

BAB III

METODOLOGI PENELITIAN

3.1 Hardware dan Software

Penelitian ini memerlukan dukungan perangkat keras (*hardware*) dan perangkat lunak (*software*) untuk menjalankan proses pengolahan data, pelatihan model, serta evaluasi hasil penelitian (Jan dkk., 2022). Spesifikasi perangkat yang digunakan disesuaikan dengan kebutuhan komputasi pada proses *preprocessing* data teks, pelatihan model *deep learning*, serta analisis hasil klasifikasi.

3.1.1 Spesifikasi Hardware

Penelitian ini memerlukan Perangkat dengan spesifikasi laptop Lenovo IdeaPad Slim 3 dengan prosesor Intel Core i5, RAM 8 GB, dan penyimpanan SSD 512 GB. Laptop digunakan untuk pengolahan data, penulisan kode, dan pengelolaan eksperimen.

3.1.2 Spesifikasi Software

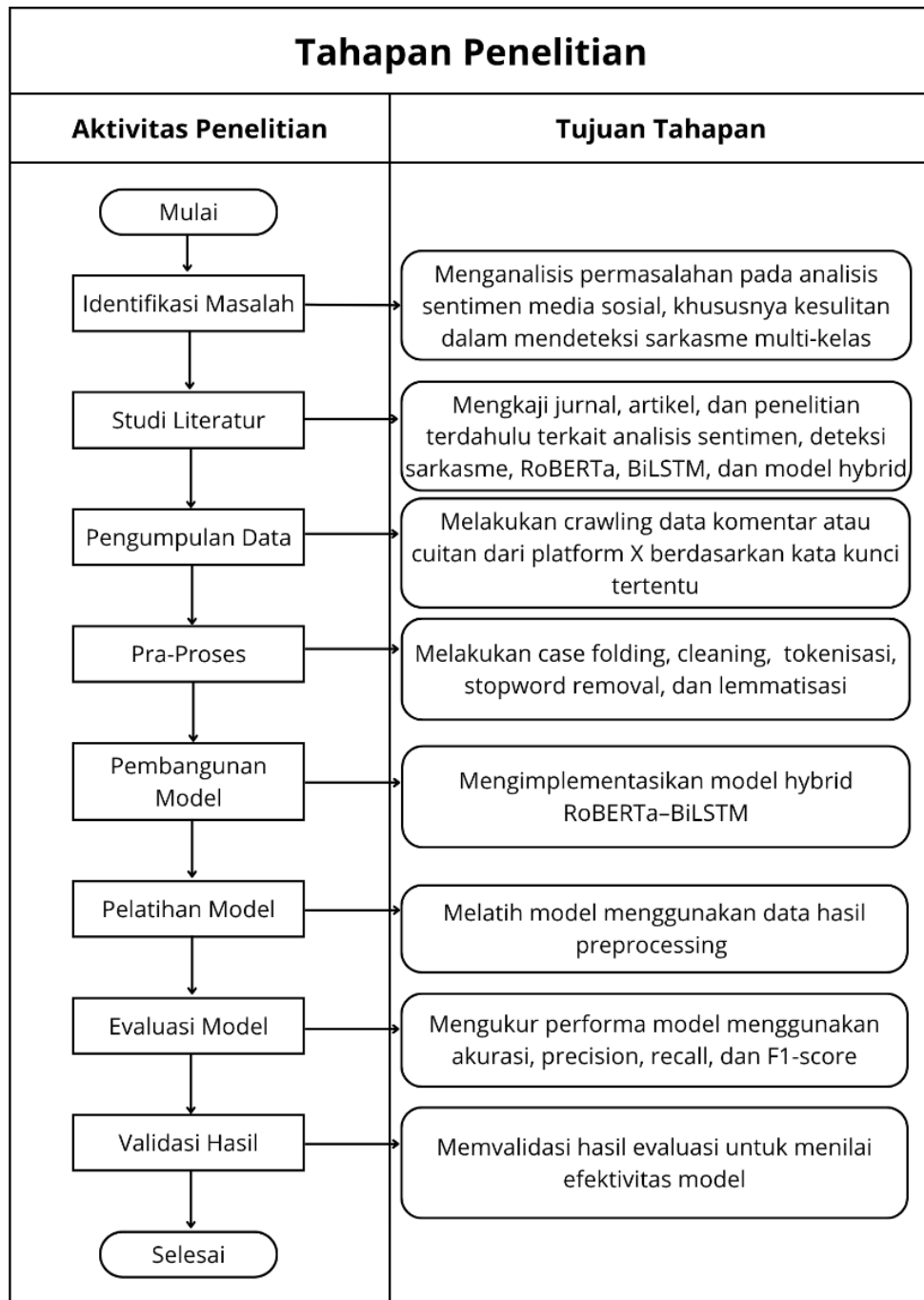
Lingkungan perangkat lunak dikonfigurasi untuk mendukung implementasi pengolahan data teks, *preprocessing* bahasa Indonesia, serta pelatihan model klasifikasi berbasis deep learning. Adapun komponen perangkat lunak yang diperlukan adalah sebagai berikut:

