

RANCANG BANGUN APLIKASI ANALISIS SENTIMEN ULASAN APLIKASI DI GOOGLE PLAY STORE MENGGUNAKAN METODE *LOGISTIC REGRESSION* DAN *LEXICON-BASED APPROACH*

Fauzan Fadhillah Arisandi¹, Muhammad Rafi Muttaqin^{2*}

¹ Informatics Engineering, Sekolah Tinggi Teknologi Wastukencana, Purwakarta, Indonesia
Email: fauzanfadhillah55@wastukencana.ac.id

² Informatics Engineering, Sekolah Tinggi Teknologi Wastukencana, Purwakarta, Indonesia*
Email: rafi@wastukencana.ac.id

*Corresponding Author: rafi@wastukencana.ac.id

Submitted: dd-mm-yyyy, Revised: dd-mm-yyyy, Accepted: dd-mm-yyyy

ABSTRAK

Di era digital yang semakin maju, ulasan pengguna terhadap aplikasi seluler pada platform seperti Google Play Store telah menjadi sumber informasi penting bagi para pengembang dan pengguna aplikasi. Ulasan tersebut dapat mencerminkan kepuasan, keluhan, dan ekspektasi pengguna terhadap suatu aplikasi. Namun, karena semakin banyaknya ulasan, sangat sulit bagi pengembang untuk membaca dan mendistribusikan semua opini secara manual. Di sisi lain, ulasan yang tersedia seringkali tidak mencerminkan kondisi sebenarnya karena ulasan tersebut palsu, bias, atau manipulatif. Oleh karena itu, diperlukan sistem otomatis yang mampu mengelompokkan ulasan berdasarkan polaritas sentimen secara akurat dan efisien. Penelitian ini bertujuan untuk merancang dan membangun aplikasi berbasis web menggunakan *framework* Streamlit yang dapat melakukan analisis sentimen terhadap ulasan aplikasi di Google Play Store. Metode yang digunakan adalah gabungan antara *Logistic Regression* dan *Lexicon-Based Approach* dengan tahapan text preprocessing yakni *cleaning*, *case folding*, *slang fix*, *tokenizing*, *stopword removal* dan *stemming*, hingga transformasi data menggunakan TF-IDF. Penelitian ini menggunakan pendekatan CRISP-DM dalam membangun model dan sistem aplikasi. Hasil akhir dari penelitian ini adalah sebuah aplikasi analisis sentimen yang mampu mengklasifikasikan ulasan ke dalam kategori positif dan negatif, serta menampilkan hasil model evaluasi dalam bentuk visualisasi metrik dan *word cloud*. Berdasarkan hasil evaluasi kinerja model dengan *confusion matrix*, model *logistic regression* yang dibangun mencapai *accuracy* sebesar 91,2%, *precision* sebesar 91,73%, *recall* sebesar 98,46%, dan *F1-score* sebesar 64,78%.

Kata Kunci: Analisis Sentimen, Ulasan Aplikasi, *Logistic Regression*, *Lexicon-Based Approach*, Google Play Store.

ABSTRACT

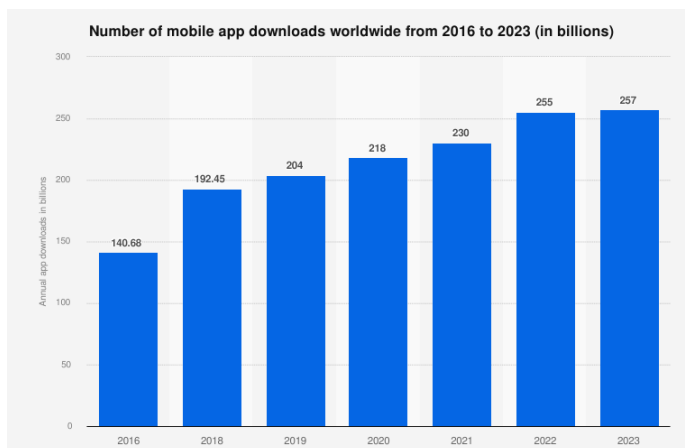
In the increasingly advanced digital era, user reviews of mobile applications on platforms such as the Google Play Store have become an important source of information for developers and application users. These reviews can reflect user satisfaction, complaints, and expectations of an application. However, due to the increasing number of reviews, it is very difficult for developers to read and distribute all opinions manually. On the other hand, the available reviews often do not reflect the actual conditions because the reviews are fake, biased, or manipulative. Therefore, an automated system is needed that is able to group reviews based on sentiment polarity accurately and efficiently. This study aims to design and build a web-based application using the Streamlit framework that can perform sentiment analysis on application reviews on the Google Play Store. The method used is a combination of Logistic Regression and Lexicon-Based Approach with text preprocessing stages, namely

cleaning, case folding, slang fix, tokenizing, stopwords removal and stemming, to data transformation using TF-IDF. This study uses the CRISP-DM approach in building application models and systems. The final result of this study is a sentiment analysis application that is able to classify reviews into positive and negative categories, and display the results of the evaluation model in the form of metric visualization and word cloud. Based on the results of the model performance evaluation with confusion matrix, the logistic regression model built achieved an accuracy of 91.2%, a precision of 91.73%, a recall of 98.46%, and an F1-score of 64.78%.

Keywords: Sentiment Analysis, App Reviews, Logistic Regression, Lexicon-Based Approach, Google Play Store.

