

BAB III METODE PENELITIAN

3.1 Objek Penelitian

Objek penelitian ini adalah komentar atau ulasan pengguna yang diambil dari aplikasi Ruangguru di *platform Google Play Store*. Ulasan-ulasan ini mencerminkan berbagai aspek yang berkaitan dengan penggunaan aplikasi, termasuk kepuasan pengguna, fitur aplikasi, serta pengalaman belajar. Penelitian ini menggunakan 5000 ulasan pengguna yang dikumpulkan secara daring.

Lokasi pengumpulan data dilakukan secara *online* melalui *Google Play Store* tanpa interaksi langsung dengan pengguna. Waktu pengumpulan data dilaksanakan pada tanggal 28 Oktober – 31 Desember 2024, yang mencakup ulasan yang diposting oleh pengguna selama periode tertentu. Fokus penelitian adalah untuk mengklasifikasikan sentimen yang terkandung dalam komentar pengguna menjadi beberapa kategori, yaitu Baik Sekali, Baik, Cukup Baik, dan Kurang Baik.

Pemilihan *Google Play Store* sebagai sumber data didasarkan pada ketersediaan ulasan yang mencakup berbagai latar belakang pengguna serta kemudahan dalam memperoleh data secara terbuka. Data dari *Google Play Store* dianggap representatif karena *platform* ini menjadi salah satu media utama bagi pengguna dalam menyampaikan evaluasi terkait aplikasi Ruangguru.

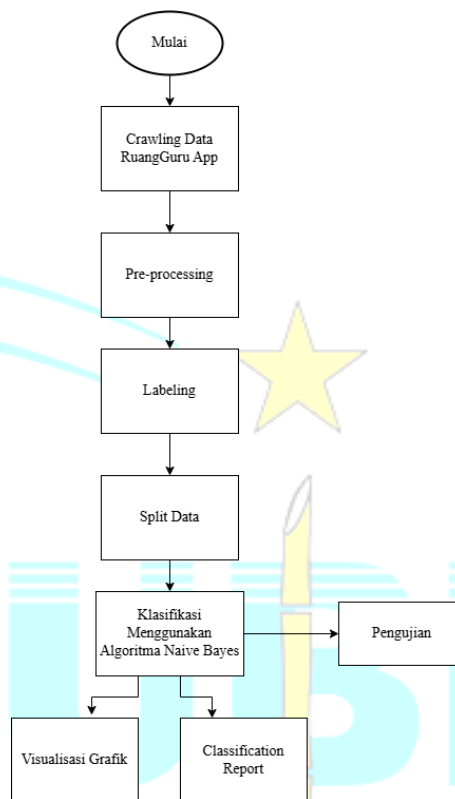
Penelitian ini bertujuan memberikan **wawasan** tentang bagaimana pengguna menilai aplikasi Ruangguru, yang pada gilirannya diharapkan dapat memberikan umpan balik penting untuk pengembangan aplikasi tersebut di masa mendatang.

Dengan demikian, penelitian ini membatasi lingkup analisis pada sentimen pengguna terhadap aplikasi Ruangguru, serta tidak mencakup analisis *fitur* teknis aplikasi secara mendalam ataupun ulasan dari *platform* lain selain *Google Play Store*.

3.2 Prosedur Penelitian

Prosedur penelitian yang digunakan dalam penelitian ini terdiri dari beberapa tahapan utama, yaitu pengumpulan data, *pre-processing* data, pelabelan data, pembagian data (*split data*), klasifikasi menggunakan *Naïve Bayes*, Evaluasi dan visualisasi hasil klasifikasi

(Sanjaya, 2024). Berikut adalah diagram alur proses penelitian beserta penjelasan tiap tahapannya:



Gambar 3.1 Flowchart Pengumpulan Data

3.2.1 Pengumpulan Data

Pengumpulan data merupakan tahapan penting dalam penelitian ini, di mana data dikumpulkan untuk dianalisis dan digunakan sebagai informasi yang mendukung penelitian (Sanjaya, 2024). Studi ini menggunakan komentar pengguna dari aplikasi Ruangguru di *Google Play Store*, yang dikumpulkan melalui proses *web scraping* untuk mengekstrak informasi relevan terkait topik penelitian. Ulasan yang diambil adalah komentar pengguna yang berkaitan langsung dengan pengalaman mereka dalam menggunakan aplikasi. Dari data yang diperoleh, hanya kolom yang memuat isi komentar (*content*) yang digunakan dalam proses analisis, karena kolom lain tidak relevan dan dapat menghambat jalannya penelitian.

Dalam penelitian ini, terdapat empat kategori sentimen, yaitu:

- a. Baik Sekali, Menampilkan ulasan yang sangat positif dan memuaskan. Biasanya ditandai dengan pujian yang kuat atau kesan yang sangat memuaskan dari pengguna.

- b. Baik, Merupakan ulasan yang mengandung sentimen positif, menunjukkan bahwa pengguna merasa cukup puas terhadap aplikasi, meskipun tidak seantusias kategori "Baik Sekali".
- c. Cukup Baik, Merupakan ulasan yang bersifat netral atau menengah, di mana pengguna merasa aplikasi cukup membantu.
- d. Kurang Baik, Merupakan ulasan yang mengandung kata-kata bernuansa negatif. Biasanya ditandai dengan keluhan, kritik, atau ketidakpuasan terhadap fitur atau layanan dalam aplikasi.

