

BAB 3

METODE PENELITIAN

3.1 Bahan Penelitian

Data penelitian yang menjadi objek dalam studi ini terdiri dari ulasan pengguna aplikasi KAI Access yang diperoleh melalui pengunduhan dari platform Google Play Store. Data tersebut dikumpulkan untuk dianalisis menggunakan pendekatan *Natural Language Processing* (NLP) dalam menentukan sentimen pengguna terhadap aplikasi. Setiap ulasan yang diperoleh berisi informasi teks yang merepresentasikan pengalaman, kepuasan, maupun keluhan pengguna terhadap layanan KAI Access. Data ini menjadi bahan utama dalam proses analisis sentimen, di mana setiap ulasan akan diproses lebih lanjut melalui tahapan pra-pemrosesan teks, ekstraksi fitur, serta klasifikasi sentimen. Pemilihan data ulasan dari Google Play Store didasarkan pada pertimbangan bahwa platform tersebut merupakan salah satu sumber utama umpan balik konsumen terhadap aplikasi berbasis mobile.

3.2 Alat Penelitian

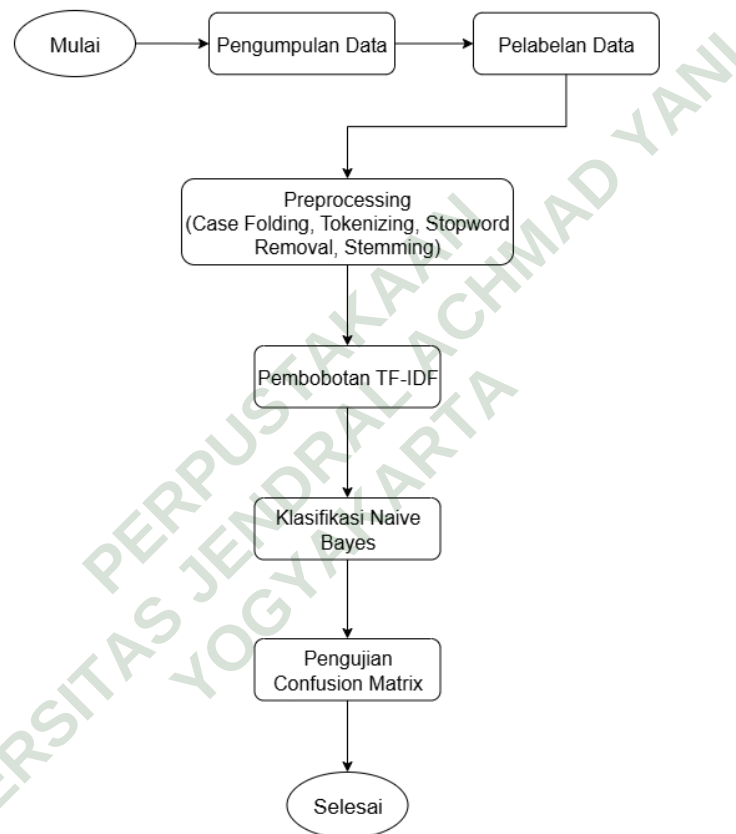
Perangkat yang dipergunakan dalam penelitian ini adalah komputer dengan spesifikasi yang cukup untuk mendukung operasional sistem operasi, perangkat lunak pengembangan, serta akses koneksi internet. Sistem operasi dan aplikasi yang digunakan dalam proses pengembangan penelitian ini mencakup:

1. Sistem Operasi: Windows 11
2. Perangkat Keras: Laptop Acer
3. Bahasa Pemrograman : Python
4. Text Editor : Google Collab
5. Pengambilan Data: Google Play Store
6. Browser : Google Chrome

3.3 Jalan Penelitian

3.3.1 Perancangan Sistem

Untuk mempermudah desain sistem, penelitian ini menggunakan diagram alir (*Flowcart*) sebagai representasi visual dari proses yang dikembangkan. Perancangan sistem ini melibatkan serangkaian langkah terstruktur.