1. PENDAHULUAN

Aplikasi seluler telah menjadi bagian penting dari kehidupan masyarakat di era teknologi modern. Aplikasi yang tersedia di perangkat seluler semakin bergantung pada aktivitas sehari-hari seperti komunikasi, belanja, hiburan, dan pekerjaan (Oktaviane & Herwanto, 2024). Sebagai salah satu platform distribusi digital terbesar, Google Play Store menawarkan berbagai aplikasi dengan fitur penilaian dan ulasan pengguna. Fitur ini memiliki manfaat banyak bagi pengembang. Ini membantu pengembang meningkatkan kualitas layanan dengan memberikan umpan balik kepada pengguna (Cahyaningtyas dkk., 2021).



Gambar 1. Jumlah aplikasi mobile yang diunduh di seluruh dunia

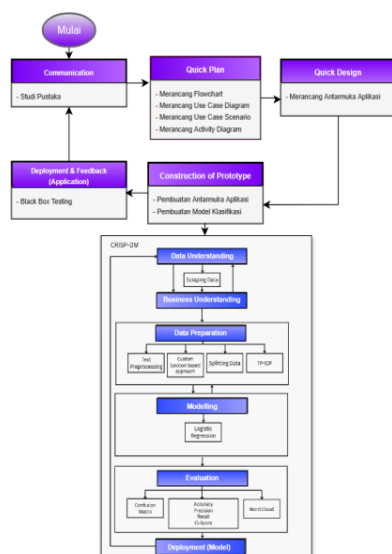
Gambar 1 adalah grafik yang disusun oleh Statista Research Department ini menunjukkan peningkatan besar dalam jumlah unduhan aplikasi hingga 2023, meningkat dari 140,68 miliar pada 2016 menjadi 257 miliar pada 2023, menunjukkan ketergantungan masyarakat global terhadap aplikasi digital. Jumlah ulasan pengguna terus meningkat seiring dengan jumlah unduhan. Setiap pengguna baru yang mengunduh aplikasi memiliki kesempatan untuk memberikan ulasan, ulasan ini kemudian menjadi komponen penting dalam ekosistem evaluasi aplikasi.

Namun, banyaknya ulasan merupakan masalah besar. Hal ini menimbulkan masalah serius bagi pengembang yang kesulitan menganalisis ulasan satu per satu secara manual untuk memahami keluhan atau masukan pengguna. Selain itu, dengan adanya ulasan palsu (*fake reviews*) semakin memperumit situasi. Penelitian (Aralikatte, 2018, sebagaimana dikutip dalam Sadiq dkk., 2021) menunjukkan bahwa 20% ulasan mengandung ketidaksesuaian antara rating dan isi teks, yang berpotensi menyesatkan pengguna dan merusak reputasi pengembang (Sadiq dkk., 2021). Ketergantungan pada skor *rating* saja juga tidak cukup akurat karena rentan terhadap manipulasi dan bias subjektif.

Penelitian sebelumnya umumnya hanya berfokus pada pengujian model analisis sentimen tanpa menyediakan antarmuka aplikasi yang mudah digunakan. (Saputra dkk., 2022) mengevaluasi performa model *Naïve Bayes* pada ulasan aplikasi Ajaib, namun tidak mengembangkan aplikasi yang dapat diakses oleh pengguna umum. Hal serupa juga terlihat pada penelitian (Basri dkk., 2024) yang hanya membandingkan algoritma tanpa merancang antarmuka yang *user-friendly*. Hasil analisis sentimen memiliki potensi besar untuk dimanfaatkan secara langsung oleh pengembang aplikasi, pemilik produk, atau tim pemasaran dalam memahami opini dan kepuasan pengguna. Untuk mengatasi masalah ini, penelitian ini mengusulkan penerapan analisis sentimen berbasis *lexicon* dan ekstraksi frekuensi kata kunci. Analisis sentimen merupakan proses otomatis untuk mengelompokkan teks berdasarkan polaritas positif, negatif, atau netral guna memahami opini pengguna terhadap produk, layanan, atau isu tertentu (Sujadi dkk., 2022). Pendekatan ini memungkinkan klasifikasi otomatis ulasan menjadi positif, atau negatif menggunakan model Logistic Regression, sekaligus mengidentifikasi kata kunci dominan yang sering muncul dalam ulasan. Dengan demikian, pengembang dapat lebih mudah memahami fitur apa yang paling banyak dibahas atau dikeluhkan oleh pengguna. Hasil analisis dengan menggunakan *custom lexicon* juga akan dapat membantu mendeteksi ulasan palsu dengan membandingkan polaritas teks dan rating, serta melihat konsistensi opini. Keunggulan metode ini terletak pada keakuratannya yang tidak hanya bergantung pada *rating*, efisiensinya dalam memproses banyak ulasan dalam waktu singkat, serta kemudahan penggunaannya melalui antarmuka berbasis Streamlit yang *user-friendly*.

Manfaat dari solusi ini bersifat dua arah. Bagi pengguna, hasil analisis sentimen dapat menjadi panduan untuk memperbaiki aplikasi berdasarkan masukan pengguna yang terstruktur, serta memperhatikan jumlah *thumbsupcount*, sehingga memudahkan mereka dalam menilai aplikasi secara lebih objektif, mengurangi risiko tertipu oleh ulasan palsu. Sementara bagi pengembang, visualisasi polaritas sentimen dan frekuensi kata kunci memudahkan mereka dalam menilai aplikasi secara lebih objektif, sehingga dapat tepat sasaran dalam pengembangan aplikasi dan meningkatkan kepuasan pengguna. Aplikasi hasil penelitian ini dapat langsung dimanfaatkan oleh pengembang dan pengguna umum berkat antarmukanya yang sederhana. Dengan demikian, diharapkan ekosistem aplikasi digital menjadi informatif, mendorong perbaikan kualitas aplikasi serta pengambilan keputusan yang lebih bijak dari pengguna.