3.2.2 Pre-processing Data

Tahap *pre-processing* atau pembersihan data bertujuan untuk menghapus komponen yang tidak diperlukan, seperti data duplikat, karakter simbol, link, maupun kata-kata tidak sesuai kaidah bahasa (Kurniawan, 2023). Adapun proses tahapannya terdiri dari :

- a. *Text Clean*, Fase awal pembersihan data melibatkan penghilangan komponen yang tidak relevan, termasuk hashtag, mention, tautan, nilai numerik, dan teks yang berulang.
- b. *Case Folding*, adalah metode yang digunakan untuk mengubah teks menjadi huruf kecil, sehingga menghilangkan perbedaan antara huruf kecil dan besar untuk memastikan keseragaman. Proses ini sangat penting untuk memastikan bahwa perbedaan makna tidak timbul hanya karena perbedaan kapitalisasi.
- c. *Tokenizing*, melibatkan pemecahan sistematis dokumen menjadi komponen yang lebih kecil, termasuk paragraf, kalimat, atau kata, untuk memudahkan pemrosesan selanjutnya.
- d. *Stemming*, Tahap ini dilakukan untuk mengubah kata turunan ke bentuk dasarnya. Contohnya, “memahami”, “dipahami”, dan “pemahaman” akan disederhanakan menjadi “paham”.
- e. *Stopword*, Tahap ini dalam proses penambangan teks adalah fase di mana istilah-istilah yang sering muncul dihilangkan karena tidak relevan dengan materi yang dianalisis.

3.2.3 Pelabelan Data

Dalam tahap ini, setiap ulasan akan dianalisis untuk kemudian diberi label sesuai dengan sentimen untuk memfasilitasi klasifikasi opini. Pemberian label ini berfokus pada memetakan opini pengguna ke dalam empat kategori sentimen, sehingga proses klasifikasi dengan *Naïve Bayes Classifier* dapat berjalan secara efektif (Rizaldi, 2023).

3.2.3 Split Data

Split data adalah tahapan dalam pemodelan yang membagi dataset ke dalam dua bagian, Model dikembangkan menggunakan data pelatihan, sedangkan data uji digunakan untuk mengevaluasi kinerja model (Poerwandono, 2023). Proporsi pembagian sebesar 80:20 memungkinkan algoritma belajar dari mayoritas data, sementara sisanya digunakan sebagai data baru untuk pengujian, sehingga hasil prediksi yang dihasilkan dapat lebih akurat.

3.2.4 Naive Bayes classifier

Tahap klasifikasi merupakan inti dari penelitian ini, di mana data komentar yang telah melalui proses *pre-processing* dan pelabelan selanjutnya dianalisis untuk mengidentifikasi sentimen secara otomatis. Model klasifikasi yang digunakan adalah algoritma *Naïve Bayes Multinomial*, sebuah metode statistik berbasis probabilitas yang banyak digunakan dalam klasifikasi teks karena kesederhanaannya, kecepatan pelatihan yang tinggi, serta kinerjanya yang baik pada data teks berskala besar.

Dalam penelitian ini, diterapkan varian *Multinomial Naïve Bayes*, yang paling cocok untuk data diskrit seperti frekuensi kata dalam dokumen teks. Model ini memperhitungkan frekuensi kemunculan kata untuk memprediksi kelas sentimen suatu komentar. Sebelum digunakan dalam proses klasifikasi, data teks terlebih dahulu diubah ke dalam bentuk numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). TF-IDF memberikan bobot pada setiap kata dalam dokumen berdasarkan seberapa sering kata tersebut muncul di dokumen tertentu dibandingkan dengan seluruh dokumen dalam dataset, sehingga kata-kata yang terlalu umum memiliki bobot lebih rendah, dan kata-kata yang lebih unik memiliki bobot lebih tinggi.

3.2.5 Evaluasi dan Visualisasi

Untuk menilai sejauh mana model mampu mengklasifikasikan sentimen secara akurat, dilakukan tahap evaluasi performa model. Penilaian ini mencakup pengukuran kinerja model klasifikasi menggunakan empat metrik utama, yaitu akurasi (*accuracy*), presisi (*precision*), *recall*, dan *F1-score*. Metrik-metrik tersebut dipilih karena dapat memberikan gambaran menyeluruh terhadap kemampuan model dalam mengenali dan membedakan setiap kategori sentimen, yaitu Baik Sekali, Baik, Cukup Baik, dan Kurang Baik.

- a. Akurasi mencerminkan tingkat keseluruhan prediksi yang benar terhadap total data uji.
- b. *Precision* menggambarkan seberapa tepat model dalam memprediksi suatu label, dengan fokus pada kebenaran hasil prediksi positif.

- c. *Recall* mengukur sejauh mana model mampu mengenali data yang benar-benar relevan dalam setiap kategori.
- d. *F1-score* menggabungkan *precision* dan *recall* untuk menilai keseimbangan antara keduanya dalam satu nilai metrik.

Selain evaluasi metrik numerik, penelitian ini juga menyertakan visualisasi untuk memperjelas dan memperkuat pemahaman terhadap hasil klasifikasi. Beberapa bentuk visualisasi yang digunakan antara lain:

- a. *Pie chart*, yang menunjukkan distribusi persentase hasil klasifikasi keempat kategori sentimen.
- b. *Bar chart*, yang menampilkan jumlah komentar dalam masing-masing label hasil prediksi model.
- c. *Confusion matrix* divisualisasikan dalam bentuk *heatmap*, agar lebih mudah mengidentifikasi pola kesalahan klasifikasi antar kategori.
- d. *Word cloud*, yang memberikan representasi visual terhadap kata-kata yang paling sering muncul dalam data komentar pengguna.

Evaluasi ini membantu dalam menilai efektivitas algoritma *Naïve Bayes Multinomial* dalam pengolahan data teks dan klasifikasi sentimen, serta memberikan wawasan tambahan melalui pendekatan visual untuk memahami pola-pola dominan yang muncul dari hasil klasifikasi.