

BAB III

METODE PENELITIAN

A. Jenis Penelitian

Penelitian ini adalah jenis penelitian kuantitatif deskriptif, yang bertujuan untuk menggambarkan serta menganalisis fenomena menggunakan data angka yang dikumpulkan secara sistematis tanpa adanya intervensi pada variabel yang diteliti. Metode kuantitatif deskriptif diterapkan untuk memberikan tampilan objektif mengenai kecenderungan pendapat pengguna terhadap aplikasi BRImo berdasarkan analisis statistik dari data ulasan pengguna (Noor, 2016).

Pendekatan analisis yang digunakan dalam penelitian ini adalah *Aspect Based Sentiment Analysis* (ABSA), yang memungkinkan identifikasi opini pengguna terhadap aspek-aspek spesifik dari suatu objek kajian. Analisis sentimen berbasis aspek ini diimplementasikan menggunakan pendekatan pemrosesan bahasa alami dengan memanfaatkan model IndoRoBERTa untuk mengklasifikasikan sentimen pengguna pada setiap aspek yang dianalisis, sehingga menghasilkan gambaran persepsi pengguna yang lebih mendalam terhadap aplikasi BRImo.

B. Setting Penelitian

Penelitian ini dilakukan pada ulasan pengguna aplikasi BRImo yang tersedia di Google Play Store. Setting penelitian bersifat online dan observasional, artinya peneliti hanya melakukan pengamatan terhadap data yang sudah ada tanpa melakukan intervensi langsung terhadap pengguna. Data yang digunakan berupa ulasan pengguna aplikasi BRImo yang diambil dari Google Play Store pada periode 24 juni 2025 – 24 Desember 2025. Rentang waktu tersebut dipilih agar data yang dianalisis mencerminkan persepsi dan pengalaman terkini pengguna terhadap versi terbaru aplikasi BRImo. Ulasan yang dikumpulkan mencakup komentar dan penilaian pengguna terkait berbagai aspek aplikasi, seperti fitur transaksi, kemudahan penggunaan, keamanan, serta layanan pelanggan. Dengan pendekatan observasional ini, peneliti dapat menangkap opini pengguna secara alami dan autentik tanpa pengaruh eksternal, sehingga hasil analisis mencerminkan persepsi asli dari pengguna aplikasi. Selain itu, pendekatan ini juga memungkinkan

pengumpulan data dalam jumlah besar secara efisien, yang kemudian digunakan untuk melatih dan mengevaluasi model IndoRoBERTa dalam mengklasifikasikan sentimen berbasis aspek secara akurat dan kontekstual. Adapun jadwal pelaksanaan penelitian secara rinci disajikan pada Tabel 3.1.

Tabel 3. 1 Jadwal Penelitian

No	Kegiatan	Bulan 1				Bulan 2				Bulan 3				Bulan 4			
		1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4
1	Studi Literatur & Penelitian Pendahuluan																
2	Pengumpulan Data (Scraping Ulasan)																
3	Pemahaman & Eksplorasi Data																
4	Pra-pemrosesan & Persiapan Data																
5	Pemodelan (Fine-tuning IndoRoBERTa)																
6	Evaluasi Model & Analisis Hasil																
7	Implementasi & Visualisasi Hasil																
8	Penyusunan Laporan Skripsi																

C. Alat dan Bahan

1. Alat

Penelitian ini memanfaatkan perangkat keras dan perangkat lunak sebagai berikut:

a. Perangkat keras

Laptop atau komputer dengan spesifikasi paling tidak:

- a) Prosesor: AMD FX™ FX-9800P
- b) RAM: 8 GB

b. Perangkat lunak

- a) Sistem operasional: Windows
- b) Bahasa coding: Python
- c) Lingkungan pengembangan: Google Colab

c. Library Python

- a) Pandas (versi ≥ 1.5) dan NumPy (versi ≥ 1.23) untuk manipulasi dan analisis data
- b) Transformers (Hugging Face) (versi $\geq 4.x$) untuk implementasi dan fine-tuning model IndoRoBERTa
- c) Scikit-learn (versi ≥ 1.2) untuk penilaian kinerja model dengan memanfaatkan metrik akurasi, presisi, *recall*, *F1-score*, serta *Confusion Matrix*

Pemilihan Python dan library tersebut didasarkan pada kemampuannya dalam mendukung pemrosesan bahasa alami, fleksibilitas dalam pengolahan data teks, serta ketersediaan dokumentasi dan dukungan komunitas yang luas.