2. METODOLOGI PENELITIAN



Gambar 2. Alur Penelitian

Dalam penelitian ini digunakan dua pendekatan metodologi yaitu *Prototype Model* untuk pengembangan aplikasi dan CRISP-DM sebagai *framework* utama dalam proses analisis data dan pengembangan model analisis sentimen. Berikut penjelasan masing-masing tahapan dalam metode penelitian:

Communication

Studi pustaka ini dilakukan oleh penulis dengan pendekatan analisis pustaka terhadap berbagai sumber yang relevan dengan topik, meliputi buku referensi, jurnal ilmiah, tesis terdahulu, dan situs terkait. Fokus utama studi pustaka ini adalah untuk memperoleh pemahaman teoritis tentang konsep dasar yang mendukung penelitian, seperti *Natural Language Processing (NLP)*, *Text Preprocessing*, *Machine Learning*, dan *algoritma Logistic Regression*.

Quick Plan

Pada tahap *Quick Plan* merupakan tahap perancangan sistem awal yang dilakukan berdasarkan hasil komunikasi dan studi literatur. Pada tahap ini dilakukan beberapa kegiatan penting, antara lain merancang *flowchart* untuk menggambarkan alur sistem secara keseluruhan, merancang *use case diagram* untuk memodelkan interaksi antara pengguna dengan sistem, dan menyusun *use case scenario* yang menjelaskan skenario penggunaan sistem secara lebih rinci. Selain itu, dibuat juga *activity diagram* untuk menggambarkan secara visual alur aktivitas dalam sistem. Semua perancangan ini bertujuan untuk memberikan gambaran awal mengenai sistem yang akan dikembangkan dan menjadi acuan pada tahap perancangan dan implementasi selanjutnya.

Quick Design

Tahap ini berfokus pada perancangan awal antarmuka pengguna aplikasi. Perancangan antarmuka dibuat berdasarkan hasil rancangan cepat dan ditujukan untuk menggambarkan bentuk visual aplikasi yang akan dibangun. Pada tahap ini dibuat sketsa tampilan aplikasi untuk memperjelas bagaimana pengguna akan berinteraksi dengan sistem dan bagaimana hasil analisis sentimen akan ditampilkan.

Construction of Prototype

Tahap *Construction of Prototype* merupakan tahap implementasi awal sistem berdasarkan rancangan sebelumnya. Kegiatan utama pada tahap ini adalah pengembangan antarmuka aplikasi dan pembuatan model klasifikasi. Pada tahap ini juga dilakukan implementasi proses analisis sentimen menggunakan pendekatan CRISP-DM.

CRISP-DM

Penelitian ini menggunakan metode CRISP-DM sebagai *framework* utama dalam mengembangkan aplikasi analisis sentimen dan yang biasanya digunakan dalam proyek data mining (Elkabalawy dkk., 2024). Metode ini terdiri dari enam tahap yang saling terkait dan dilakukan secara berurutan yaitu *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling* menggunakan algoritma *Logistic Regression*, *Evaluation*, dan *Deployment*. Pada tahap *Data Understanding* dilakukan proses pengumpulan data dengan menggunakan teknik *web scraping* pada user review dari Google Play Store, dengan total 5000 review berhasil dikumpulkan dan dikonversi menjadi csv (*Comma Separated Value*). Selanjutnya pada tahap *Data Preparation* dilakukan beberapa proses penting seperti *text preprocessing* yaitu *cleaning*, *case folding*, *tokenizing*, *stopword removal*, dan *stemming*, penerapan pendekatan custom lexicon based untuk memperkuat identifikasi sentimen, pembagian data menjadi *data training* dan *data test* dengan perbandingan 80:20, dan transformasi data ke dalam bentuk numerik menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*). Pada tahap Evaluasi, kinerja model dianalisis menggunakan beberapa metrik evaluasi, yaitu *accuracy*, *precision*, *recall* dan *f1-score* untuk mengukur

sejauh mana model mampu mengklasifikasikan sentimen secara akurat dan seimbang, juga menggunakan *confusion matrix* dan *wordcloud*.

Deployment & Feedback

Tahap ini meliputi penerapan prototipe aplikasi ke lingkungan pengujian dan pengujian fungsionalitas sistem. Pengujian dilakukan menggunakan metode pengujian *black-box* untuk memverifikasi bahwa semua fitur berjalan sesuai kebutuhan. Selain itu, masukan dikumpulkan dari pengguna atau penguji untuk mengidentifikasi kekurangan sistem dan rencana untuk iterasi perbaikan berikutnya.

