

BAB III

METODE PENELITIAN

A. Jenis Penelitian

Penelitian ini menggunakan pendekatan kuantitatif berbasis *data mining* untuk menganalisis sentimen 2289 ulasan pengguna *ChatGPT* di *Google Play Store*. Sentimen diklasifikasikan menjadi positif, negatif, dan netral menggunakan empat algoritma *Naive Bayes*, *SVM*, *Decision tree*, dan *Random forest* dengan pembobotan TF-IDF. Kinerja model dievaluasi menggunakan *Confusion matrix* dengan metrik akurasi, presisi, recall, dan F1-score. Selain itu, data hasil pelabelan divisualisasikan dalam bentuk *wordcloud* untuk mengidentifikasi kata dominan pada tiap kategori sentimen.

B. Setting Penelitian

Penelitian ini dilaksanakan secara daring dengan menggunakan data ulasan aplikasi *ChatGPT* yang diambil dari *Google Play Store*. Fokus utama penelitian mencakup pengumpulan data, tahapan prapemrosesan (*preprocessing*), pelabelan dan visualisasi hasil pelabelan menggunakan *wordcloud*, penerapan algoritma klasifikasi, serta evaluasi kinerja model. Data dikumpulkan dalam rentang waktu 3 bulan dimulai bulan Januari – Maret 2025 dengan ulasan yang paling relevan (*most relevan*).

C. Alat dan Bahan

1. Alat

Penelitian ini dilakukan menggunakan laptop dengan spesifikasi memadai yang digunakan sebagai perangkat utama untuk pengolahan data. Disini *Python* digunakan sebagai alat utama mulai dari pengumpulan data menggunakan *web scraping* hingga implementasi dan evaluasi algoritma yang didukung oleh *library* seperti *pandas*, *numpy*, dan *nlTK*, serta data diambil dari aplikasi *Play Store* yang menyediakan informasi ulasan pengguna mengenai aplikasi *ChatGPT*.

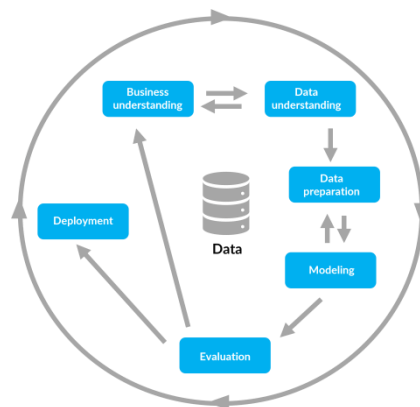
2. Bahan

Dataset yang digunakan dalam penelitian ini adalah kumpulan ulasan pengguna aplikasi *ChatGPT* yang diambil dari *Google Play Store*. Data dikumpulkan sebanyak 2289 sampel ulasan yang paling relevan yang mencerminkan pengalaman terkini pengguna terhadap aplikasi *ChatGPT* melalui teknik *web scraping*. Variabel yang terdapat dalam dataset ini terdiri dari:

- a. Teks ulasan: Berisi opini atau komentar pengguna. Variabel ini digunakan sebagai fitur utama (variabel independen) dalam klasifikasi sentimen.
- b. Skor *rating* (1–5): Berisi penilaian pengguna terhadap aplikasi. Digunakan sebagai dasar pelabelan sentimen (variabel dependen).
- c. Tanggal ulasan: Berisi informasi waktu ulasan dikirimkan. Digunakan untuk proses filter data berdasarkan periode waktu tertentu, namun tidak digunakan dalam proses klasifikasi sentimen.

D. Metode Implementasi

Metode implementasi dalam penelitian ini disusun berdasarkan pendekatan CRISP-DM (*Cross-Industry Standard Process for Data Mining*) yang terdiri dari enam tahap utama (Amri, 2020). Pendekatan ini digunakan untuk menganalisis data ulasan aplikasi *ChatGPT* di *Google Play Store* menggunakan empat algoritma klasifikasi yaitu *Naive bayes*, *Support vector machine (SVM)*, *Random forest*, dan *Decision tree*. Adapun langkah-langkah rinci tiap tahap sebagai berikut:



Gambar 3. 1 Tahapan CRISP-DM

1. *Business Understanding*

Tahap ini bertujuan untuk memahami latar belakang, permasalahan, dan tujuan penelitian. *ChatGPT* sebagai aplikasi berbasis AI telah digunakan secara luas di Indonesia untuk berbagai keperluan, namun ulasan pengguna di *Google Play Store* belum dianalisis secara sistematis. Penelitian ini menggunakan 2.289 data ulasan (Januari–Maret) yang dikelompokkan ke dalam tiga kategori sentimen: positif, netral, dan negatif. Selain itu, dilakukan visualisasi *wordcloud* untuk menyoroti kata-kata kunci pada tiap kategori. Penelitian ini membandingkan performa algoritma *Naive bayes*, *Support vector machine* (SVM), *Random forest*, dan *Decision tree* menggunakan metrik accuracy, precision, recall, dan F1-score, dengan harapan memberikan informasi bagi pengembang serta kontribusi pada penelitian analisis sentimen berbasis teks ulasan.