2. Bahan

Bahan penelitian berupa dataset ulasan pengguna aplikasi BRImo yang diperoleh dari Google Play Store menggunakan teknik *web scraping*. Data ulasan diseleksi berdasarkan periode waktu penelitian dan kriteria relevansi,

sehingga diperoleh dataset yang merepresentasikan pengalaman pengguna terkini.

Dataset yang digunakan terdiri dari teks ulasan, *score* pengguna (1–5), dan tanggal ulasan. Teks ulasan digunakan sebagai variabel independen, sedangkan sentimen sebagai variabel dependen. Pelabelan sentimen dilakukan berdasarkan *score* pengguna dengan ketentuan: *score* 1–2 dikategorikan sebagai sentimen negatif, *score* 3 sebagai netral, dan *score* 4–5 sebagai positif, sebagaimana umum digunakan dalam penelitian analisis sentimen berbasis ulasan pengguna.

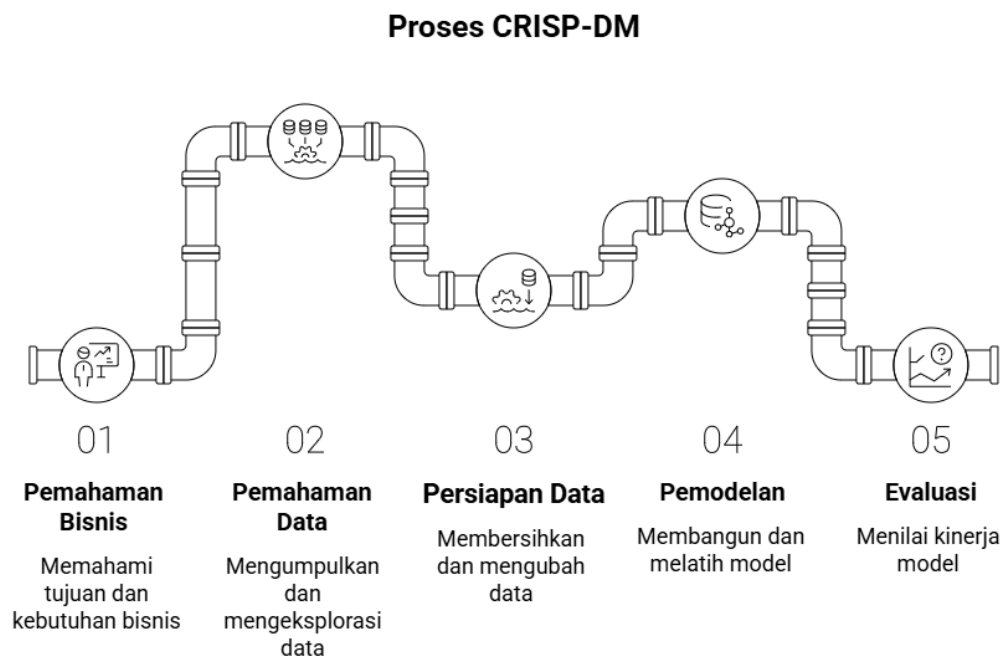
Jumlah data yang digunakan ditentukan dengan mempertimbangkan kecukupan data untuk proses fine-tuning model IndoRoBERTa serta keterwakilan distribusi sentimen pada setiap aspek yang dianalisis. Proses seleksi data dilakukan dengan menyaring ulasan yang ditulis dalam Bahasa Indonesia, memiliki isi yang relevan dengan penggunaan aplikasi, serta tidak bersifat duplikasi atau kosong.

D. Metode Penelitian

Metode yang digunakan pada penelitian ini disusun berdasarkan metode CRISP-DM (*Cross-Industry Standard Process for Data mining*), adalah standar proses yang sering kali dalam pengembangan model *data mining*. Pendekatan ini dipilih karena bersifat iteratif dan sistematis, sehingga memungkinkan peneliti untuk melaksanakan tahapan penelitian secara terstruktur mulai dari memahami bisnis, memahami data, menyiapkan data, pemodelan, sampai mengevaluasi.