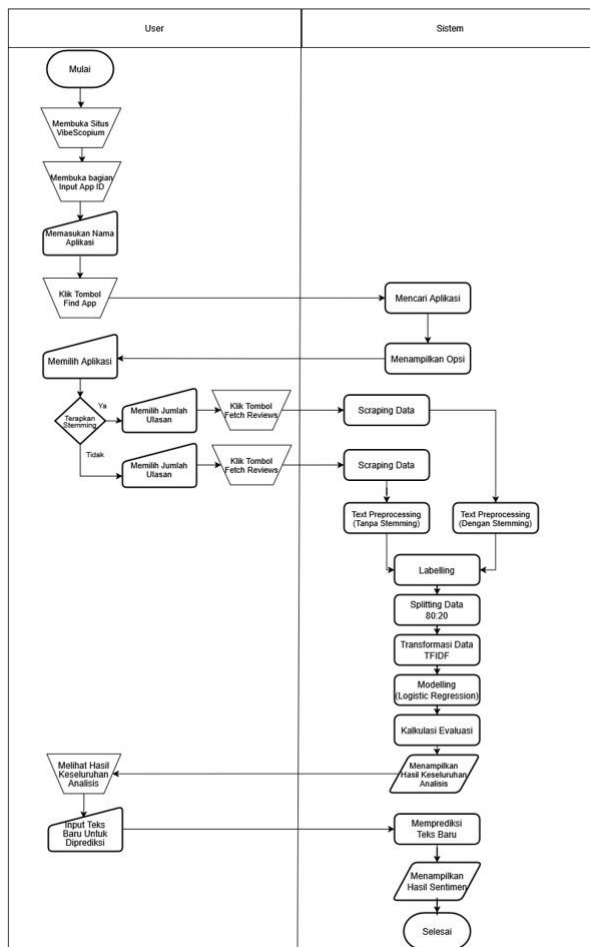
3. HASIL DAN PEMBAHASAN

Communication

Tahap *communication* dalam pengembangan aplikasi ini dilakukan melalui studi pustaka untuk memperoleh pemahaman mengenai konsep-konsep dasar yang mendasari sistem analisis sentimen, serta metode dan pendekatan yang digunakan. Studi pustaka mencakup penelusuran terhadap literatur terkait analisis sentimen, *machine learning*, pendekatan *lexicon-based*, serta algoritma klasifikasi seperti *Logistic Regression*. Selain itu, juga dilakukan penelaahan terhadap proses *text preprocessing*, transformasi fitur menggunakan TF-IDF, dan teknik evaluasi performa model.

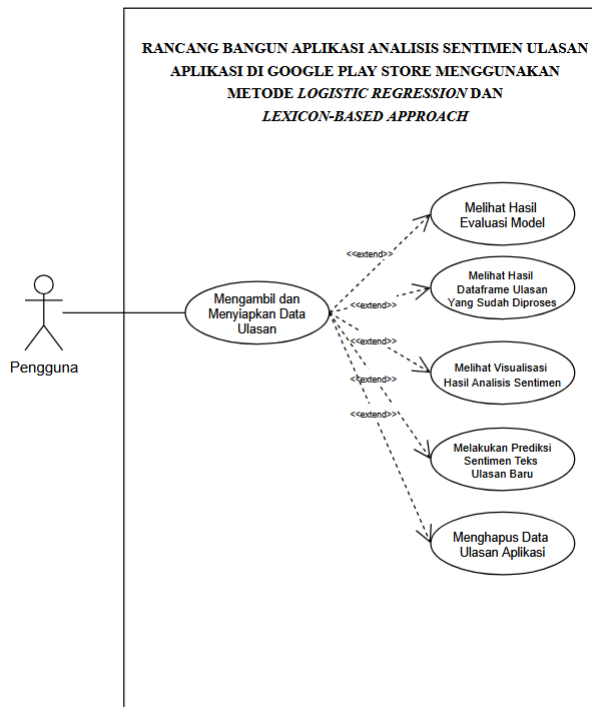
Quick Plan

Berikut *Flowchart* Sistem dapat dilihat pada Gambar 3.



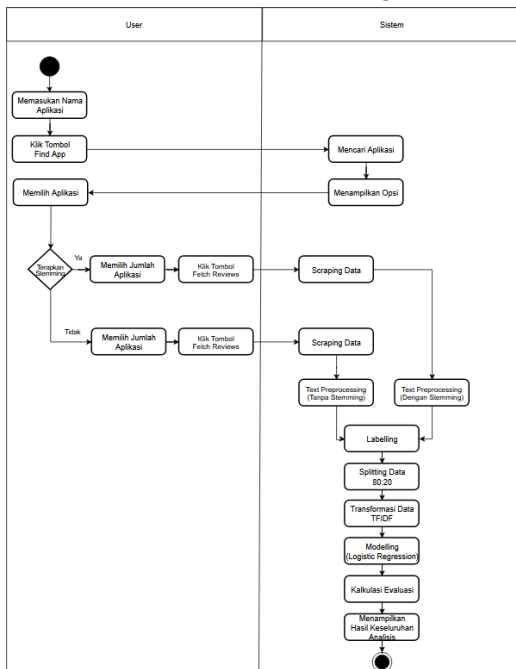
Gambar 3. *Flowchart Sistem*

Berikut *Use Case Diagram* dapat dilihat pada Gambar 4.



Gambar 4. *Use Case Diagram*

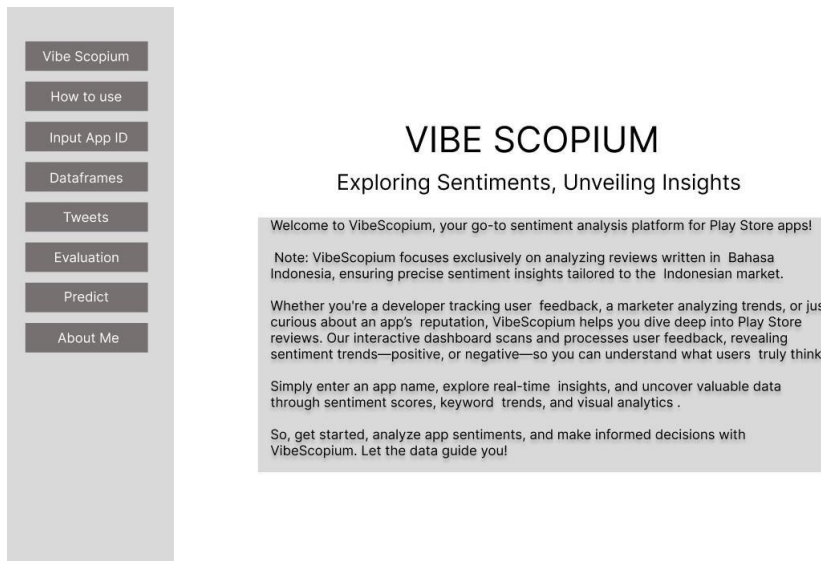
Berikut *Activity Diagram* mengambil dan menyiapkan ulasan dapat dilihat pada Gambar 5.



