

BAB I

PENDAHULUAN

1.1 Latar Belakang

Perkembangan teknologi informasi dan komunikasi telah mendorong pertumbuhan signifikan pada industri *e-commerce* di Indonesia dalam beberapa tahun terakhir. Beragam platform, seperti Shopee, Tokopedia, dan Lazada, semakin banyak dimanfaatkan masyarakat karena menyediakan proses transaksi yang praktis, cepat, serta menawarkan akses yang luas terhadap berbagai jenis produk dan layanan. Meningkatnya aktivitas pengguna pada platform tersebut menghasilkan volume interaksi digital yang sangat besar, salah satunya berupa ulasan aplikasi yang dipublikasikan melalui Google Play Store. Ulasan tersebut memuat beragam informasi, mulai dari pengalaman penggunaan, pendapat, kritik, saran, hingga tingkat kepuasan terhadap layanan yang diberikan. Tingginya adopsi platform digital menjadikan ulasan pengguna di Google Play Store sebagai sumber data bernilai tinggi untuk mengevaluasi kualitas layanan dan memahami persepsi konsumen secara sistematis (Nurian et al., 2023).

Ulasan yang diberikan pengguna memiliki peranan penting sebagai bahan evaluasi bagi perusahaan dalam memahami kebutuhan serta harapan pelanggan terhadap layanan yang disediakan. Informasi tersebut memungkinkan pengembang aplikasi mengenali berbagai kendala yang dialami pengguna, menilai efektivitas fitur yang telah tersedia, serta menentukan aspek yang perlu diperbaiki untuk meningkatkan kualitas layanan. Pengguna cenderung menyampaikan pengalaman mereka secara spontan melalui gaya bahasa informal, singkatan, variasi kosakata tidak baku, hingga penggunaan emoji sebagai penegas ekspresi emosional. Namun, eksponensialnya volume ulasan harian menyebabkan proses analisis secara manual menjadi tidak efisien, membutuhkan alokasi waktu yang besar, serta rentan terhadap bias subjektivitas (Putri et al., 2025).

Salah satu pendekatan yang banyak diterapkan untuk menganalisis opini pengguna adalah analisis sentimen. Untuk mengatasi tantangan tersebut, analisis sentimen hadir sebagai cabang *Natural Language Processing* (NLP) yang berfokus pada identifikasi, ekstraksi, dan klasifikasi emosi ke dalam kategori sentimen seperti positif, netral, dan negatif (Mao et al., 2024). Informasi dari analisis ini

menjadi acuan strategis bagi perusahaan dalam melakukan evaluasi layanan, penyempurnaan fitur aplikasi, serta mendukung pengambilan keputusan bisnis yang efektif (Tania Puspa Rahayu Sanjaya et al., 2023).

Dalam praktiknya, sebagian besar penelitian terdahulu masih mengandalkan metode klasifikasi tradisional seperti Naive Bayes karena efisiensi komputasinya yang tinggi dan proses pelatihan yang sederhana (Safira et al., 2023). Pendekatan ini umumnya dikombinasikan dengan pembobotan fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengidentifikasi kata-kata penting dalam dokumen (Adi Widiyanto et al., 2025). Meskipun efektif pada kasus umum, Naive Bayes mengasumsikan independensi antar-fitur (*bag-of-words*), sehingga gagal memperhitungkan keterkaitan konteks semantik, struktur sintaksis, konstruksi kalimat negasi, maupun variasi leksikal informal.

Walaupun telah menunjukkan kinerja yang cukup baik pada berbagai penelitian, pendekatan yang mengandalkan frekuensi kata masih memiliki keterbatasan dalam memahami makna serta konteks bahasa. Hal tersebut disebabkan oleh asumsi dasar Naive Bayes yang menganggap setiap fitur bersifat independen, sehingga hubungan antar kata di dalam sebuah kalimat tidak diperhitungkan secara langsung (Agus Hermanto, 2021). Akibatnya, model sering mengalami kesulitan untuk memahami konteks kalimat yang mengandung hubungan semantik antar kata, negasi, ambiguitas, dan ironi. Sebagai contoh, kalimat “aplikasinya tidak buruk” dan “aplikasinya buruk” memiliki makna sentimen yang berbeda, tetapi model berbasis frekuensi kata cenderung memberikan perhatian yang sama terhadap kata “buruk”. Keterbatasan tersebut menjadi semakin kompleks ketika diterapkan pada ulasan pengguna yang bersifat informal dan tidak mengikuti kaidah bahasa baku.