1. Sistem operasi (Windows 11 dari Microsoft sebagai lingkungan kerja utama.)
2. Google *Colaboratory*
3. Bahasa pemrograman (Python 3 sebagai bahasa utama dalam pengembangan dan pemodelan sistem

4. *Framework deep learning (PyTorch (torch) yang dikembangkan oleh Meta Platforms untuk membangun dan melatih model neural network.)*
5. Library NLP dan machine learning (*pandas, re, emoji, Sastrawi, transformers, torch*)

3.2 Tahapan Penelitian

Tahapan penelitian yang akan dilakukan dijelaskan secara rinci dalam Gambar 3.1, mulai dari identifikasi masalah hingga evaluasi. Setiap langkah memiliki tujuan yang jelas. sehingga keseluruhan proses diharapkan dapat memberikan kontribusi terhadap bidang penelitian yang diteliti.



Gambar 3. 1 Tahapan Penelitian

3.2.1 Identifikasi Masalah

Kesalahan klasifikasi sentimen pada teks yang mengandung sarkasme merupakan permasalahan yang sering muncul dalam analisis sentimen (Fitrianto &

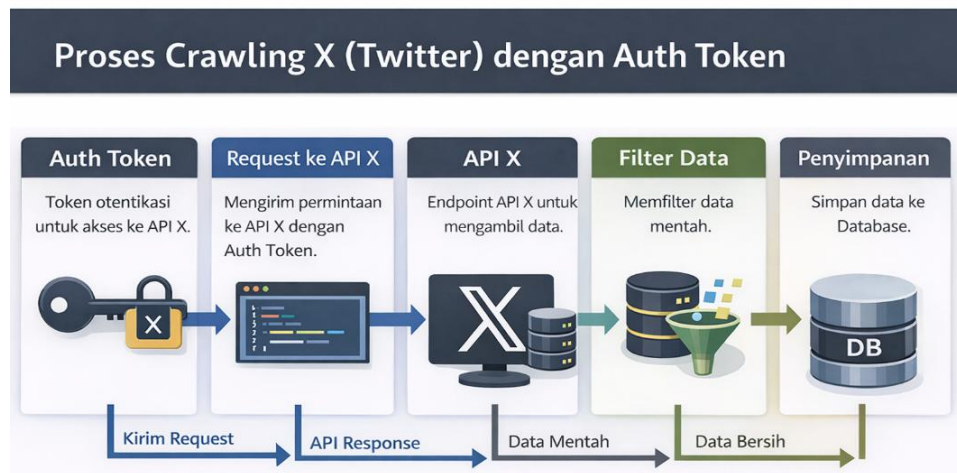
Editya, 2024). Menurut (Cahyo dkk., 2020) Sarkasme menyebabkan ketidaksesuaian antara makna literal kata-kata yang digunakan dengan maksud sebenarnya dari penulis. Sebuah komentar dapat mengandung kata-kata bernada positif, tetapi secara konteks justru menyampaikan sentimen negatif, Permasalahan ini semakin kompleks karena model klasifikasi sentimen umumnya mengandalkan pola kata dan distribusi fitur secara eksplisit, tanpa memahami konteks pragmatik, ironi, atau maksud tersembunyi di balik suatu pernyataan (Warakmulya & Putra, 2025). Akibatnya, komentar sarkastik sering diklasifikasikan sebagai sentimen positif atau netral, padahal makna sebenarnya bersifat negatif. Oleh karena itu, diperlukan pendekatan yang mampu mempertimbangkan konteks kalimat secara menyeluruh dan memahami hubungan antar kata dalam satu pernyataan. Penggunaan model berbasis konteks, penggabungan arsitektur model, serta penambahan proses deteksi sarkasme untuk mengurangi kesalahan klasifikasi sentimen dan meningkatkan akurasi hasil analisis.