Gambar 5. *Activity Diagram* Mengambil dan Menyiapkan Ulasan

Quick Design

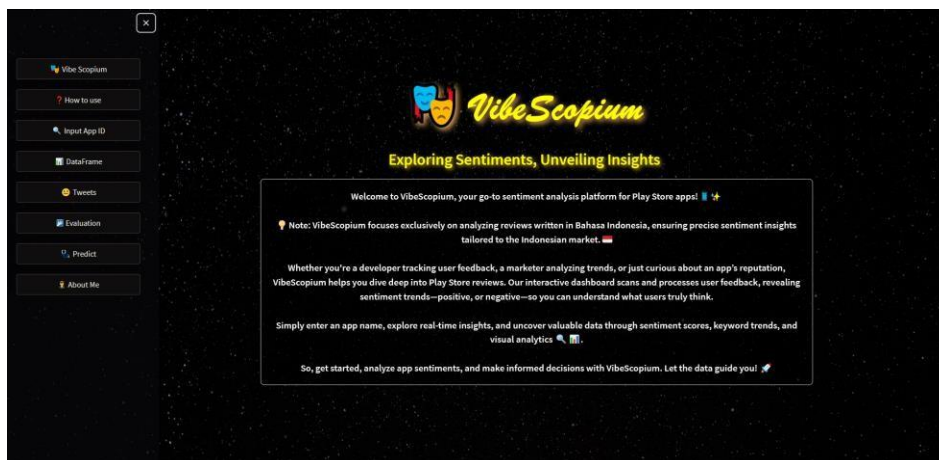
Berikut *Wireframe* halaman utama dapat dilihat pada Gambar 6.



Gambar 6. *Wireframe* Halaman Utama

Construction of Prototype

Berikut antarmuka halaman utama dapat dilihat pada Gambar 7.



Gambar 7. Antarmuka Halaman Utama

CRISP-DM

Hasil dan pembahasan selanjutnya pada penelitian ini akan dibahas menggunakan metode *Cross Industry Standard Process for Data Mining* (CRISP-DM). Metode CRISP-DM terdiri dari enam tahap utama yang membentuk kerangka kerja sistematis dalam proses data mining, yakni *Business Understanding* untuk menentukan tujuan dan kebutuhan bisnis, *Data Understanding* untuk mengumpulkan serta memahami karakteristik data awal, *Data Preparation* yang mencakup pembersihan dan transformasi data agar siap untuk pemodelan, *Modeling* dengan menerapkan algoritma yang sesuai guna membangun model prediktif, *Evaluation* untuk menilai kinerja model dan kesesuaiannya dengan tujuan bisnis, serta *Deployment* sebagai tahap penerapan model ke sistem (Lestari dkk., 2024).

Business Understanding

Studi kasus dalam penelitian ini difokuskan pada aplikasi **DANA**, sebuah dompet digital populer di Indonesia. Karena DANA memiliki jumlah ulasan yang banyak serta relevan dalam konteks analisis opini pengguna terhadap layanan dompet digital.

Data Understanding

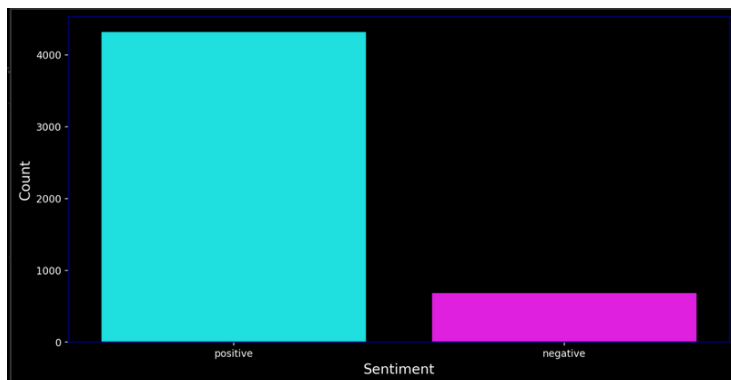
Teknik *scraping* dilakukan menggunakan *library* Python *google-play-scraper* untuk mengambil data ulasan secara terstruktur berdasarkan *Application ID*. Peneliti mengembangkan antarmuka input guna mempermudah pencarian aplikasi dan memicu proses *scraping*, yang secara otomatis mengunduh 5000 ulasan dan menyimpannya dalam format CSV. Selanjutnya, dilakukan eksplorasi awal terhadap data yang telah dikumpulkan, yang mencakup berbagai kolom atribut ulasan dari Google Play Store.

Tabel 1. Seluruh Kolom Atribut Data Ulasan Google Play Store

No	Kolom Atribut	Keterangan
1	<i>reviewId</i>	ID unik dari ulasan
2	<i>userName</i>	Nama pengguna yang menulis ulasan
3	<i>userImage</i>	URL ke gambar profil pengguna
4	<i>content</i>	Isi teks ulasan
5	<i>Score</i>	Rating yang diberikan pengguna
6	<i>thumbsUpCount</i>	Jumlah pengguna yang menganggap ulasan tersebut bermanfaat
7	<i>reviewCreatedVersion</i>	Versi aplikasi saat ulasan dibuat
8	<i>at</i>	Tanggal dan waktu saat ulasan diposting
9	<i>replyContent</i>	Isi balasan dari pengembang
10	<i>repliedAt</i>	Tanggal dan waktu saat pengembang memberikan balasan
11	<i>appVersion</i>	Versi aplikasi

Data Preparation

Text preprocessing merupakan tahap penting dalam analisis teks yang bertujuan untuk membersihkan dan menyusun ulang data agar lebih siap untuk diproses oleh model *machine learning* (Hadiprakoso dkk., 2021). *Text preprocessing* ini mencakup *cleansing*, *casefolding*, *slang fix*, *tokenizing*, *stopword removal*, dan *stemming*. Kemudian dilakukannya proses penerapan *Custom Lexicon* untuk analisis sentimen dari data tersebut.



