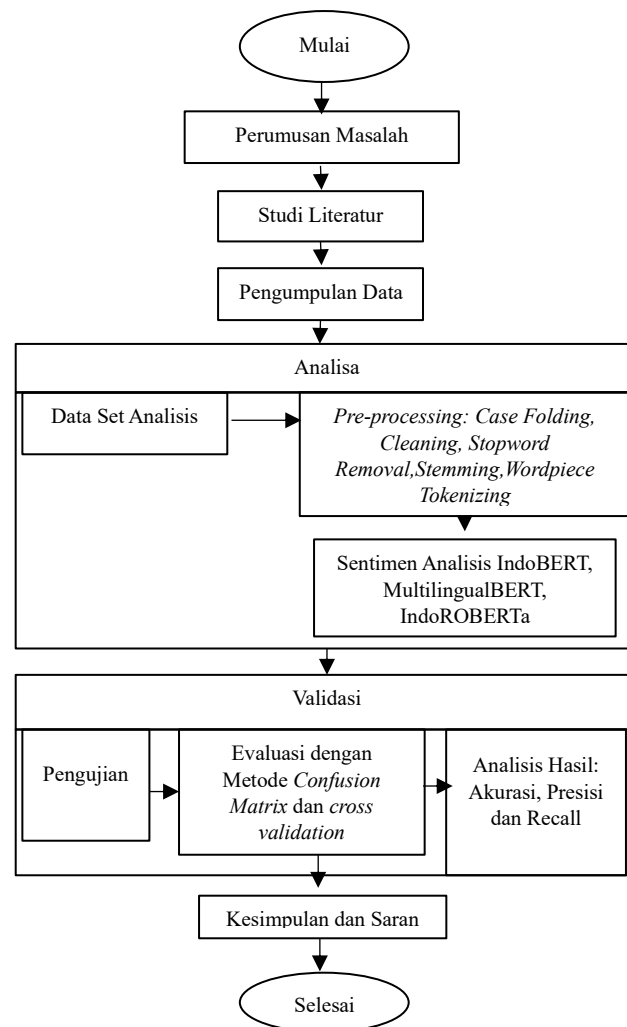


BAB III

METODOLOGI PENELITIAN

3.1 Metode Penelitian

Penelitian ini bersifat kuantitatif, karena menggunakan pendekatan berbasis data numerik yang dianalisis melalui eksperimen model-model dengan bantuan komputasi. Penelitian ini juga menerapkan pendekatan *comparative model evaluation*, yaitu membandingkan performa beberapa model berbasis Transformer seperti IndoBERT, IndoROBERTa, dan MultilingualBERT dalam melakukan analisis sentimen terhadap data teks dari media sosial X. Evaluasi dilakukan berdasarkan metrik kuantitatif seperti akurasi, *presisi*, *recall*, dan *F1-score*. Berikut adalah tahapan-tahapan dalam penelitiannya :



Gambar 3.1 Metodologi Penelitian

3.2 Perumusan Masalah

Perumusan masalah yaitu langkah awal dalam metodologi penelitian yang akan dilakukan, tahapan ini merumuskan masalah dan mempelajari masalah yang terjadi serta pada tahapan ini akan ditemukannya latar belakang permasalahan dari penelitian yang dilakukan. Perumusan masalah yang akan dilakukan dalam penelitian ini adalah bagaimana respon sentimen masyarakat terhadap Program Jabar *Quick Response*. Mengetahui dan mengkategorikan tanggapan dari Masyarakat terhadap program tersebut, apakah cenderung menghasilkan sentimen

yang positif, netral ataupun negatif. Melalui analisis ini, diharapkan dapat diperoleh pemahaman yang lebih mendalam mengenai persepsi publik terhadap program tersebut dan faktor-faktor yang mempengaruhi sikap masyarakat. Kemudian menguji bagaimana model dari IndoBERT, MultilingualBERT, dan IndoROBERTa dapat menghasilkan performa dan tingkat akurasi yang terbaik pada teks yang berkaitan dengan pandangan atau opini Masyarakat mengenai program Jabar *Quick Response*.

3.3 Studi Literatur

Studi Literatur merupakan cara yang digunakan untuk menghimpun data atau sumber yang berhubungan dengan judul yang dipakai dalam penelitian ini. Studi literatur dapat diperoleh dengan berbagai sumber antara lain:

1. Buku/*e-Book* yang digunakan untuk membahas tentang *text mining* khususnya *Bidirectional Encoder Representations from Transformers* (BERT).
2. Jurnal yang mengandung tentang klasifikasi dan analisis sentimen menggunakan *Bidirectional Encoder Representations from Transformers* (BERT).