3.2.2 Studi Literatur

Studi literatur dilakukan dengan menelaah berbagai publikasi ilmiah yang diperoleh dari platform akademik seperti *ResearchGate*, *Google Scholar*, dan *Academia.edu*. Fokus penelusuran mencakup topik analisis sentimen, deteksi sarkasme, model transformer seperti BERT dan RoBERTa, model sekuensial seperti LSTM dan BiLSTM, serta penerapan arsitektur *hybrid* dalam pengolahan bahasa alami. Tinjauan literatur ini bertujuan untuk mengidentifikasi *State Of The Art (SOTA)*, metode yang telah digunakan, tingkat keberhasilan penelitian sebelumnya, serta celah penelitian yang masih dapat dikembangkan. Hasil studi

literatur ini digunakan sebagai dasar konseptual dalam perancangan metodologi penelitian yang diusulkan.

3.2.3 Pengumpulan Data



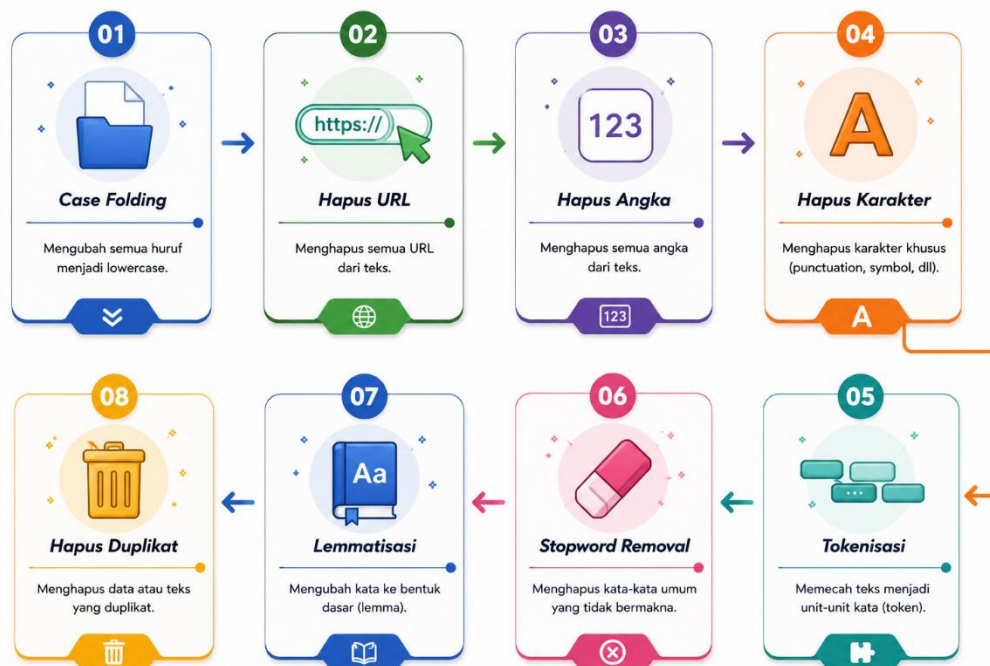
Gambar 3. 2 Pengumpulan Data

Data yang digunakan pada penelitian ini adalah data primer. Data primer adalah data atau informasi yang diperoleh dari sumber yang asli (Undang Sulung & Mohamad Muspawi, 2024). Pengumpulan data dilakukan menggunakan teknik *Crawling* data pada aplikasi X (twitter) menggunakan “*Twitter Auth Token*” seperti pada Gambar 3.3 untuk mengakses komentar, yaitu proses pengambilan data secara otomatis dari platform tersebut dengan memanfaatkan kata kunci yang relevan dengan topik penelitian (Lestari & Febrianti, 2025) kata kunci yang digunakan seperti "makan bergizi gratis", "program makan gratis", "makan siang gratis", "gizi gratis", "makan gratis prabowo", "program prabowo makan gratis", "program makan bergizi", "mbg prabowo", "prabowo makan gratis", "prabowo gizi gratis". Komentar yang dikumpulkan berisi tanggapan, opini, maupun reaksi pengguna terhadap isu yang berkaitan dengan topik “Program Makan Bergizi Gratis”. Proses

pengumpulan data dilakukan dalam rentang waktu 1 Januari 2024 hingga 1 September 2025.

3.2.4 Pra Proses

Pada penelitian ini, tahap pra-proses atau pembersihan data dilakukan untuk menyiapkan data teks hasil *Crawling* dari media sosial X (twitter) agar dapat digunakan secara optimal dalam proses deteksi sarkasme multi-kelas. Data teks dari media sosial umumnya bersifat tidak terstruktur dan mengandung berbagai elemen yang tidak relevan, URL, simbol, serta variasi penulisan kata. Tahapan pra-proses yang dilakukan mengacu pada alur yang ditunjukkan pada Gambar 3.4.