Gambar 8. Distribusi Polaritas Sentimen

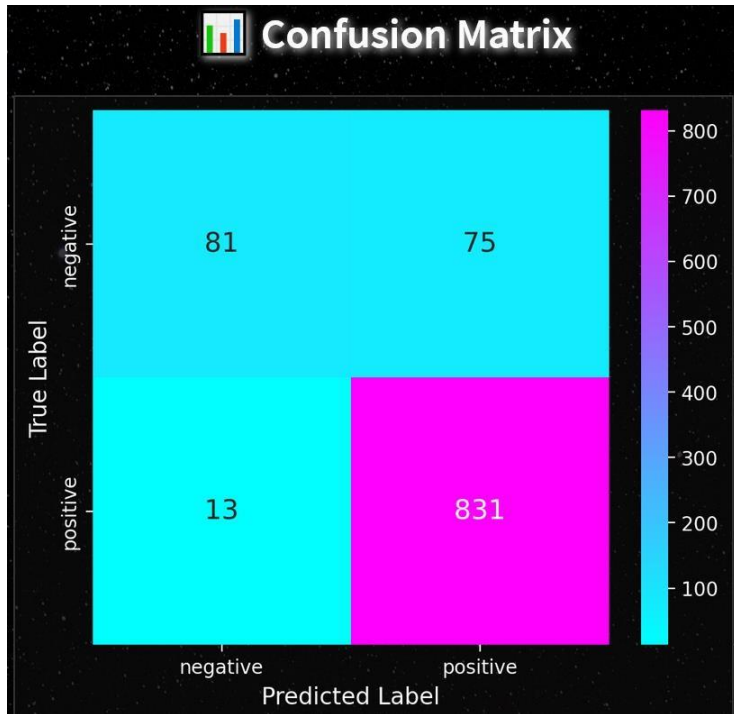
Gambar 8 menunjukkan distribusi polaritas sentimen dari hasil analisis terhadap 5000 ulasan aplikasi yang telah diklasifikasikan menggunakan pendekatan *Custom Lexicon-Based*. Dari total ulasan tersebut, sebanyak 4322 ulasan teridentifikasi mengandung sentimen positif, sementara 678 ulasan tergolong sentimen negatif. Dapat disimpulkan bahwa distribusi ini menunjukkan dominasi opini positif dari pengguna terhadap aplikasi yang dianalisis, yang dapat mengindikasikan tingkat kepuasan yang tinggi.

Modelling

Modelling merupakan tahap penting dalam proses analisis sentimen, di mana algoritma *machine learning* digunakan untuk membangun model prediksi berdasarkan data latih yang telah dipersiapkan sebelumnya. Dalam penelitian ini, algoritma yang digunakan adalah **Logistic Regression**. Implementasi *Logistic Regression* dilakukan menggunakan library *scikit-learn*.

Evaluation

Dari 1000 data uji, menghasilkan 831 TP, 81 TN, 75 FP dan 13 FN. Berikut *Confusion Matrix* dari pelatihan model dapat dilihat pada Gambar 9.



Gambar 9. *Confusion Matrix*

Berikut *Wordcloud* dari ulasan aplikasi Dana dapat dilihat pada Gambar 10.

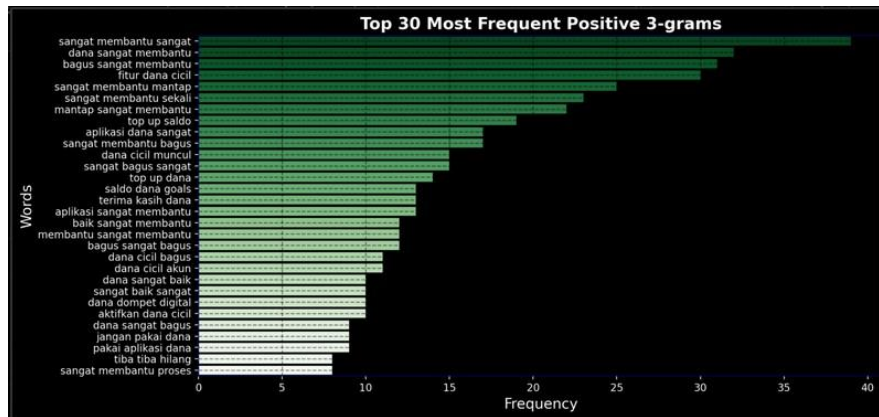


Gambar 10. *WordCloud*

Gambar 10 menampilkan visualisasi kata-kata yang paling sering muncul dalam ulasan pengguna terhadap aplikasi DANA. Kata-kata seperti "dana", "bagus", "sangat", "membantu",

muncul dalam ukuran besar, menunjukkan bahwa kata-kata tersebut paling sering disebutkan dan memiliki frekuensi tinggi.

Berikut grafik frekuensi kata kunci dapat dilihat pada Gambar 11.