Gambar 3.1 *Flowchart* Penelitian

Sistem dalam penelitian dirancang dengan metode:

1. Pengumpulan Data

Proses pengumpulan data berupa ulasan dari pengguna aplikasi KAI Access pada platform Google Play Store dilakukan melalui proses scraping pada tahapan ini. *Web scraping* merupakan teknik otomatis untuk mengekstrak data atau informasi dari dokumen semi-terstruktur di internet, seperti halaman website dengan format HTML atau XHTML, yang kemudian digunakan untuk tujuan tertentu (Kusumo &

Somya, 2022). Teknik *web scraping* memungkinkan perolehan data yang cepat, akurat, dan otomatis dari situs web yang menjadi target (Djiwadikusumah dkk., 2021) .

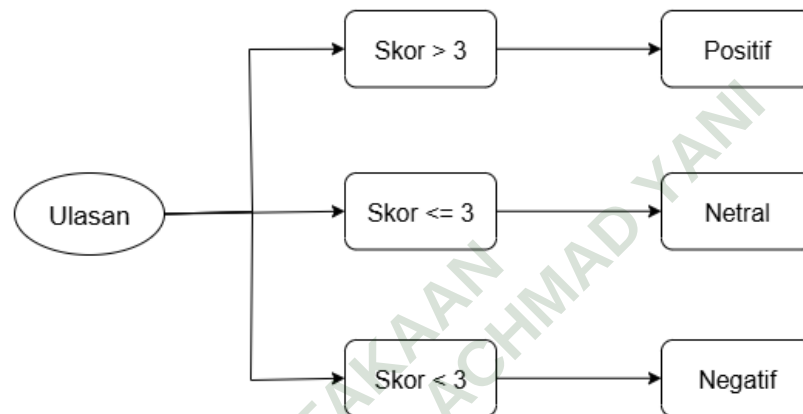
Pada Google Play Store, tersedia fitur rating dan ulasan yang digunakan untuk merepresentasikan performa sebuah aplikasi. Skor 1 bintang menandakan performa yang sangat buruk, 2 bintang menunjukkan performa yang kurang baik, 3 bintang menunjukkan performa yang cukup, 4 bintang menandakan performa yang baik, dan 5 bintang menunjukkan performa yang sangat baik. Pada penelitian ini, ulasan yang dipilih adalah ulasan yang paling relevan dengan *rating* berbintang 1, 2, 4, dan 5. Pemilihan kategori ulasan relevan ini dilakukan karena ulasan tersebut berisi tanggapan yang secara langsung berkaitan dengan aplikasi, dengan isi yang jelas menggambarkan keluhan, kepuasan pengguna, maupun saran untuk pengembangan aplikasi di masa mendatang

2. Pelabelan Data

Dalam tahap ini, setiap tanggapan dari pengguna diklasifikasikan ke dalam kelompok sentimen positif, netral, maupun negatif berdasarkan nilai skornya. Ulasan positif berisi komentar yang mendukung serta merekomendasikan penggunaan aplikasi, sementara ulasan negatif memuat kritik atau saran yang ditujukan untuk aplikasi tersebut (Musfiroh dkk., 2024). Ulasan dengan skor kurang dari 3 diberi tanda sebagai sentimen negatif, skor 4 dan 5 diklasifikasikan ke dalam kategori sentimen positif, sedangkan skor 3 dikategorikan sebagai ulasan netral. Proses pelabelan ini bertujuan untuk mengelompokkan ulasan secara otomatis agar dapat digunakan dalam tahap pelatihan (*training*) dan pengujian (*testing*) model. Selain itu, pelabelan secara manual dilakukan untuk menentukan klasifikasi sentimen pada masing-masing ulasan menjadi positif, negatif, atau netral.