Selain Naive Bayes, berbagai algoritma klasifikasi juga telah banyak diterapkan dalam penelitian analisis sentimen, di antaranya *Support Vector Machine* (SVM), *Random Forest*, *Long Short-Term Memory* (LSTM), serta Multilingual BERT. Masing-masing metode memiliki karakteristik dan keunggulan yang berbeda dalam mengolah serta mengklasifikasikan data teks. Namun, penelitian ini memilih untuk menggunakan IndoBERT karena model tersebut dikembangkan dan dilatih secara khusus. Model ini dikembangkan menggunakan korpus bahasa Indonesia sehingga mampu memahami karakteristik linguistik serta konteks penggunaan bahasa

Indonesia dengan lebih baik, termasuk penggunaan ulasan pengguna menunjukkan bahasa informal, singkatan, variasi kosakata, dan konteks kalimat yang umum. Namun, Naive Bayes Sebagai model pembandingan (*baseline*), Naive Bayes dipilih karena menawarkan proses klasifikasi yang sederhana, efisien, dan ringan dari sisi komputasi. Penggunaannya yang masih luas dalam penelitian analisis sentimen juga menjadikannya sebagai acuan yang tepat untuk membandingkan performa model berbasis *Transformer*. Metode lain seperti SVM, Random Forest, LSTM, maupun BERT Multilingual tidak dipilih karena penelitian ini bertujuan membandingkan dua pendekatan yang memiliki karakteristik berbeda, yaitu metode klasifikasi tradisional berbasis probabilistik dan model *Transformer* yang menggunakan representasi kontekstual. Dengan demikian, pengaruh kemampuan pemahaman konteks bahasa terhadap hasil klasifikasi sentimen dapat dianalisis secara lebih jelas (Wirayudha et al., 2025).

Berbagai penelitian sebelumnya menunjukkan bahwa model berbasis *Transformer*, khususnya BERT dan berbagai variannya, secara konsisten menghasilkan nilai akurasi dan *F1-score* yang lebih tinggi dibandingkan metode *machine learning* tradisional pada berbagai tugas analisis sentimen. Keunggulan tersebut diperoleh karena model *Transformer* mampu membentuk representasi kata berdasarkan konteks penggunaannya dalam kalimat (Devlin et al., 2020; Wilie et al., 2020; Maxmiliano et al., 2025). Meskipun kinerjanya secara umum di bawah model *Transformer*, Naive Bayes tetap relevan digunakan sebagai *baseline* karena efisiensinya secara komputasi serta rekam jejak penggunaannya yang luas pada penelitian klasifikasi teks (Safira et al., 2023). Berdasarkan uraian tersebut, penelitian ini membandingkan kinerja IndoBERT dan Naive Bayes dalam melakukan analisis sentimen terhadap ulasan aplikasi *e-commerce* berbahasa Indonesia guna mengetahui model yang memberikan performa terbaik yang dihasilkan oleh representasi kontekstual dapat dianalisis secara objektif tanpa melibatkan terlalu banyak variasi metode klasifikasi lainnya (Wilie et al., 2020; Maxmiliano et al., 2025).

Karakteristik ulasan pengguna pada platform digital umumnya berbeda dengan teks formal karena sering mengandung singkatan, bahasa gaul, kesalahan penulisan, campuran bahasa, serta ekspresi emosional yang beragam (Muzaki et al., 2024).

Bentuk komunikasi seperti ini banyak ditemukan pada ulasan aplikasi *e-commerce* di Google Play Store karena pengguna cenderung menuliskan pengalaman mereka secara spontan sesuai kondisi yang dialami. Variasi bahasa tersebut menyebabkan proses pemahaman makna menjadi lebih sulit apabila hanya mengandalkan frekuensi kemunculan kata. Dengan demikian, proses identifikasi sentimen pada ulasan memerlukan metode yang mampu menangkap konteks bahasa secara lebih baik sehingga makna yang disampaikan pengguna dapat dikenali secara lebih tepat (Averina et al., 2022).