Gambar 3. 3 Pra Proses

Adapun tahapan pra-proses yang dilakukan pada penelitian ini meliputi *Case Folding*, penghapusan URL, penghapusan simbol, tokenisasi, *Stopword Removal*, dan *lemmatization*. *Case Folding* dilakukan dengan mengubah seluruh

huruf pada teks menjadi huruf kecil (Ridwansyah, 2022). Tahapan ini bertujuan untuk menyamakan format teks sehingga perbedaan penggunaan huruf kapital tidak memengaruhi hasil yang berbeda analisis sentimen dan deteksi sarkasme seperti “Program MBG INI BAGUS BANGET ya, RAKYAT pasti makin SEJAHTERA” diubah menjadi “program mbg ini bagus banget ya, rakyat pasti makin sejahtera” . Setelah itu dilakukan penghapusan URL atau tautan yang terdapat dalam teks karena tidak memberikan informasi sentimen yang relevan dan dapat mengganggu proses tokenisasi. Selain URL, simbol seperti *, #, %, \$, dan @ juga dihapus karena tidak memiliki makna semantik dan tidak berkontribusi terhadap proses analisis sentimen pada dataset. Tahap berikutnya adalah tokenisasi, yaitu proses memecah teks menjadi unit-unit kata (token). Pada penelitian ini, tokenisasi dilakukan terhadap komentar media sosial terkait program MBG setelah melalui proses pembersihan data. Sebagai contoh, komentar “Wah MBG hebat banget, rakyat jadi kenyang ya” dipecah menjadi beberapa token seperti “wah”, “mbg”, “hebat”, “banget”, “rakyat”, “jadi”, “kenyang”, dan “ya”. Selanjutnya dilakukan *Stopword Removal* dengan menghapus kata-kata umum yang sering muncul namun tidak memiliki kontribusi signifikan terhadap penentuan sentimen, seperti kata penghubung dan kata depan. Setelah itu dilakukan *lemmatization* atau lemmatisasi, yaitu proses preprocessing teks yang bertujuan untuk mengubah kata ke bentuk dasarnya (lemma) berdasarkan struktur bahasa dan makna kata tersebut (Mutiasari dkk., 2025). Pada tahap analisis sentimen, data teks kemudian diberi label ke dalam tiga kelas, yaitu positif, negatif, dan netral. Kelas positif digunakan untuk teks yang mengandung opini, dukungan, atau emosi positif terhadap suatu topik. Kelas negatif