Tahapan berikutnya meliputi:

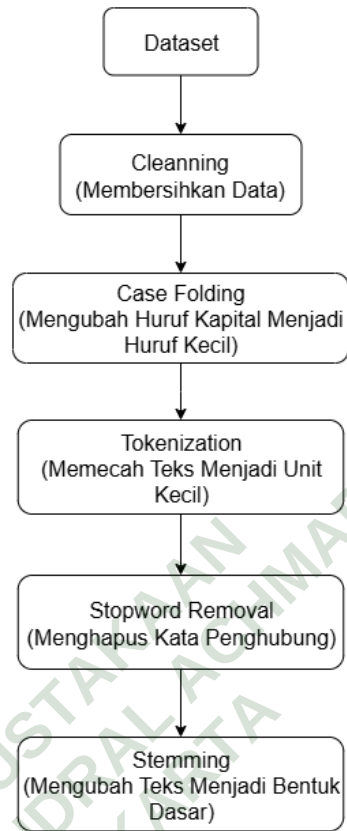
- a. Pembersihan data dilakukan melalui proses *preprocessing*.
- b. Pemberian bobot fitur dilakukan menggunakan metode TF-IDF.
- c. Seleksi fitur dilakukan berdasarkan nilai *Document Frequency* (DF).



Gambar 3.2 Pelabelan Kelas

3. Tahap *Preprocessing*

Data ulasan umumnya memiliki format yang tidak seragam, termasuk penggunaan singkatan dan simbol, sehingga memerlukan proses *preprocessing* untuk standarisasi. *Preprocessing* merupakan pelabelan dilakukan dengan memberi tanda pada ulasan yang mengandung komentar positif atau negatif secara manual, dengan mengacu pada Kamus Besar Bahasa Indonesia (KBBI). Tahap ini merupakan proses awal dalam pengolahan teks, yang melibatkan beberapa langkah, yaitu pembersihan data untuk menghilangkan simbol-simbol tertentu, pemrosesan teks (*text processing*) guna menganalisis dan menyaring data teks tidak terstruktur untuk memperoleh informasi penting, *case folding* digunakan untuk mengubah semua huruf menjadi huruf kecil, *stopword removal* bertujuan menghapus kata-kata yang kurang berpengaruh dalam kalimat, *tokenization* memisahkan kalimat menjadi kata-kata tunggal, dan *stemming* menghilangkan imbuhan agar mendapatkan kata dasar (Nurwanda dkk., 2024).



Gambar 3.3 Tahap *Preprocessing*

4. *Splitting Data*

Splitting data adalah metode yang digunakan untuk membagi data dataset, yang menjadi salah satu faktor penting dalam menentukan kualitas performa model klasifikasi pada algoritma pembelajaran mesin. Proses ini dilakukan dengan memisahkan data menjadi bagian data latih dan data uji. Proses validasi dalam split data sangat penting untuk memastikan bahwa setiap data memiliki kesempatan yang sama untuk digunakan baik dalam tahap pelatihan maupun pengujian model (Oktafiani dkk., 2023).

5. TF-IDF

TF-IDF (*Term Frequency-Inverse Document Frequency*) merupakan salah satu pendekatan populer dalam pengolahan bahasa alami yang

digunakan untuk mengevaluasi seberapa signifikan suatu kata dalam suatu dokumen relatif terhadap seluruh korpus. Komponen *Term Frequency* (TF) mengukur seberapa sering sebuah kata muncul dalam satu dokumen, biasanya dengan membandingkan jumlah kemunculan kata tersebut terhadap total kata dalam dokumen. Dalam beberapa pendekatan, nilai TF ini dapat dimodifikasi menggunakan metode pembobotan lanjutan. Di sisi lain, *Inverse Document Frequency* (IDF) berfungsi mengidentifikasi katakata yang lebih unik dalam konteks koleksi dokumen—semakin jarang sebuah kata muncul di berbagai dokumen, semakin tinggi nilai IDF-nya. Nilai ini diperoleh dengan membagi jumlah total dokumen dengan jumlah dokumen yang mengandung kata tersebut, kemudian hasilnya dihitung dalam bentuk logaritmik untuk menghasilkan skala yang lebih seimbang. Bobot akhir suatu kata ditentukan dengan mengalikan nilai TF dan IDF, sehingga menghasilkan skor yang menggambarkan relevansi kata tersebut dalam dokumen tertentu dibandingkan seluruh kumpulan dokumen (Septiani & Isabela, 2022).