Selain penggunaan bahasa informal, pengguna juga sering memanfaatkan emoji sebagai sarana untuk mengekspresikan emosi dalam komunikasi digital. Emoji tidak hanya berfungsi sebagai pelengkap teks, tetapi juga dapat memperkuat, memperjelas, bahkan mengubah makna sentimen yang terkandung dalam suatu kalimat. Sebagai contoh, kalimat yang sama dapat memiliki interpretasi sentimen yang berbeda ketika disertai emoji tersenyum atau emoji marah. Studi yang telah dilakukan sebelumnya oleh (Simanungkalit et al., 2024) mengungkapkan bahwa penggunaan teks yang disertai emoji merupakan cara yang umum dilakukan pengguna untuk memperkuat penyampaian sentimen pada ulasan aplikasi. Namun, sebagian penelitian sebelumnya masih menghapus emoji pada tahap pre-pemrosesan data. Kondisi tersebut menyebabkan informasi emosional yang terdapat pada emoji tidak dimanfaatkan secara optimal. Selain itu, menghapus emoji dapat menghilangkan informasi yang berperan dalam membantu model memahami sentimen pengguna dengan lebih baik.

Dengan kemajuan teknologi pengolahan bahasa natural (NLP), ada banyak. Dibandingkan pendekatan tradisional, metode berbasis *deep learning* lebih efektif dalam menangkap konteks bahasa yang terdapat pada suatu teks. Salah satu perkembangan penting dalam bidang NLP adalah hadirnya arsitektur *Transformer* yang diperkenalkan oleh (Vaswani, 2020). Arsitektur ini Dengan memanfaatkan mekanisme *self-attention*, model mampu mengenali keterkaitan antar kata dalam kalimat sehingga pemahaman terhadap konteks menjadi lebih efektif. Kemampuan tersebut menjadikan Transformer sebagai fondasi bagi berbagai model bahasa modern yang banyak digunakan dalam berbagai fungsi proses pemrosesan bahasa natural, salah satunya adalah analisis sentimen.

Salah satu model berbasis Transformer yang banyak digunakan adalah *Bidirectional Encoder Representations from Transformers* (BERT), yang dikembangkan oleh (Devlin et al., 2020). BERT memiliki kemampuan untuk memahami konteks kalimat dari dua arah (*bidirectional context*), yaitu dengan mempertimbangkan informasi leksikal sebelum dan sesudah suatu kata secara simultan. Pendekatan tersebut menghasilkan representasi kata yang lebih kontekstual sehingga model dapat memahami makna setiap kata berdasarkan keseluruhan isi kalimat. Kemampuan ini membuat BERT lebih efektif dalam menangani berbagai permasalahan yang sulit diselesaikan oleh metode yang hanya mengandalkan frekuensi kemunculan kata.

IndoBERT merupakan model yang dikembangkan menggunakan korpus multibahasa (Baldwin, 2020). Berkat proses pelatihannya, model ini memiliki kemampuan yang baik dalam memahami karakteristik bahasa Indonesia, termasuk penggunaan bahasa tidak formal, singkatan, serta berbagai variasi ekspresi yang umum dijumpai pada ulasan pengguna di platform digital. Kemampuan ekstraksi konteks secara komprehensif tersebut menjadikan IndoBERT—sebagai varian BERT yang telah dilatih pada korpus bahasa Indonesia—sangat ideal untuk diterapkan pada tugas analisis sentimen ulasan aplikasi berbahasa lokal.

Sejumlah penelitian sebelumnya melaporkan bahwa IndoBERT memberikan hasil yang kompetitif dalam berbagai tugas klasifikasi teks berbahasa Indonesia (Simanungkalit et al., 2024; Mursidah et al., 2025). Di sisi lain, Naive Bayes masih digunakan secara luas sebagai metode pembandingan (*baseline*) karena proses klasifikasinya yang sederhana, mudah diimplementasikan, serta membutuhkan sumber daya komputasi yang relatif rendah (Safira et al., 2023; Muzaki et al., 2024). Penelitian yang dilakukan oleh (Maxmiliano et al., 2025). menunjukkan bahwa model BERT memperoleh nilai *accuracy* dan *F1-score* yang lebih tinggi dibandingkan Naive Bayes dalam klasifikasi sentimen ulasan *e-commerce*. Namun, sebagian besar penelitian terdahulu masih terbatas pada aspek tertentu, menerapkan pendekatan prapemrosesan yang seragam untuk arsitektur model yang berbeda karakteristik, serta belum mempertimbangkan keberadaan emoji secara eksplisit sebagai representasi penting dari fitur sentimen. Selain itu, penelitian yang mengevaluasi dan membandingkan kemampuan Naive Bayes serta IndoBERT

dalam menganalisis sentimen multi-kelas (positif, netral, dan negatif) pada ulasan aplikasi *e-commerce* berbahasa Indonesia masih sangat terbatas.