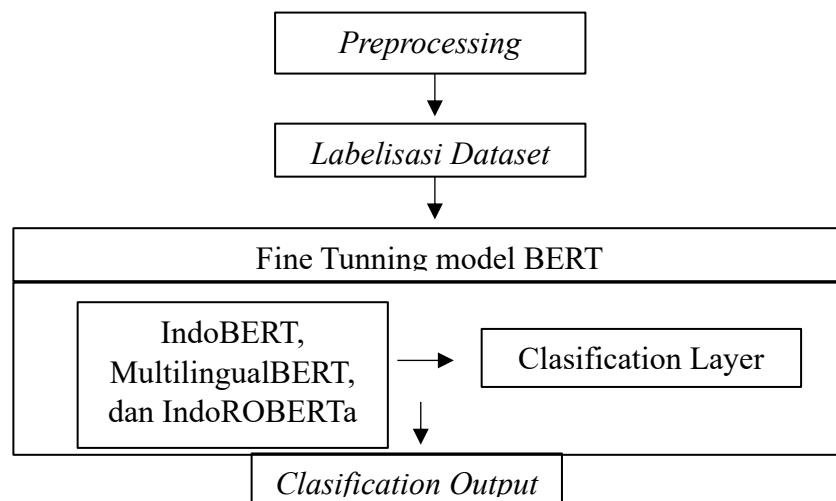
3.4 Pengumpulan Data

Pada penelitian ini dilakukan pengumpulan data dengan menggunakan metode *Crawling* untuk mendapatkan data di sosial media yang mencerminkan opini masyarakat terkait respons sentimen masyarakat terhadap Program Jabar *Quick Response*. Proses *crawling* menggunakan *Tweet Harvest* berfungsi sebagai *Twitter Crawler*, yaitu alat yang dapat mengumpulkan *tweet* berdasarkan kata

kunci, rentang waktu, atau parameter lainnya dengan cara memiliki *API Key* dan *Access Token* agar bisa mengakses data dari X.

3.4 Proses Data dan Perhitungan

Setelah dilakukannya identifikasi masalah, studi literatur, dan pengumpulan data, tahapan selanjutnya adalah analisa. Analisa ini bertujuan untuk mempelajari dan mengevaluasi permasalahan yang ada, dengan tujuan mendapatkan gambaran yang jelas tentang penelitian yang dilakukan. Berikut adalah tahapan proses dan perhitungan data yang dilakukan.



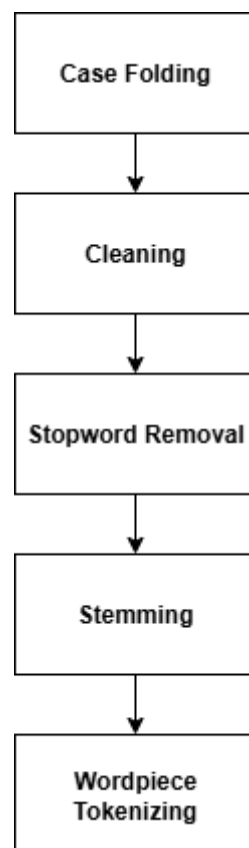
Gambar 3.2 Tahapan Analisa Model BERT

3.4.1 *Preprocessing Dataset*

Preprocessing adalah proses persiapan dan pengolahan teks sebelum dilakukan analisis atau pemodelan. Tujuan dari *pre-pocessing* adalah untuk membersihkan dan mengubah teks mentah menjadi bentuk yang lebih terstruktur dan siap digunakan dalam tahap selanjutnya. *Pre-processing* merupakan langkah yang bertujuan untuk membersihkan data dari unsur-unsur yang tidak dibutuhkan dan dapat digunakan pada proses klasifikasi.

Case Folding, cleaning, stopwords removal, wordpiece tokenizing, stemming merupakan berbagai tahapan *Preprocessing* yang dilakukan pada penelitian ini untuk mengubah *dataset* yang tidak terstruktur menjadi terstruktur dan membuat pengolahan data menjadi lebih sederhana. Selain itu, *Preprocessing* akan meningkatkan hasil analisis sentimen.