6. Pengklasifikasian

Pada tahap ini, metode *Naive Bayes Classifier* diterapkan. *Naive Bayes* merupakan salah satu teknik klasifikasi yang umum digunakan dan terbukti efektif dalam bidang *machine learning* serta *data mining*. *Naive Bayes* merupakan metode sederhana yang dikembangkan berdasarkan prinsip *Teorema Bayes*, dengan memperhatikan kondisi yang ada serta peluang untuk setiap kondisi tersebut. Metode ini berfungsi sebagai solusi pendukung dalam pengambilan keputusan, baik di dunia bisnis maupun pendidikan. Dengan penerapan *Naive Bayes*, informasi yang dibutuhkan dalam penelitian dapat diperoleh secara lebih rinci. *Naive Bayes Classifier* sendiri merupakan teknik klasifikasi berbasis statistik yang mengandalkan *Teorema Bayes* dan mengasumsikan adanya independensi antar fitur. Teknik ini

menghitung peluang suatu data tergolong dalam kategori tertentu dengan mempertimbangkan distribusi nilai fitur pada data latih. Meskipun tergolong sederhana, *Naive Bayes Classifier* mampu menghasilkan performa yang baik dalam berbagai permasalahan klasifikasi, termasuk dalam mengidentifikasi fitur-fitur yang paling berpengaruh terhadap hasil prediksi berdasarkan nilai probabilitasnya (Julkarnain & Yustiardin, 2024).

3.3.2 Metode Pengujian Sistem

Dalam penelitian ini, metode yang digunakan untuk pengujian sistem mencakup beberapa teknik sebagai berikut:

1. Pengujian *Confusion Matrix*

Teknik *Confusion Matrix* adalah metode evaluasi yang digunakan untuk mengukur keakuratan suatu model klasifikasi dalam mengidentifikasi objek secara tepat (Arifqi dkk., 2024). Evaluasi sebagai ukuran kinerja sistem ini memperlihatkan perbandingan antara data asli dengan data yang diprediksi oleh sistem.

Untuk mengevaluasi model analisis sentimen, digunakan *Confusion Matrix* yang memungkinkan evaluasi menggunakan metrik seperti akurasi, presisi, *recall*, dan *F1-score*.

Tabel 3.1 *Confusion Matrix*

Kelas	Prediksi Negatif	Prediksi Netral	Prediksi Positif
Kelas Negatif	TN	FNt	FP
Kelas Netral	FN	TNt	FP
Kelas Positif	FN	FNt	TP

Keterangan:

- True Positive* (TP) merupakan data yang diprediksi sebagai kelas positif, dan dalam kenyataannya memang benar termasuk dalam kelas tersebut.

- b. *False Positive* (FP) adalah data yang diklasifikasikan sebagai positif oleh model, padahal sebenarnya berasal dari kelas negatif atau kelas lainnya.
- c. *True Negative* (TN) mengacu pada data yang diklasifikasikan sebagai kelas negatif, dan secara faktual memang berasal dari kelas negatif.
- d. *False Negative* (FN merujuk pada data yang diprediksi sebagai negatif, namun sebenarnya tergolong ke dalam kelas positif atau netral.
- e. *True Netral* (TNt) adalah data yang diprediksi berada dalam kelas netral dan faktanya memang tergolong ke dalam kelas netral.
- f. *False Netral* (FNt) menunjukkan data yang diklasifikasikan sebagai netral oleh model, padahal sebenarnya berasal dari kelas positif atau negatif.