digunakan untuk teks yang mengandung kritik, penolakan, atau emosi negatif. Sementara itu, kelas netral digunakan untuk teks yang bersifat informatif, deskriptif, atau tidak menunjukkan kecenderungan emosi tertentu.

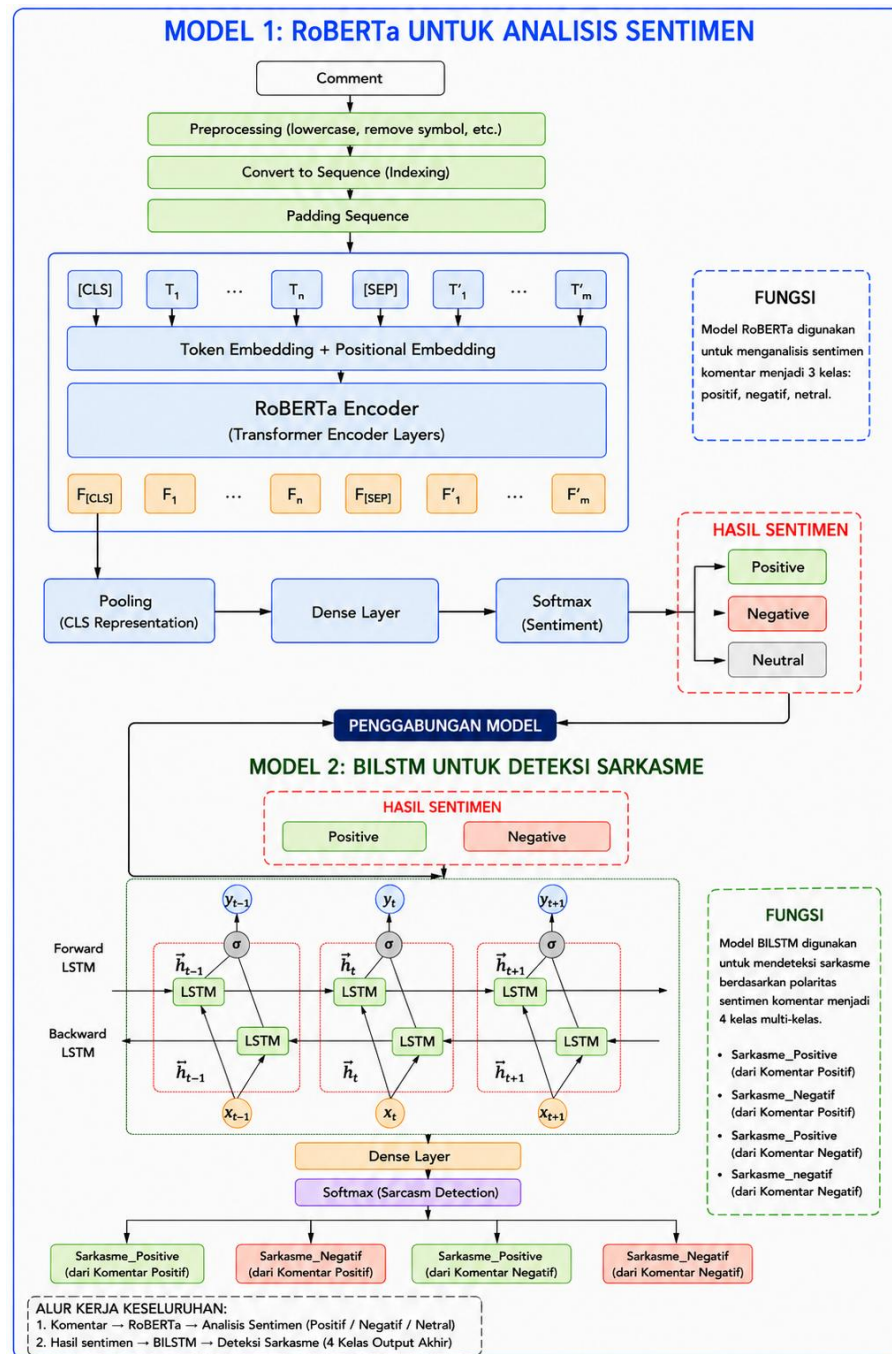
3.2.5 Pelabelan Sarkasme (Dua Kelas)

