. Pastikan untuk senantiasa memantau dan menampung umpan balik pengguna sehingga aplikasi tetap relevan dan berguna bagi mereka.

DAFTAR PUSTAKA

- [1] B. K. Widodo, *PENERAPAN ALGORITMA NAÏVE BAYES UNTUK ANALISIS SENTIMEN PENGGUNAAN APLIKASI JOBSTREET DI GOOGLE PLAY STORE*. repository.upnvj.ac.id, 2022. [Online]. Available: <https://repository.upnvj.ac.id/19665/>
- [2] P. Paramita and A. Ibrahim, "ANALISIS SENTIMEN TERHADAP PENGGUNA QRIS (QUICK RESPOND CODE INDONESIAN STANDART) PADA TWITTER MENGGUNAKAN METODE NAÏVE ...," *JOISIE (Journal ...)*, 2023, [Online]. Available: <https://ejournal.pelitaIndonesia.ac.id/ojs32/index.php/JOISIE/article/view/2963>
- [3] Y. A. V Gunawana, N. A. S. ERa, I. B. M. Mahendraa, and ..., "Analisis Sentimen Ulasan Aplikasi Transportasi Online Menggunakan Multinomial Naïve Bayes dan Query Expansion Ranking," *Jurnal Elektronik Ilmu ...*. ojs.unud.ac.id. [Online]. Available: <https://ojs.unud.ac.id/index.php/JLK/article/download/88930/47224>
- [4] J. Elektronik and D. Komputer, "Analisa Sentimen Aplikasi Pedulilindungi Dengan Metode Nbc Dan Svm," vol. 16, no. 1, pp. 100–108, 2023, [Online]. Available: <https://journal.stekom.ac.id/index.php/elkom/page100>
- [5] W. Wardianto, F. Farikhin, and D. M. Kusumo Nugraheni, "Analisis Sentimen Berbasis Aspek Ulasan Pelanggan Restoran Menggunakan LSTM Dengan Adam Optimizer," *JOINTECS (Journal Inf. Technol. Comput. Sci.)*, vol. 8, no. 2, p. 67, 2023, doi: 10.31328/jointecs.v8i2.4737.
- [6] A. P. Nardilasari, A. L. Hananto, and ..., "Analisis Sentimen Calon Presiden 2024 Menggunakan Algoritma SVM Pada Media Sosial Twitter," *JOINTECS ...*. ojs.publishing-widyagama.ac.id, 2023. [Online]. Available: <https://ojs.publishing-widyagama.ac.id/index.php/jointecs/article/download/4265/2464>
- [7] A. Saepulrohman, S. Saepudin, and ..., "Analisis Sentimen Kepuasan Pengguna Aplikasi Whatsapp Menggunakan Algoritma Naïve Bayes Dan Support Vector Machine," *@ is Best Account. ...*, 2021, [Online]. Available: <http://ojs.unikom.ac.id/index.php/aisthebest/article/view/4919>
- [8] M. I. Putri and I. Kharisudin, "Analisis Sentimen Pengguna Aplikasi Marketplace Tokopedia Pada Situs Google Play Menggunakan Metode Support Vector Machine (SVM), Naïve Bayes, dan ...," *Prism. Pros. Semin. Nas. ...*, 2022, [Online]. Available: <https://journal.unnes.ac.id/sju/index.php/prisma/article/view/54577>
- [9] R. O. Mardiyanto, K. Kusriani, and ..., "ANALISIS SENTIMEN PENGGUNA APLIKASI BANK SYARIAH INDONESIA DENGAN MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE (SVM)," ... *Teknol. Inf. dan ...*, 2023, [Online]. Available: <https://www.jurnal.stmiksznw.ac.id/index.php/teknimedia/article/view/85>
- [10] R. I. Syah, H. Hoiriyah, and M. Walid, "ANALISIS SENTIMEN PENGGUNA MEDIA SOSIAL TERHADAP APLIKASI M-HEALTH PEDULI LINDUNG DENGAN METODE LEXICON BASED DAN NAÏVE BAYES," *Indones. J. ...*, 2023, [Online]. Available: <https://ejournal.almaata.ac.id/index.php/IJUBI/article/view/3275>
- [11] A. A. Munandar, F. Farikhin, and C. E. Widodo, "Sentimen Analisis Aplikasi Belajar Online Menggunakan Klasifikasi SVM," *JOINTECS (Journal Inf. Technol. Comput. Sci.)*, vol. 8, no. 2, p. 77, 2023, doi: 10.31328/jointecs.v8i2.4747.
- [12] F. R. Irawan, A. Jazuli, and T. Khotimah, "Analisis Sentimen Terhadap Pengguna Gojek Menggunakan Metode K-Nearset Neighbors,"

- JIKO (Jurnal Inform. ...*, 2022, [Online]. Available: <http://ejournal.unkhair.ac.id/index.php/jiko/article/view/4267>
- [13] R. Puspita and A. Widodo, "Perbandingan Metode KNN, Decision Tree, dan Naïve Bayes Terhadap Analisis Sentimen Pengguna Layanan BPJS," *J. Inform. Univ. Pamulang*, 2021, [Online]. Available: <https://media.neliti.com/media/publications/467010-none-7d5afb74.pdf>
- [14] F. N. Majid, "ANALISA SENTIMEN APLIKASI PEDULILINDUNGI DENGAN METODE NBC DAN SVM," *Elkom J. Elektron. dan Komput.*, 2023, [Online]. Available: <https://journal.stekom.ac.id/index.php/elkom/article/view/1000>
- [15] R. T. Handayanto, H. Herlawati, P. D. Atika, and ..., "Analisis Sentimen Pada Situs Google Review dengan Naïve Bayes dan Support Vector Machine," *J. Komtika ...*, 2021, [Online]. Available: <https://journal.unimma.ac.id/index.php/komtika/article/view/6280>
- [16] R. Slamet, W. Gata, A. Novtariany, and ..., "Analisis sentimen Twitter terhadap penggunaan artis Korea Selatan sebagai brand ambassador produk kecantikan lokal," *INTECOMS J. ...*, 2022, [Online]. Available: <https://journal.ipm2kpe.or.id/index.php/INTECOM/article/view/3933>