Penelitian ini bertujuan untuk menganalisis sentimen ulasan aplikasi *e-commerce* berbahasa Indonesia menggunakan model IndoBERT serta membandingkan kinerjanya terhadap algoritma Naive Bayes sebagai model *baseline*. Dengan memanfaatkan data ulasan pengguna yang diperoleh dari Google Play Store, penelitian ini berfokus pada klasifikasi sentimen multikelas yang mencakup kategori positif, netral, dan negatif. Pemilihan kedua pendekatan ini didasari oleh perbedaan fundamental pada paradigma pembentukan fiturnya; Naive Bayes mewakili metode probabilistik tradisional berbasis frekuensi kata, sedangkan IndoBERT mewakili arsitektur *Transformer* berbasis konteks semantik dua arah. Penggunaan Naive Bayes sebagai model dasar difungsikan untuk menyediakan tolok ukur yang objektif dalam mengevaluasi sejauh mana peningkatan kinerja yang mampu dicapai oleh pemodelan berbasis *Transformer* (Jayadianti et al., 2022).

Sejumlah penelitian terdahulu, seperti yang dilakukan oleh (Safira et al., 2023); (Nurian et al., 2023); serta (Simanungkalit et al., 2024) masih mengandalkan algoritma *machine learning* konvensional—khususnya Naive Bayes—dalam menganalisis sentimen ulasan aplikasi. Pendekatan tersebut umumnya terfokus pada pemodelan berbasis frekuensi kata atau probabilitas sederhana, tanpa mengeksplorasi pemahaman konteks semantik secara mendalam sebagaimana yang ditawarkan oleh arsitektur *Transformer*. Sebaliknya, penelitian (Kireyna Cindy Pradhisa, 2024); (Putra, A., et al., 2025) dan (Mursidah et al., 2025) telah menerapkan IndoBERT sebagai model klasifikasi sentimen, namun penelitian tersebut belum melakukan perbandingan langsung dengan metode *machine learning* klasik menggunakan dataset yang sama. Sementara itu, penelitian (Maxmiliano et al., 2025) telah membandingkan Naive Bayes dan BERT, tetapi hanya dilakukan pada ulasan produk Shopee dengan fokus pada atribut produk. Akibatnya, penelitian ini dilakukan untuk mengevaluasi efektivitas Naive Bayes dan IndoBERT pada dataset gabungan ulasan Tokopedia, Shopee, dan Lazada yang didominasi teks informal, kedua pendekatan dalam analisis sentimen berbahasa Indonesia.

1.2 Rumusan Masalah

Berdasarkan latar belakang masalah yang telah diuraikan, rumusan masalah dalam penelitian ini disusun sebagai berikut:

1. Bagaimana proses prapemrosesan dan pengolahan teks informal yang menyertakan emoji pada ulasan *e-commerce* berbahasa Indonesia menggunakan arsitektur *Transformer* IndoBERT?
2. Bagaimana perbandingan performa antara IndoBERT dan Naive Bayes sebagai model baseline dalam kategori ulasan *e-commerce* berbahasa Indonesia berdasarkan perasaan?
3. Bagaimana kemampuan IndoBERT dan Naive Bayes dalam mengkategorikan perasaan dalam ulasan *e-commerce* berbahasa Indonesia yang mengandung karakteristik bahasa informal?

1.3 Manfaat Penelitian

Hasil penelitian ini diharapkan dapat memberikan kontribusi bermakna, baik secara teoritis maupun praktis, bagi berbagai pihak sebagai berikut:

1. Bagi Penulis
Untuk meningkatkan pemahaman dan pengalaman penulis tentang penerapan *Natural Language Processing* (NLP), khususnya analisis sentimen menggunakan arsitektur *Transformer* IndoBERT dan metode Naive Bayes pada teks Bahasa Indonesia.
2. Bagi Peneliti Selanjutnya
Menjadi landasan rujukan dan pembandingan untuk pengembangan riset analisis sentimen lanjutan, khususnya dalam menerapkan teknik penanganan data tidak seimbang (*class weighting*, *Focal Loss*) atau pengayaan data (*text augmentation*).