Gambar 11. Frekuensi Kata Kunci Positif

Gambar 11 menunjukkan bahwa frasa seperti "dana sangat membantu", "aplikasi sangat bagus", dan "saldo dana goals" mencerminkan kepuasan pengguna terhadap aplikasi Dana. Pengguna merasa aplikasi ini memudahkan aktivitas finansial mereka, memberikan pengalaman yang baik secara keseluruhan, serta menyediakan fitur spesifik seperti "Dana Goals" yang membantu dalam perencanaan keuangan.

Tabel 2. Evaluasi Metrik

	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-score</i>
Accuracy	$\frac{831 + 81}{831 + 81 + 75 + 13} = 91.2\%$	-	-	-
Negatif	-	$\frac{81}{81 + 13} = 86.17\%$	$\frac{81}{81 + 75} = 51.92\%$	$2 \times \frac{0.9173 \times 0.9846}{0.9173 + 0.9846} = 94.9\%$
Positif	-	$\frac{831}{831 + 75} = 91.73\%$	$\frac{831}{831 + 13} = 98.46\%$	$2 \times \frac{0.8617 \times 0.5192}{0.8617 + 0.5192} = 64.78\%$

Dari Tabel 2. menunjukkan bahwa model memiliki *accuracy* sebesar 91.2%, yang berarti mayoritas prediksi model sudah tepat. Untuk kelas positif, model menunjukkan performa sangat baik dengan *precision* 91.73% dan *recall* 98.46%, menghasilkan *F1-score* sebesar 94.99%, menandakan keseimbangan yang kuat antara ketepatan dan kemampuan mendeteksi ulasan positif. Namun, pada kelas negatif, *precision* nya cukup baik sebesar 86.17%, tetapi *recall*-nya rendah di angka 51.92%, yang menyebabkan *F1-score* turun menjadi 64.78%. Hal ini mengindikasikan bahwa model lebih andal dalam mengenali ulasan positif dibandingkan negatif.

Deployment

Setelah model *Logistic Regression* berhasil dibangun dan dievaluasi, tahap selanjutnya adalah implementasi agar model dapat digunakan secara interaktif oleh pengguna. Proses deploy dilakukan menggunakan Streamlit yang dihosting melalui Streamlit Cloud, dengan source code terlebih dahulu diunggah ke GitHub. File aplikasi utama (.py) berisi antarmuka pengguna dan logika backend, disertai dengan file pendukung seperti *requirements.txt* untuk

mendefinisikan dependensi. Repositori GitHub ini kemudian dihubungkan ke Streamlit Cloud, dan sistem akan

secara otomatis membangun dan menjalankan aplikasi yang akhirnya dapat diakses secara publik melalui URL yang tersedia.

Setelah aplikasi berhasil *dideploy*, dilakukan pengujian untuk memastikan semua fitur berjalan sesuai desain. Pengujian dilakukan dengan metode *black-box testing* yang menilai fungsi aplikasi berdasarkan *output* dan *input* yang diberikan, tanpa melihat struktur internal program. Aplikasi yang diberi nama VibeScopium ini tersedia sebagai platform berbasis web yang dapat diakses melalui browser, dan hasil pengujian berfungsi sebagai umpan balik untuk memastikan sistem berjalan optimal dari sudut pandang pengguna. Pengujian menggunakan *black-box* dapat dilihat pada Tabel 3.

Tabel 3. *Black-Box Testing*

No	Kondisi	Hasil yang Diharapkan	Hasil yang Didapatkan	Sukses/Gagal
1	Pengguna memasukkan nama aplikasi dan klik "Find App"	Aplikasi menampilkan 5 hasil pencarian yang relevan	Aplikasi menampilkan 5 hasil sesuai dengan kata kunci	Sukses
2	Pengguna tidak mengisi kolom input lalu klik "Find App"	Muncul peringatan bahwa input tidak boleh kosong	Muncul notifikasi "Silakan isi nama aplikasi terlebih dahulu."	Sukses
3	Pengguna memilih aplikasi dan klik "Fetch Reviews" dengan jumlah 1000	Sistem melakukan <i>scraping</i> dan menyiapkan data, lalu menampilkan notifikasi berhasil	Proses <i>scraping</i> berjalan dan data siap ditampilkan	Sukses
4	Pengguna membuka halaman <i>Dataframe</i>	Data ulasan yang telah diproses ditampilkan dalam tabel	Data tampil sesuai <i>preprocessing</i>	Sukses
5	Pengguna mengunduh <i>Dataframe</i>	Sistem memberikan <i>Dataframe</i> pada pengguna untuk diunduh	Pengguna dapat mengunduh <i>Dataframe</i>	Sukses
6	Pengguna membuka halaman <i>Evaluation</i>	Sistem menampilkan metrik <i>accuracy</i> , <i>precision</i> , <i>recall</i> , <i>F1-score</i> , dan <i>confusion matrix</i>	Semua metrik evaluasi ditampilkan lengkap	Sukses
7	Pengguna membuka halaman <i>Tweets</i>	Ditampilkan <i>word cloud</i> dan frekuensi kata kunci sesuai kategori sentimen	Visualisasi tampil dengan jelas dan sesuai kategori sentimen	Sukses
8	Pengguna mengisi kolom prediksi dengan teks ulasan baru	Sistem mengeluarkan hasil sentimen dari input teks tersebut	Sentimen tampil (positif/negatif)	Sukses
9	Pengguna mengosongkan kolom prediksi dan klik tombol prediksi	Muncul notifikasi error bahwa input tidak boleh kosong	Notifikasi muncul dan prediksi dibatalkan	Sukses
10	Pengguna klik "Reset Data" setelah proses selesai	Data <i>scraping</i> dan hasil olahan dihapus dari tampilan dan memori aplikasi	Data berhasil dihapus, tampilan dikosongkan	Sukses