Tahapan *preprocessing* yang dilakukan pada penelitian ini :



Gambar 3.3 Tahapan *Preprocessing*

a. *Case folding*

Case folding merupakan proses mengubah semua huruf dalam teks menjadi huruf kecil agar tidak ada perbedaan antara huruf besar dan kecil dalam analisis.

b. *Cleaning*

Cleaning merupakan proses menghilangkan karakter khusus, tanda baca, atau simbol yang tidak relevan atau mengganggu dalam analisis.

c. *Stopword Removal*

Stopword removal merupakan tahapan penghapusan kata-kata umum yang tidak memberikan kontribusi signifikan dalam pemahaman teks, seperti "dan", "atau", "yang", dll.

d. *Stemming*

Stemming adalah proses mengubah variasi bentuk kata ke bentuk kata bakunya semisal menulis menjadi tulis atau pengaturan menjadi atur.

e. *Wordpiece tokenizing*

Wordpiece tokenizing adalah metode tokenisasi teks yang digunakan dalam pemrosesan bahasa alami. Pendekatan ini membagi teks menjadi unit-unit yang lebih kecil yang disebut *wordpieces*. *Wordpieces* adalah potongan-potongan kata yang dapat berupa kata lengkap atau subkata. Metode *wordpiece tokenizing* berguna dalam beberapa konteks, termasuk pemodelan bahasa, mesin terjemahan, dan tugas-tugas pemrosesan bahasa alami lainnya. Keuntungan utamanya adalah kemampuan untuk mengatasi kata yang tidak umum atau kata-kata yang tidak pernah ditemukan dalam kamus

3.4.2 Labelisasi Dataset

Dataset yang sudah memiliki label atau *annotator* diperlukan untuk analisis sentimen menggunakan metode pembelajaran terawasi. Pendekatan pembelajaran terawasi memerlukan contoh, sehingga pelabelan ini harus dilakukan. Inti dari pembelajaran terbimbing adalah pengembangan mekanisme di mana model dapat melihat contoh dan menghasilkan generalisasi sehingga keluaran model adalah prediksi yang sesuai dengan label yang diinginkan. Selain itu, model dapat mengamati dan memahami cara di mana komentar positif, netral, dan negatif berfungsi. Dengan pemberian nilai sebagai penanda, maka pelabelan yang dilakukan bertujuan untuk mengklasifikasikan komentar sebagai positif, netral, atau negatif. Komentar dengan sentimen positif diberi skor 2, komentar negatif diberi skor 0 dan komentar netral diberi skor 1. Teknik labelisasi menggunakan *TextBlob*, yang merupakan pustaka pemrosesan bahasa alami untuk menentukan label kelas sentimen secara otomatis. Namun, terdapat keterbatasan hasil dikarenakan tidak bisa langsung melabelisasi teks yang berbahasa Indonesia.

3.4.3 *Fine-Tuning* Model BERT

Tahapan *Fine-Tuning* pada penelitian ini:

a. Pembagian Data *Training* dan Data *Testing*

Data *Training* dibagi menjadi 70% dan data *Testing* dibagi menjadi 30% setelah data dilakukan proses pelabelan. Selain train-test split penelitian ini menggunakan *Stratified K-Fold Cross Validation* dengan $k=3$ berguna untuk menjaga keseimbangan antara bias dan *variance*.

b. *Fine Tuning Hyperparameters*

Khusus untuk data *Training* dilakukan *fine tuning hyperparameters* dengan menggunakan parameter *epoch*, *learning rate*, dan *batch size*.

c. Model

Model menggunakan *IndoBERT*, *MultilingualBERT*, dan *IndoROBERTa*.

d. *Tools* dan *Framework*

Hugging Face Transformers, *Pytorch*, untuk komputasi analisis sentimennya menggunakan *CUDA (Compute Unified Device Architecture)* dengan memanfaatkan *GPU* agar lebih cepat dalam proses komputasinya.

e. Matriks Evaluasi

1. *Accuracy* memberikan gambaran umum terhadap kinerja model.
2. *Precision* mengevaluasi seberapa tepat model dalam memprediksi suatu kelas.
3. *Recall* mengukur kemampuan model dalam menemukan semua data dari suatu kelas.
4. *F1-Score* memberikan keseimbangan antara *precision* dan *recall*.

3.5 Pengujian

Pada tahap ini, dilakukan pengujian kinerja model untuk 3 kelas sentimen yaitu positif, netral, negatif dengan menghitung tingkat akurasi menggunakan *Confusion Matrix*. Selanjutnya, model diterapkan berdasarkan perancangan skenario eksperimen yang telah dirancang sebelumnya. Seluruh eksperimen akan dilakukan dengan mengikuti metode yang telah direncang sebelumnya, termasuk pemisahan data menjadi *set* pelatihan dan pengujian. Hasil dari eksperimen ini akan memberikan wawasan tentang kemampuan model *IndoBERT*, *MultilingualBERT*,

dan IndoROBERTa dalam melakukan analisis sentimen dengan performa dan Tingkat akurasi yang terbaik tentang kinerja model yang berbeda.