2. *Data Understanding*

Penelitian dimulai dengan scraping data ulasan aplikasi *ChatGPT* menggunakan library *google-play-scraper* pada Python yang dikumpulkan dalam rentang waktu Januari hingga Maret 2025. Data yang dikumpulkan mencakup teks ulasan, skor *rating*, dan tanggal ulasan. Data diambil berdasarkan ulasan paling relevan dan difilter untuk wilayah dan bahasa Indonesia. Distribusi *rating* divisualisasikan terlebih dahulu untuk memberikan gambaran awal terhadap sebaran opini pengguna.

3. *Data Preparation*

Pada tahap ini, data mentah diproses dan disiapkan untuk pemodelan melalui tahapan berikut:

a. *Preprocessing*, terdiri dari:

1) *Case Folding dan Delimiter*:

Mengubah semua huruf menjadi huruf kecil (*case folding*) untuk menyeragamkan format teks. Proses ini dikombinasikan dengan pembersihan (*delimiter*) yang mencakup penghapusan

URL, mention, hashtag, simbol, emotikon, tanda baca, angka, whitespace berlebih, dan komentar duplikat, guna menghilangkan elemen non-linguistik dari data teks.

2) *Tokenizing*:

Memecah kalimat atau teks ulasan menjadi potongan kata (token) berdasarkan spasi. Tahap ini bertujuan mempersiapkan teks agar dapat dianalisis pada level kata.

3) *Filtering*:

Menghapus kata-kata tidak penting atau yang memiliki informasi rendah (*stopword*), menggunakan *stopword* list dari library NLTK dan daftar tambahan khusus.

4) *Stemming*:

Mengubah kata-kata hasil *filtering* menjadi bentuk dasarnya menggunakan library Sastrawi. Tahapan ini bertujuan untuk mengurangi variasi kata dan menyederhanakan representasi teks.

b. Labelling Data

Proses pemberian atau pembagian label pada dataset yang sudah selesai melalui tahap *preprocessing*. Pada penelitian ini pelabelan menggunakan tiga kategori, yaitu : sentimen positif, sentiment negatif dan sentimen netral. Penentuan label mengaju pada *score/rating* dari pengguna *google Play Store* dengan kondisi sebagai berikut (Puh & Bagić Babac, 2023) (Fang & Zhan, 2015):

Tabel 3. 1 Kategori Pelabelan

Score (rating)	Label
>3	Positif
<3	Negatif
=3	Netral

c. *TF-IDF Vectorization and Handling Inbalance*

Data teks hasil *preprocessing* dikonversi ke dalam bentuk numerik menggunakan metode TF-IDF sebagai input model. Untuk mengatasi data yang tidak seimbang digunakan teknik Handling inbalance dengan metode SMOTE (*Synthetic Minority Over-sampling Technique*), yang bekerja dengan membuat sampel sintetis pada kelas minoritas sehingga distribusi kelas lebih seimbang tanpa menghilangkan informasi penting dari data.

4. *Modelling*

Tahap ini data dibagi dengan rasio 80% untuk pelatihan (training) dan 20% untuk pengujian (testing). Rasio 80:20 merupakan standar umum dalam praktik *machine learning*, terutama ketika ukuran dataset cukup besar dan penyebaran datanya merata (Zavarella, 2025). Selanjutnya, dilakukan pembangunan model klasifikasi dengan memanfaatkan beberapa algoritma, di mana masing-masing dijalankan menggunakan parameter bawaan (*default*) dari library Python. Dalam penelitian ini, parameter tidak diubah dan tetap menggunakan (*default*) dari library Python agar semua algoritma diuji dengan kondisi yang sama sehingga hasil perbandingan lebih adil (*API Reference*, n.d.):

- a. *Naive bayes* menggunakan MultinomialNB.
- b. *Support vector machine* (SVM) menggunakan SVC.
- c. *Decision tree* menggunakan DecisionTreeClassifier.
- d. *Random forest* menggunakan RandomForestClassifier.

5. *Evaluation*

Evaluasi model dilakukan menggunakan *Confusion matrix* bertipe multi-class. Setiap matrix dikonversi menjadi bentuk biner untuk masing-masing kelas (positif, netral, negatif), dan dihitung metrik berikut:

- a. *Accuracy*: Persentase prediksi yang benar dari total data.
- b. *Precision*: Ketepatan prediksi untuk masing-masing kelas.
- c. *Recall*: Kemampuan model dalam mendeteksi kelas yang benar.

d. *F1-score*: Rata-rata harmonis dari *precision* dan *recall*.

6. *Deployment*

Pada tahap ini, data hasil pelabelan awal divisualisasikan dalam bentuk *Wordcloud* untuk setiap label sentimen (positif, netral, negatif) guna memberikan gambaran kata-kata yang sering muncul pada masing-masing kategori. Selain itu, dilakukan perbandingan performa keempat algoritma klasifikasi berdasarkan hasil evaluasi, untuk menentukan model dengan kinerja terbaik.