Model ini digunakan karena mampu memberikan alur kerja yang jelas dan fleksibel untuk membangun serta mengevaluasi model analisis data. Setiap tahap saling berhubungan dan dapat dilakukan secara berulang sesuai kebutuhan penelitian (Ruswanti dkk., 2024). Dengan menggunakan CRISP-DM, proses pembangunan model IndoRoBERTa untuk analisis sentimen berbasis aspek pada ulasan pengguna BRI_{mo} dapat dilakukan secara konsisten, terkontrol, dan dapat dipertanggungjawabkan secara ilmiah.

Namun, pada penelitian ini tahapan CRISP-DM dibatasi hanya sampai pada tahap *Evaluation*. Tahap deployment tidak dilakukan karena penelitian ini berfokus pada pembangunan dan evaluasi performa model, bukan pada implementasi sistem secara langsung ke lingkungan pengguna. Adapun langkah-langkah dalam penelitian ini meliputi: *business understanding*, *data understanding*, *data preparation*, *modeling*, dan *evaluation*.



Gambar 3. 1 Tahapan CRISP-DM

1. Business Understanding

Tahap *Business Understanding* bertujuan untuk mengidentifikasi permasalahan penelitian serta menetapkan tujuan yang ingin dicapai. Permasalahan dalam penelitian ini adalah bagaimana mengidentifikasi dan menganalisis sentimen pengguna terhadap aplikasi BRImo berdasarkan ulasan yang diberikan pada Google Play Store.

Penelitian ini menggunakan pendekatan *Aspect Based Sentiment Analysis* (ABSA) untuk memperoleh informasi yang lebih spesifik mengenai opini pengguna terhadap aspek-aspek tertentu dari aplikasi BRImo. Aspek yang dianalisis dalam penelitian ini meliputi tampilan, fitur dan performa, serta layanan. Penentuan aspek

dilakukan menggunakan pendekatan berbasis kata kunci (*keyword-based aspect extraction*) yang diadaptasi dari penelitian terdahulu.

Model yang digunakan dalam penelitian ini adalah IndoRoBERTa, yaitu model berbasis Transformer yang telah dilatih menggunakan korpus bahasa Indonesia. Model tersebut digunakan untuk mengklasifikasikan sentimen ulasan pengguna ke dalam tiga kategori, yaitu positif, netral, dan negatif.

Untuk mengukur kinerja model digunakan beberapa metrik evaluasi, yaitu accuracy, precision, recall, F1-score, dan confusion matrix. Penggunaan metrik tersebut bertujuan untuk memastikan bahwa model mampu melakukan klasifikasi sentimen secara akurat dan konsisten.

2. Data Understanding

Tahap *Data Understanding* merupakan proses pengumpulan dan eksplorasi data yang digunakan dalam penelitian. Dataset yang digunakan berupa ulasan pengguna aplikasi BRImo yang diperoleh dari Google Play Store melalui teknik *web scraping* menggunakan bahasa pemrograman Python dengan bantuan pustaka *google-play-scraper*.

Data yang dikumpulkan mencakup 11 atribut yang tersedia pada Google Play Store, seperti identitas ulasan, informasi pengguna, teks ulasan, *score*, tanggal ulasan, dan atribut pendukung lainnya. Namun, penelitian ini hanya menggunakan atribut *content* dan *score* karena keduanya relevan dengan proses analisis sentimen yang dilakukan.

Selanjutnya dilakukan eksplorasi data untuk memahami karakteristik dataset, yang meliputi jumlah data, struktur atribut, serta distribusi *score*. Analisis distribusi *score* dilakukan untuk mengetahui kecenderungan penilaian pengguna terhadap aplikasi BRImo.

Tahap ini bertujuan untuk memastikan bahwa data yang digunakan memiliki kualitas yang baik, relevan, serta representatif untuk digunakan dalam proses analisis selanjutnya.