2. Akurasi

Akurasi merupakan ukuran persentase kelas yang berhasil diprediksi dengan benar oleh model yang telah dibuat (Amaliah dkk., 2022) . Berikut rumus untuk menghitungnya:

$$Accuracy = \frac{TP + TNt + TN}{TP + TNt + TN + FP + FNt + FN} \quad (10)$$

3. Presisi

Presisi adalah ukuran ketepatan sistem dalam memberikan jawaban yang relevan terhadap permintaan informasi dari pengguna (Clara dkk., 2021). Berikut rumus untuk menghitungnya:

$$Precision(positif) = \frac{TP}{TP + FP} \quad (11)$$

4. Recall

Recall merupakan persentase prediksi program terhadap data yang tidak sesuai dengan kelas aslinya, dan *recall* juga dikenal dengan

istilah sensitivitas (Sadeli & Lawanda, 2023). Berikut rumus untuk menghitungnya:

$$Recall(positif) = \frac{TP}{TP + FN = FNT} \quad (12)$$

5. *F1-Score*

F1-score merupakan metrik evaluasi yang mengintegrasikan recall dan presisi, serta menunjukkan rata-rata berbobot antara keduanya (Ramdloni dkk., 2022). Berikut rumus untuk menghitungnya:

$$F1(positif) = 2x \frac{Precision(positif) \times Recall(positif)}{Precision(positif) + Recall(positif)} \quad (13)$$

3.3.3 Teknik Analisis Data

Analisis data merupakan tahapan pengolahan dan pengorganisasian hasil dari observasi, wawancara, serta berbagai data lainnya secara terstruktur agar peneliti dapat memperoleh pemahaman yang lebih mendalam mengenai kasus yang sedang diteliti, sekaligus menyajikannya dalam bentuk temuan bagi pihak lain. Tujuan dari analisis data adalah mengatur data sedemikian rupa sehingga memiliki makna dan mudah dipahami. Para peneliti menyatakan bahwa tidak ada satu metode yang dianggap mutlak benar dalam mengorganisasi, menganalisis, atau menginterpretasi data. Oleh karena itu, prosedur analisis data harus disesuaikan dengan tujuan dari penelitian yang dilakukan. Tahapan analisis data yang diterapkan dalam penelitian ini meliputi beberapa proses sebagai berikut :

1. Pengumpulan Data

Pengumpulan data merupakan tahap krusial dalam penelitian, terutama dalam analisis sentimen. Data yang dikumpulkan dapat berupa teks, dokumen, atau ulasan yang mengandung opini atau sentimen. Sumber data ini bisa berasal dari berbagai platform media sosial, platform digital, atau melalui wawancara. Langkah ini sangat penting karena kualitas dan relevansi data yang dikumpulkan akan memengaruhi hasil

analisis dan interpretasi dalam penelitian (Aryanti & Mahendra, 2023).

2. Klasifikasi Sentimen

Dalam analisis sentimen, teks atau dokumen yang telah dikumpulkan diklasifikasikan ke dalam kategori sentimen, yaitu positif, negatif, atau netral. Proses klasifikasi ini dapat dilakukan dengan menggunakan teknik pemrosesan bahasa alami (*Natural Language Processing/NLP*) maupun melalui pemanfaatan perangkat lunak khusus untuk analisis sentimen.

3. Penghitungan Frekuensi

Setelah teks atau dokumen dikelompokkan berdasarkan kategori sentimen, langkah selanjutnya adalah menghitung jumlah frekuensi masing-masing kategori. Proses penghitungan ini dapat dilakukan secara absolut, yang menunjukkan jumlah total sentimen positif dan negatif, atau secara relatif, yaitu dengan menghitung persentase sentimen positif dan negatif dari seluruh dokumen yang dianalisis.

4. Interpretasi Hasil

Selanjutnya, hasil analisis data dianalisis lebih lanjut untuk mengungkapkan pemahaman mengenai sentimen atau pandangan yang tersirat dalam dokumen atau teks. Melalui proses interpretasi ini, membantu memberikan pemahaman yang lebih mendalam tentang pandangan publik terhadap suatu topik atau produk tertentu (Bria & Witanti, 2023).