1.4 Batasan Masalah

Agar penelitian ini tetap berfokus pada ruang lingkup yang terukur dan mencapai sasaran secara substantif, batasan masalah yang diterapkan dirumuskan sebagai berikut:

1. Objek Penelitian: Data yang digunakan adalah ulasan teks pengguna dari aplikasi *e-commerce* (Tokopedia, Shopee, Lazada) yang dapat diakses melalui Google Play Store. Ulasan yang dianalisis berupa teks Bahasa Indonesia, dengan mempertahankan elemen emoji sebagai fitur intrinsik pembawa sentimen, bukan dibuang sebagai *noise*.
2. Fokus Analisis: Penelitian ini berfokus pada klasifikasi sentimen, (*Sentiment Classification*) ada tiga jenis: Positif, Negatif, dan Netral.
3. Model Yang Digunakan: Model utama yang digunakan adalah arsitektur Transformer IndoBERT. Sebagai model baseline, digunakan metode klasifikasi Naive Bayes berdasarkan representasi fitur TF-IDF.
4. Ruang Lingkup Bahasa: Ulasan yang dianalisis terbatas pada Bahasa Indonesia, termasuk penggunaan bahasa yang tidak formal seperti slang, singkatan, kesalahan ketik, dan variasi bahasa yang umum ditemukan dalam ulasan daring.
5. Evaluasi Model: *Confusion matrix* dengan *accuracy*, *precision*, *recall*, dan *F1-Score* digunakan untuk mengevaluasi performa model.

1.5 Tujuan Penelitian

Sejalan dengan rumusan masalah yang telah dipaparkan, penelitian ini dilaksanakan dengan tujuan sebagai berikut:

1. Menganalisis dan mengevaluasi performa arsitektur Transformer IndoBERT dalam melakukan klasifikasi untuk teks ulasan *e-commerce* berbahasa Indonesia yang informal.
2. Menganalisis dan membandingkan kinerja arsitektur *Transformer* IndoBERT terhadap algoritma *baseline* Naive Bayes dalam mengklasifikasikan sentimen multikelas (positif, netral, dan negatif) pada ulasan *e-commerce* berbahasa Indonesia, berdasarkan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*.
3. Menganalisis kemampuan kedua model dalam menangani karakteristik teks informal yang mengandung singkatan, bahasa tidak baku, dan emoji.

1.6 Metodologi Penelitian

Penelitian ini menerapkan pendekatan kuantitatif dengan metode eksperimen (*experimental research*) untuk mengevaluasi dan membandingkan kinerja model klasifikasi sentimen pada korpus ulasan aplikasi *e-commerce* berbahasa Indonesia yang berkarakteristik informal. Arsitektur *Transformer* IndoBERT dimanfaatkan sebagai model utama (*proposed model*), sedangkan algoritma Naive Bayes berbasis pembobotan *Term Frequency–Inverse Document Frequency* (TF-IDF) difungsikan sebagai model dasar (*baseline*). Evaluasi dilakukan secara multikelas (positif, netral, dan negatif) dengan mempertahankan eksistensi emoji sebagai indikator pemaknaan sentimen yang esensial. Berdasarkan eksperimen dan pengujian komparatif yang telah dieksekusi, hasil utama penelitian dirangkum sebagai berikut:

1. Pengumpulan Data (*Data Collection*)

Data penelitian ini berasal dari ulasan pengguna dari aplikasi *e-commerce* berbahasa Indonesia yang dikumpulkan melalui teknik web scraping pada platform Google Play Store. Aplikasi ini termasuk Tokopedia, Shopee, dan Lazada, yang memiliki ulasan berbahasa Indonesia, termasuk emoji. Sebelum memasuki tahap pra-pemrosesan, data yang telah dikumpulkan disiapkan untuk pelabelan sentimen.

2. Pelabelan Sentimen (*Sentiment Labeling*)

Pelabelan sentimen dilakukan secara otomatis berdasarkan nilai rating—juga dikenal sebagai rating bintang yang diberikan pengguna di platform Google Play Store. Pembagian kelas sentimen dibagi menjadi tiga kelompok multikelas. Ulasan dengan skor 4 hingga 5 dianggap positif, skor 3 netral, dan skor 1 hingga 2 dianggap negatif. Hasil pelabelan ini digunakan sebagai target klasifikasi untuk metode pembelajaran yang diawasi.

3. Tahap Prapemrosesan Data (*Data Preprocessing*)

Sebelum dimasukkan ke dalam proses pelatihan model, tahap prapemrosesan data bertujuan untuk membersihkan, mengubah, dan mengondisikan data ulasan mentah untuk mencapai tingkat kualitas dan konsistensi terbaik. Pembersihan ini sangat penting dalam pemrosesan bahasa alami (NLP) untuk mengurangi kebisingan seperti penulisan

karakter yang tidak relevan, typo, dan aksentuasi berlebihan tanpa menghilangkan komponen yang membawa sentuhan seperti emoji.

