

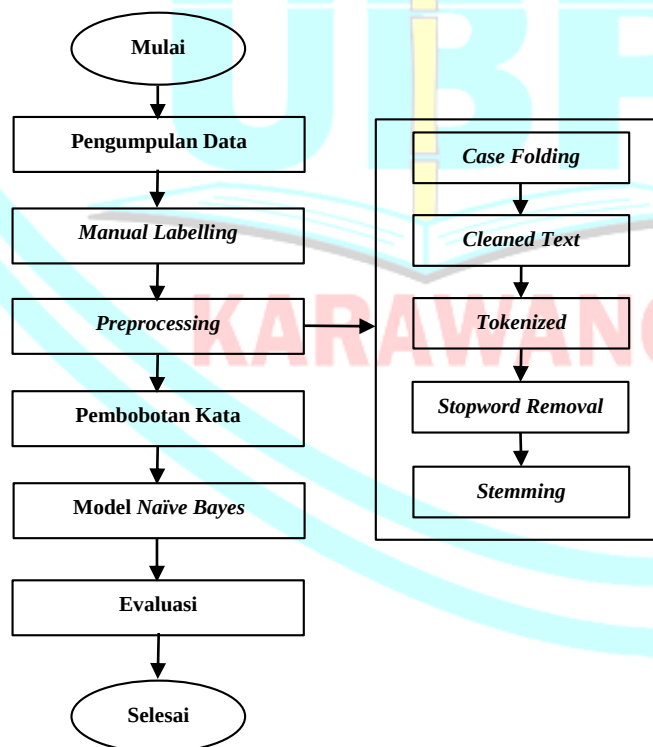
BAB III

METODE PENELITIAN

3.1 Objek Penelitian

Objek penelitian ini adalah ulasan aplikasi DANA Dompet Digital Indonesia yang diperoleh melalui layanan *Google Play Store* menggunakan teknik *scraping* dengan data yang diperoleh 1.500 ulasan dari bulan September 2024 - Desember 2024. Ulasan tersebut kemudian dikelompokkan ke dalam lima kategori yang telah ditentukan sebelumnya. Kategorinya, meliputi: (1) Transaksi, (2) Keamanan, (3) Performa Aplikasi, (4) Pelayanan, serta (5) Aktivasi dan Verifikasi. Adapun tahapan-tahapan penelitian ini akan diuraikan pada alur berikut (Helmayanti et al., 2023).

3.2 Prosedur Penelitian



Gambar 3.1 Alur Penelitian (Ramadhan & Sugianto, 2024)

3.2.1 Pengumpulan data

Pada tahap ini, data diperoleh dari platform *Google Play Store* melalui teknik *web scraping*. Teknik ini adalah metode otomatis yang memanfaatkan program atau skrip komputer untuk mengekstraksi data dari situs web. Teknik

web scraping pada penelitian ini diterapkan menggunakan *Google Colab* dengan bahasa pemrograman Python serta library *google-play-scraper* yang perlu diinstal terlebih dahulu ke dalam *Python*. Ulasan yang berhasil diambil dari proses *scraping* disimpan dalam bentuk *dataframe*, lalu dikonversi ke dalam format CSV untuk memudahkan proses pengolahan data.

3.2.2 *Manual Labelling*

Tahap selanjutnya adalah *manual labelling* atau pelabelan *manual*, di mana setiap ulasan akan dikelompokkan ke dalam 5 kategori berikut:

1. Transaksi

Kategori transaksi ini berfokus pada ulasan yang membahas seputar proses pembayaran, pengiriman uang, pengisian saldo.

2. Keamanan

Ulasan pada kategori keamanan yaitu ulasan yang membahas aspek keamanan data, keamanan transaksi, perlindungan akun, privasi, serta kasus penipuan.

3. Performa Aplikasi

Kategori yang berisi ulasan tentang kecepatan, kestabilan, *bug*, *error*, atau kemudahan penggunaan aplikasi.

4. Pelayanan

Kategori yang berisi ulasan mengenai *respons* dari *customer service*, pengalaman pengguna dengan dukungan teknis, serta kualitas layanan yang diberikan.

5. Aktivasi dan Verifikasi

Ulasan yang berkaitan dengan proses pendaftaran akun, verifikasi identitas, serta kendala selama proses aktivasi maupun verifikasi.

Manual labelling atau pelabelan manual yaitu pemberian label atau kategori secara manual oleh pelabel, berdasarkan sudut pandang masing-masing serta kriteria tertentu yang telah ditetapkan sebelumnya. Proses ini dilakukan tanpa bantuan otomatisasi atau algoritma, melainkan melalui pengamatan dan interpretasi pelabel terhadap setiap ulasan.