3. Data Preparation

Tahap ini bertujuan untuk menyiapkan data agar dapat digunakan dalam proses pemodelan. Pada penelitian ini, tahapan yang dilakukan meliputi seleksi atribut, pelabelan sentimen, penentuan aspek ulasan, pembentukan input model, dan pembagian dataset.

a. Seleksi Atribut (Filtering Data)

Dataset hasil scraping awal memiliki 11 atribut. Namun, tidak seluruh atribut digunakan dalam penelitian ini. Oleh karena itu dilakukan proses seleksi atribut untuk memilih data yang relevan dengan tujuan penelitian. Atribut yang digunakan dalam penelitian ini terdiri dari:

- 1) *content* (ulasan pengguna)
- 2) *score* (rating)

Kolom *content* digunakan sebagai data utama dalam proses analisis sentimen karena berisi opini dan pengalaman pengguna terhadap aplikasi BRImo. Sementara itu, kolom *score* digunakan sebagai dasar dalam proses pelabelan sentimen, yaitu untuk mengelompokkan ulasan ke dalam kategori sentimen positif, netral, dan negatif.

Atribut lainnya tidak digunakan karena tidak memiliki pengaruh langsung terhadap proses klasifikasi sentimen yang dilakukan dalam penelitian ini. Hasil dari proses seleksi atribut kemudian digunakan pada tahap pelabelan sentimen, penentuan aspek, dan pemodelan menggunakan algoritma IndoRoBERTa.

b. Pemeriksaan dan Persiapan Data

Sebelum proses pelabelan sentimen dilakukan, dataset terlebih dahulu diperiksa untuk mengetahui kualitas data yang akan digunakan dalam penelitian. Pemeriksaan dilakukan terhadap keberadaan URL, username, data kosong, dan karakter yang berpotensi mengganggu proses analisis.

Berdasarkan hasil pemeriksaan, tidak ditemukan data yang mengandung URL dan hanya ditemukan sejumlah kecil data yang mengandung username. Oleh karena itu, tidak dilakukan proses pembersihan data secara ekstensif seperti penghapusan URL, username, stopword removal, maupun stemming.

Keputusan ini didukung oleh penelitian (Khairani dkk., 2024) yang menunjukkan bahwa model IndoBERT dan IndoBERTweet tanpa tahapan *remove stopwords* dan *stemming* menghasilkan performa yang lebih baik dibandingkan model yang menggunakan tahapan tersebut. Penelitian tersebut memperoleh akurasi sebesar 92,54% pada model IndoBERTweet dan 88,81% pada model IndoBERT ketika tidak menerapkan *remove stopwords* dan *stemming*.

Selain itu, model IndoRoBERTa yang digunakan dalam penelitian ini merupakan model berbasis Transformer yang telah dilatih pada korpus bahasa Indonesia dan memiliki tokenizer sendiri untuk memproses teks secara kontekstual. Oleh karena itu, data ulasan dapat langsung diproses menggunakan tokenizer IndoRoBERTa sebelum tahap pelatihan model.

c. Pelabelan Sentimen

Pelabelan sentimen dilakukan menurut nilai penilaian yang diberikan oleh pengguna. Klasifikasi sentimen dibagi menjadi tiga kategori, yaitu positif, netral, dan negatif. Aturan pelabelan sentimen adalah sebagai berikut:

- 1) *Score* 1–2 termasuk dalam kategori sentimen negatif.
- 2) *Score* 3 termasuk dalam kategori sentimen netral.
- 3) *Score* 4–5 termasuk dalam kategori sentimen positif.

Pendekatan ini digunakan karena *score* yang diberikan pengguna mencerminkan tingkat kepuasan terhadap aplikasi.

d. Penentuan Aspek Ulasan

Penelitian ini menggunakan pendekatan *Aspect Based Sentiment Analysis (ABSA)*. Oleh karena itu setiap ulasan perlu dipetakan ke dalam aspek tertentu sebelum dilakukan klasifikasi sentimen.