Proses yang dilakukan termasuk:

- a. *Cleaning*, yaitu menghapus karakter khusus, URL, tag HTML, dan elemen lain yang tidak relevan, sementara emoji tetap dipertahankan.
- b. *Case Folding*, yaitu mengubah seluruh huruf menjadi huruf kecil (*lowercase*).
- c. *Tokenization*, yaitu memecah teks menjadi token-token yang akan diproses oleh masing-masing model.
- d. *Stopword Removal*, Tahap *Stopword Removal* (penghapusan kata-kata umum yang dianggap tidak membawa muatan sentimen signifikan) secara khusus hanya diterapkan pada pemodelan Naive Bayes. Kebijakan ini didasarkan pada perbedaan karakteristik fundamental kedua arsitektur; model *Transformer* seperti IndoBERT membutuhkan keutuhan struktur kalimat, urutan kata, dan keterkaitan sintaksis untuk mengoptimalkan mekanisme *self-attention*. Penghapusan *stopwords* pada IndoBERT justru berpotensi merusak representasi kontekstual alami dari teks informal. Sedangkan pada IndoBERT tidak dilakukan agar konteks kalimat tetap terjaga.

4. Representasi Fitur (*Feature Representation*)

Representasi fitur dilakukan sesuai karakteristik masing-masing model.

- a. Pada model Naive Bayes Selanjutnya, untuk mentransformasi teks ulasan menjadi masukan numerik yang dapat diproses oleh algoritma Naive Bayes, digunakan metode *Term Frequency–Inverse Document Frequency* (TF-IDF). Metode ini menghitung bobot relevansi setiap kata berdasarkan frekuensi kemunculannya dalam sebuah ulasan relatif terhadap seluruh korpus dokumen.
- b. Pada model IndoBERT, teks diproses menggunakan WordPiece Tokenizer untuk menghasilkan model masukan ke model Transformer, identitas input dan penghalang perhatian.

5. Eksperimen Model (*Model Experimentation*)

Agar proporsi, dataset dibagi menjadi 70% data latih dan 30% data uji menggunakan metode *stratified sampling* agar proporsi setiap kelas sentimen tetap terjaga.

6. Evaluasi Model (*Model Evolution*)

Data uji independen, yang belum pernah digunakan selama proses pelatihan dan validasi, digunakan untuk mengevaluasi kinerja kedua model secara terpisah. Empat metrik utama, keakuratan, recall, dan skor F1, digunakan untuk menghitung evaluasi kuantitatif berdasarkan pembentukan matrix confusion. Selanjutnya, hasil evaluasi dari kedua pemodelan dievaluasi dan dibandingkan untuk menentukan arsitektur mana yang dapat mengklasifikasikan sentimen dengan paling andal, tepat, dan umum untuk ulasan e-commerce berbahasa Indonesia.

1.7 Sistematika Penulisan

Guna memberikan gambaran yang sistematis mengenai alur pemikiran dan penyusunan penelitian ini, sistematika penulisan dibagi ke dalam lima bab utama sebagai berikut:

BAB I PENDAHULUAN

Membahas latar belakang masalah analisis sentimen ulasan aplikasi *e-commerce*, rumusan masalah, batasan masalah, tujuan dan manfaat penelitian, metodologi penelitian, serta sistematika penulisan.

BAB II TINJAUAN PUSTAKA

Membahas teori-teori pendukung seperti *Natural Language Processing* (NLP), arsitektur *Transformer*, IndoBERT, Naive Bayes, TF-IDF, serta penelitian terdahulu sebagai landasan penelitian.

BAB III METODOLOGI PENELITIAN

Menjelaskan tahapan penelitian yang meliputi pengumpulan data, prapemrosesan data, representasi fitur, perancangan eksperimen model, pembagian dataset, evaluasi model, serta analisis kebutuhan sistem.

BAB IV HASIL DAN PEMBAHASAN

Menyajikan hasil eksperimen dan diskusi mengenai performa model Naive Bayes sebagai *baseline* dan IndoBERT sebagai model utama dalam melakukan klasifikasi sentimen ulasan aplikasi *e-commerce* berbahasa Indonesia.

BAB V PENUTUP

Menyajikan kesimpulan dari penelitian analisis sentimen ulasan aplikasi *e-commerce* menggunakan model Naive Bayes dan IndoBERT, serta saran yang dapat digunakan sebagai acuan untuk penelitian yang akan datang