3.2.3 *Preprocessing*

Preprocessing merupakan tahapan untuk menyiapkan atau membersihkan data agar siap untuk diproses pada tahapan selanjutnya, hal tersebut dikarenakan data ulasan yang diambil dari Layanan *Google Play Store* tidak terstruktur. Tahapan pertama dalam proses *preprocessing* data pada penelitian ini adalah

1. *Case Folding* : Mengkonversi seluruh huruf kapital menjadi huruf kecil.
2. *Cleaning Text* : Tahap menghapus karakter yang tidak bermakna dari data, seperti tanda baca, angka, maupun simbol yang tidak memiliki nilai informasi.
3. *Tokenized* : Tahap pemisahan teks menjadi komponen-komponen yang lebih kecil seperti kata atau frasa, yang disebut sebagai token.
4. *Stopword Removal* : Kata-kata yang tidak relevan atau terlalu umum akan dihilangkan, contoh *stopword* dalam bahasa Indonesia yaitu “dan”, “yang”, “dari”.
5. *Stemming* : Bertujuan untuk mengganti kata imbuhan dengan kata dasar, sehingga kata tersebut dapat diubah dari kata turunan menjadi bentuk kata dasarnya (kata baku).

3.2.4 **Pembobotan Kata**

Teknik terminologi pembobotan frekuensi, yang dikenal sebagai *Term Frequency-Inverse Document Frequency* (TF-IDF), digunakan untuk mengkonversi data berbentuk teks menjadi data berbentuk angka atau numerik. Tujuan dari metode ini yaitu memberikan bobot pada kata-kata atau fitur tertentu dalam sebuah dataset. TF-IDF merupakan metode yang memberikan nilai bobot pada kata-kata berdasarkan berapa kali kata-kata tersebut muncul dalam dokumen. Suatu term, yang dapat berupa kata, frasa, atau elemen indeks dalam naskah, ditunjukkan dengan *Term Frequency* (TF). Semakin sering sebuah istilah digunakan, semakin tinggi bobotnya (Cindy Catherine Yolanda et al., 2024).

3.2.5 Model Naïve Bayes

Dalam penelitian ini, pemodelan yang diterapkan adalah *Multinomial Naïve Bayes*. Model ini dipilih karena kesederhanaannya serta popularitasnya yang tinggi. Model ini memungkinkan setiap atribut berkontribusi secara setara pada keputusan akhir. Salah satu keunggulan dari metode *Naïve Bayes* adalah kemampuannya untuk memperkirakan parameter rata-rata dan variasi variabel yang diperlukan untuk klasifikasi, hanya dengan memerlukan jumlah data latihan yang relatif kecil. Teorema bayes mempunyai rumus:

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)} \quad (1)$$

Keterangan :

$P(c|x)$ merepresentasikan probabilitas posterior untuk kelas c berdasarkan fitur x . Sementara itu, $P(x|c)$ menyatakan probabilitas kemunculan fitur x dalam kelas c . Adapun $P(c)$ menggambarkan probabilitas prior untuk kelas c , dan $P(x)$ menunjukkan probabilitas keseluruhan dari fitur x .

Sebelum tahap pemodelan dilakukan, dataset perlu dipisahkan menjadi dua bagian. Bagian pertama, yaitu data latih, digunakan untuk melatih model dalam mengenali pola klasifikasi, sementara bagian kedua, yaitu data uji, dimanfaatkan untuk menilai seberapa baik kinerja model tersebut (Sanjaya et al., 2024). Akurasi model yang lebih baik dapat dicapai dengan jumlah data latih yang lebih besar. Pemodelan ini bertujuan untuk melatih model agar dapat mengklasifikasikan dan memprediksi data.

3.2.6 Evaluasi

Langkah berikutnya adalah evaluasi, tujuan dari tahap evaluasi adalah untuk mengidentifikasi dan menetapkan hasil dari penerapan model yang telah dilakukan pada tahap sebelumnya. Dalam evaluasi model, terdapat berbagai metrik yang digunakan, termasuk akurasi, presisi, *recall*, dan *f1-score*. Akurasi mengacu pada persentase prediksi yang benar dari keseluruhan data. Presisi menunjukkan seberapa tepat model dalam memprediksi suatu kelas tertentu, *recall* mengukur kemampuan model dalam menemukan seluruh data yang termasuk dalam suatu kelas tertentu, dan *f1-score*

menggambarkan keseimbangan antara presisi dan *recall*. Selain itu, dalam tahap evaluasi *confusion matrix* juga digunakan untuk memberikan gambaran yang lebih mendetail tentang performa model. *Confusion matrix* merupakan suatu alat yang dimanfaatkan untuk mengevaluasi kinerja model dengan cara menunjukkan label hasil prediksi model dan label manual. *Matrix* ini menunjukkan seberapa akurat model memprediksi kelas yang benar, serta menggambarkan kesalahan yang dilakukan model dalam proses klasifikasi dan prediksinya.