Berbeda dengan analisis sentimen, pada tugas deteksi sarkasme pelabelan dilakukan ke dalam dua kelas, yaitu sarkasme positif dan sarkasme negatif. Pelabelan ini bertujuan untuk mengidentifikasi komentar yang mengandung makna sindiran atau ironi berdasarkan konteks kalimat yang digunakan. Sarkasme positif digunakan untuk teks yang secara eksplisit tampak positif, namun sebenarnya mengandung makna implisit berupa sindiran atau ironi negatif (Gunawan & Utami, 2025). Pada jenis ini, komentar terlihat seperti pujian atau dukungan, tetapi makna sebenarnya cenderung menyindir suatu kondisi atau keadaan tertentu. Sementara itu, sarkasme negatif digunakan untuk teks yang tampak negatif secara eksplisit, tetapi digunakan secara ironi untuk menyampaikan makna sebaliknya atau sebagai bentuk sindiran tajam (Nadifa Zaelani & Mulyani, 2026). Jenis sarkasme ini umumnya menggunakan ungkapan bernada kritik atau keluhan untuk menyampaikan pesan tertentu secara tidak langsung.

Menurut penelitian yang dilakukan oleh (Liu dkk., 2022) Kelas netral tidak disertakan pada pelabelan sarkasme karena sarkasme secara linguistik selalu mengandung muatan sikap atau evaluasi, sehingga tidak relevan untuk dikategorikan sebagai netral penelitian ini menyatakan bahwa sarkasme melibatkan konflik antara sentimen literal dan implisit karena dalam teks sarkastik, sentimen

yang tersurat (literal) dan yang tersirat (implisit) seringkali berlawanan, misalnya positif vs negatif, bukan netral.

3.2.6 Arsitektur Model



Gambar 3.4 Arsitektur Model

Tahapan pembangunan model pada penelitian ini dilakukan secara bertahap (*two-stage*) sebagaimana ditunjukkan pada Gambar 3.5 arsitektur model. Tahap pertama bertujuan untuk melakukan analisis sentimen, sedangkan tahap kedua difokuskan pada deteksi sarkasme. Pendekatan bertahap ini dilakukan agar model memahami polaritas sentimen terlebih dahulu sebelum mengidentifikasi sarkasme yang bersifat implisit. Proses diawali dengan memasukkan teks komentar media sosial X (twitter) sebagai input ke dalam sistem. Teks komentar kemudian diproses menggunakan RoBERTa *Tokenizer* untuk memecah kalimat menjadi token-token subword yang sesuai dengan format masukan model RoBERTa.