4. KESIMPULAN DAN SARAN

Kesimpulan dari penelitian ini mencakup beberapa hal penting. Pertama, model *Logistic Regression* yang digunakan menunjukkan kinerja yang cukup baik dengan *accuracy* sebesar 91,20%. Performa model menunjukkan optimal dalam mengklasifikasikan sentimen positif dengan *precision* sebesar 91,73%, *recall* sebesar 98,46%, dan *F1-score* sebesar 94,90%, namun masih kurang konsisten dalam mendeteksi sentimen negatif dengan *precision* sebesar 86,17%, *recall* sebesar 51,92%, dan *F1-score* sebesar 64,78%. Kedua, hasil analisis *trigram* mengungkap pola-pola yang sering muncul, yakni "dana sangat membantu" untuk sentimen positif dan "saldo hilang sendiri" untuk sentimen negatif. Temuan ini menyoroti bahwa pengguna merasa puas terhadap kemudahan dan fitur aplikasi, namun juga menunjukkan adanya keluhan serius terkait teknis dan layanan yang perlu segera ditindaklanjuti oleh pengembang. Ketiga,

aplikasi VibeScopium yang dikembangkan mampu melakukan analisis sentimen secara otomatis, mulai dari proses *scraping*, *preprocessing*, pelabelan sentimen, transformasi TF-IDF,

hingga klasifikasi menggunakan model *Logistic Regression*. Aplikasi ini diharapkan dapat menjadi alat bantu yang efektif bagi pengembang, peneliti, dan pelaku industri dalam memahami opini pengguna dan mendukung pengambilan keputusan berbasis data.

Adapun beberapa saran untuk penelitian selanjutnya adalah mencoba berbagai algoritma klasifikasi dan metode transformasi data lainnya untuk menemukan model yang lebih optimal. Selain itu, disarankan untuk menambah jumlah data ulasan sehingga model dapat mengenali pola sentimen yang lebih bervariasi dan kompleks. Penelitian selanjutnya juga dapat difokuskan pada pengembangan metode deteksi ulasan palsu yang lebih akurat lagi, seperti dengan memanfaatkan pendekatan berbasis *machine learning* yang mampu menganalisis fitur linguistik yang mencurigakan secara lebih mendalam.

UCAPAN TERIMA KASIH

Terima kasih kepada Ketua dan para dosen STT Wastukencana Purwakarta serta LPPM Yarsi Pratama University yang telah banyak membantu menyelesaikan penelitian ini hingga dapat menjadi sebuah karya ilmiah yang diterbitkan.

DAFTAR PUSTAKA

- Basri, H., Junianto, M. B. S., & Kusyadi, I. (2024). Enhancing Usability Testing Through Sentiment Analysis: A Comparative Study Using SVM, Naive Bayes, Decision Trees and Random Forest. *Jurnal Teknologi Sistem Informasi dan Aplikasi*, 7(4), 1603-1610. <https://doi.org/10.32493/jtsi.v7i4.45117>
- Cahyaningtyas, C., Nataliani, Y., & Widiyarsari, I. R. (2021). Analisis sentimen pada rating aplikasi Shopee menggunakan metode Decision Tree berbasis SMOTE. *AITI: Jurnal Teknologi Informasi*, 18(Agustus), 173-184.
- Elkabalawy, M., Al-Sakkaf, A., Mohammed Abdelkader, E., & Alfalah, G. (2024). CRISP-DM-Based Data-Driven Approach for Building Energy Prediction Utilizing Indoor and Environmental Factors. *Sustainability*, 16(17), 7249. <https://doi.org/10.3390/su16177249>
- Hadiprakoso, R. B., Setiawan, H., Yasa, R. N., & Girinoto, G. (2021). Text Preprocessing for Optimal Accuracy in Indonesian Sentiment Analysis Using a Deep Learning Model with Word Embedding. *AIP Conference Proceedings*. <https://www.researchgate.net/publication/365799401>
- Lestari, R. I., Andrawina, L., & Mufidah, I. (2024). Optimization of throughput rate prediction in animal feed industry using crisp-dm and operational research approaches. *IJIO*, 6(1), 87-102. <https://doi.org/10.12928/ijio.v6i1.11357>
- Oktaviane, S. P., & Herwanto, P. (2024). Dampak Penggunaan Perangkat Mobile dalam Mendukung Kegiatan Pembelajaran Mandiri Siswa Kelas IX di SMP PGRI Rancaekek. *Uranus : Jurnal Ilmiah Teknik Elektro, Sains dan Informatika*, 2(4), 165-174. <https://doi.org/10.61132/uranus.v2i4.491>
- Sadiq, S., Umer, M., Ullah, S., Mirjalili, S., Rupapara, V., & Nappi, M. (2021). Discrepancy detection between actual user reviews and numeric ratings of Google App store using deep learning. *Expert Systems with Applications*, 181. <https://doi.org/10.1016/j.eswa.2021.115111>
- Saputra, S. A., Rahmatullah, B., & Budiyo, P. (2022). Sentiment Analysis User Ajaib Application Using Naïve Bayes Algorithm. *Journal of Information System, Informatics and Computing*, 6(2), 497-505. <https://doi.org/10.52362/jisicom.v6i2.964>
- Sujadi, H., Fajar, S., & Roni, C. (2022). ANALISIS SENTIMEN PENGGUNA MEDIA SOSIAL TWITTER TERHADAP WABAH COVID-19 DENGAN METODE NAIVE BAYES CLASSIFIER DAN SUPPORT VECTOR MACHINE. *INFOTECH journal*, 8(1), 22-27. <https://doi.org/10.31949/infotech.v8i1.1883>