Penentuan aspek dilakukan menggunakan metode *keyword-based aspect extraction*. Kata kunci aspek diadaptasi dari penelitian terdahulu yang membahas analisis ulasan aplikasi mobile. Aspek yang digunakan dalam penelitian ini terdiri dari:

- 1) Tampilan
- 2) Fitur dan Performa
- 3) Layanan

Setiap ulasan diperiksa berdasarkan kemunculan kata kunci yang mewakili masing-masing aspek. Ulasan yang tidak mengandung kata kunci dari aspek yang telah ditentukan tidak digunakan dalam proses analisis. Oleh karena itu, dilakukan proses filtering data sehingga hanya ulasan yang berhasil dipetakan ke dalam aspek tertentu yang digunakan dalam tahap analisis selanjutnya.

e. Pembentukan Input Model

Setelah aspek berhasil ditentukan, dilakukan pembentukan representasi input yang akan digunakan oleh model IndoRoBERTa. Input dibentuk dengan menggabungkan label aspek dan isi ulasan menggunakan format “Aspect : Content” dengan contoh “fitur_performa : aplikasi sering mengalami error saat transfer”. Penggabungan aspek dan ulasan dilakukan agar model dapat memahami konteks aspek yang sedang dianalisis pada setiap ulasan.

f. Pembagian Dataset

Dataset yang telah mengalami proses preprocessing dan penandaan sentimen selanjutnya dibagi menjadi tiga bagian, yaitu data untuk pelatihan, data untuk validasi, dan data untuk pengujian.

Rasio pembagian dataset yang diterapkan adalah 80% untuk data pelatihan, 10% untuk data validasi, dan 10% untuk data pengujian. Pembagian dataset dilakukan dengan cara stratified untuk mempertahankan keseimbangan distribusi kelas sentimen di setiap bagian dataset.

4. Modeling

Tahap ini merupakan proses pembangunan dan pelatihan model untuk melakukan klasifikasi sentimen. Model yang digunakan dalam penelitian ini adalah

IndoRoBERTa yang telah dilatih sebelumnya (*pre-trained model*), kemudian dilakukan proses *fine-tuning* menggunakan dataset penelitian.

a. Arsitektur Model

Model IndoRoBERTa digunakan dengan penambahan lapisan klasifikasi (*classification layer*) untuk melakukan prediksi terhadap tiga kelas sentimen. Representasi kalimat diperoleh dari token khusus pada awal input yang digunakan sebagai representasi keseluruhan kalimat dan diteruskan ke lapisan klasifikasi untuk menghasilkan prediksi sentimen.

b. Tokenisasi Data

Sebelum diproses oleh model, data teks terlebih dahulu diubah menjadi representasi numerik menggunakan tokenizer IndoRoBERTa. Tahapan tokenisasi meliputi:

- 1) Pemecahan teks menjadi token.
- 2) Konversi token menjadi token ID.
- 3) Penerapan padding untuk menyeragamkan panjang input.
- 4) Penerapan truncation untuk memotong teks yang melebihi batas maksimum.
- 5) Penetapan panjang maksimum input sebesar 128 token.

c. Fine-tuning Model

Proses *fine-tuning* dilakukan untuk menyesuaikan model IndoRoBERTa dengan dataset penelitian. Model dilatih menggunakan data latih dan dievaluasi menggunakan data validasi.

Parameter pelatihan yang digunakan meliputi:

- 1) *Epoch*: 10
- 2) *Batch size*: 16
- 3) *Learning rate*: $2e-5$
- 4) *Early Stopping Patience*: 2

Nilai *epoch* merupakan jumlah maksimum iterasi pelatihan yang dilakukan oleh model. Namun, pada penelitian ini diterapkan mekanisme *early stopping*

sehingga proses pelatihan dapat dihentikan lebih awal apabila tidak terdapat peningkatan performa pada data validasi.

Selama proses pelatihan, dilakukan pemantauan terhadap nilai *loss* dan performa model untuk memastikan model dapat belajar secara optimal serta menghindari terjadinya *overfitting*.

5. Evaluation

Tahap Evaluasi memiliki tujuan untuk mengukur sejauh mana model bekerja dengan menggunakan ukuran akurasi, presisi, *recall*, *F1-score*, *F1-score* dipilih sebagai ukuran utama karena dapat menggambarkan keseimbangan antara presisi dan *recall* pada data yang mungkin tidak seimbang (Sokolova & Lapalme, 2009). Evaluasi juga dilakukan secara per aspek untuk mengetahui distribusi dan kecenderungan sentimen pengguna. Apabila performa belum memenuhi kriteria keberhasilan, dilakukan iterasi pada tahap *Data Preparation* atau *Modeling*.