Token hasil tokenisasi selanjutnya dimasukkan ke dalam *pre-trained* RoBERTa yang terdiri dari beberapa Transformer Block. Pada tahap ini, setiap token direpresentasikan dalam bentuk embedding vektor, misalnya E_1 untuk token pertama, E_2 untuk token kedua, hingga E_n untuk token terakhir. Embedding ini diproses menggunakan mekanisme *self-attention*, sehingga RoBERTa mampu mempelajari hubungan antar kata secara kontekstual dalam satu kalimat. Melalui proses encoding tersebut, RoBERTa menghasilkan representasi kontekstual untuk seluruh token, termasuk representasi khusus token [CLS] yang berfungsi sebagai ringkasan informasi dari keseluruhan kalimat. Representasi token [CLS], yang dinyatakan sebagai vektor konteks kalimat C , kemudian diteruskan ke sentiment classifier untuk menentukan polaritas sentimen komentar. Berdasarkan nilai vektor C , komentar diklasifikasikan ke dalam sentimen positif, negatif atau netral, yang merepresentasikan sentimen literal dari teks komentar. Pada tahap ini, RoBERTa

belum secara eksplisit mendeteksi sarkasme, melainkan berfokus pada pemahaman sentimen umum dari komentar.

Hasil dari tahap analisis sentimen selanjutnya digunakan sebagai masukan untuk tahap deteksi sarkasme. Representasi tersebut kemudian diproses oleh *Bidirectional Long Short-Term Memory* (BiLSTM). BiLSTM digunakan karena kemampuannya dalam mempelajari pola urutan kata secara dua arah, yaitu dari awal ke akhir dan dari akhir ke awal kalimat (Sitorus dkk., 2025). Kemampuan ini penting dalam mendeteksi sarkasme, karena sarkasme sering muncul akibat konflik makna yang tersebar di beberapa bagian kalimat (Adelawati & Arimi, 2025). *Output* dari BiLSTM selanjutnya diteruskan ke *sarcasm classifier* untuk mengklasifikasikan komentar ke dalam dua kelas, yaitu sarkasme non-sarkasme dan sarkasme positif atau negatif.

3.2.7 Pelatihan Model

Tahap pelatihan model bertujuan untuk melatih model agar mampu mengenali pola sentimen dan sarkasme secara akurat berdasarkan data yang telah dipersiapkan. Pada tahap ini, dataset yang telah melalui proses pra-pemrosesan dibagi menggunakan rasio 70% sebagai data latih, 20% sebagai data uji, dan 10% sebagai data validasi. Menurut (Brando & Anggai, 2025) pembagian skema ini dilakukan untuk memastikan bahwa model memperoleh data yang cukup dalam proses pembelajaran, sekaligus menyediakan data yang representatif untuk mengevaluasi performa model selama pelatihan dan pengujian. Model RoBERTa, BiLSTM, serta model *hybrid* RoBERTa–BiLSTM dilatih secara terpisah untuk memperoleh hasil perbandingan kinerja yang objektif.

3.2.8 Evaluasi Model

Tahap evaluasi model dilakukan untuk menilai kinerja model RoBERTa, BiLSTM, dan model *hybrid* RoBERTa-BiLSTM dalam melakukan analisis sentimen dan deteksi sarkasme multi-kelas pada data media sosial. Kinerja model diukur menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-Score* (Sujon dkk., 2025). Hasil evaluasi kemudian dibandingkan untuk menentukan model dengan performa terbaik serta menilai efektivitas penggunaan model *hybrid* dalam meningkatkan ketepatan dan keandalan deteksi sarkasme multi-kelas.

3.2.9 Validasi Hasil

Tahap validasi hasil dilakukan dengan meninjau kembali hasil klasifikasi model terhadap beberapa contoh komentar media sosial untuk memastikan kesesuaian antara prediksi model dan makna sentimen yang sebenarnya. Sebagai contoh, komentar “*MBG sangat bermanfaat, terima kasih*” divalidasi untuk memastikan model mengklasifikasikannya sebagai sentimen positif secara tepat. Proses validasi ini didukung oleh hasil pengukuran kinerja model berupa nilai probabilitas yang diperoleh dari data uji.