

BAB II

KAJIAN LITERATUR

2.1 Tinjauan Pustaka

Pada bagian ini berisikan landasan teori dan akan dijelaskan tinjauan pustaka yang berkaitan dengan penelitian yang dilakukan.

2.1.1 ABSA (*Aspect-Based Sentiment Analysis*)

Analisis sentimen merupakan salah satu cabang dari pemrosesan bahasa alami NLP yang berfokus pada bagaimana komputer memahami respons manusia melalui teks. Tujuannya adalah untuk mengenali emosi atau sikap yang terkandung dalam sebuah pernyataan, apakah bersifat positif, negatif, atau netral. Menurut Sembiring dan Dewa [16], analisis ini mampu membantu mengubah data teks mentah menjadi informasi yang lebih bermakna, misalnya untuk mengetahui kecenderungan respons publik terhadap isu sosial atau kebijakan pemerintah. Dengan begitu, analisis sentimen tidak hanya memetakan arah respons masyarakat, tetapi juga menjadi dasar bagi pengambilan keputusan berbasis data [14].

Seiring berkembangnya penelitian di bidang NLP, muncul pendekatan ABSA yang berfungsi untuk menggali sentimen secara lebih mendalam. Meskipun ABSA menawarkan analisis sentimen yang lebih granular, penerapannya pada teks media sosial menghadapi tantangan metodologis yang signifikan, seperti penggunaan bahasa tidak baku, opini implisit, ironi, serta ketidakseimbangan distribusi kelas pada level aspek. Tantangan ini menuntut pendekatan ABSA yang mampu menangkap representasi kontekstual dan hubungan sekuensial secara terpadu untuk memahami evaluasi publik terhadap kebijakan yang bersifat kompleks dan multidimensional. Berbeda dengan analisis sentimen biasa, pendekatan ini tidak hanya menilai arah emosinya, namun mencari tahu aspek apa yang menjadi pusat pembahasan [17]. Cahyaningtyas et al. menjelaskan bahwa metode ini menghubungkan antara aspek dan sentimen dalam satu kalimat untuk memberikan pemahaman yang lebih rinci terhadap isi respons [18]. Dengan begitu, pendekatan ini memungkinkan peneliti mengetahui aspek mana yang mendapatkan tanggapan positif dan aspek mana yang justru menuai kritik. Dalam opini publik terhadap kebijakan yang bersifat *multidimensional*, pendekatan ABSA tidak lagi sekadar menjadi alternatif metodologis, melainkan kebutuhan analitis untuk menangkap kompleksitas hubungan antara aspek kebijakan dan sentimen masyarakat. Selain itu, aspek yang menjadi objek evaluasi sering

kali tidak disebutkan secara eksplisit, sehingga menuntut pendekatan ABSA yang mampu menangkap semantik lokal dan global [14].

Secara teoretis, ABSA dibangun atas dua tahapan utama yang saling berkaitan, yaitu mengidentifikasi aspek-aspek apa saja yang dibicarakan dalam teks (*aspect extraction*) dan menentukan sentimen terhadap setiap aspek yang telah diidentifikasi (*aspect-level sentiment classification*). Kedua tahapan ini bekerja untuk menghasilkan pemahaman komprehensif tentang struktur opini dalam teks. Pada tahapan *aspect extraction*, aspek sendiri dapat didefinisikan sebagai entitas, atribut, atau topik tertentu yang menjadi subjek evaluasi atau kritik. Dalam literatur, terdapat beberapa pendekatan yang digunakan untuk melakukan ekstraksi aspek. Pendekatan *rule-based* menggunakan pola linguistik, seperti *part-of-speech tagging* dan *dependency parsing* untuk mengidentifikasi aspek berdasarkan struktur gramatikal kalimat. Pendekatan ini relatif sederhana dan *interpretable*, namun terbatas pada domain dengan pola bahasa yang konsisten dan *predictable*. Pendekatan berbasis *machine learning* menggunakan algoritma, seperti *Conditional Random Fields* (CRF), *Hidden Markov Model* (HMM), atau model *deep learning* seperti LSTM dan BERT untuk mendeteksi aspek dari data yang telah berlabel [14], [19]. Keunggulan pendekatan ini adalah fleksibilitasnya, namun memerlukan data pelatihan yang cukup besar dan berkualitas tinggi. Pendekatan yang ketiga adalah berbasis *topic modeling*, khususnya menggunakan LDA, yang dapat mengidentifikasi aspek-aspek laten dari kumpulan dokumen tanpa memerlukan data berlabel. Pendekatan ini sangat berguna dalam fase eksplorasi awal, ketika peneliti belum mengetahui dengan pasti aspek-aspek apa yang akan menjadi perhatian utama dalam dataset baru.

Sebagai contoh, pada kalimat “*Danantara punya tujuan bagus, tapi transparansi pengelolaannya masih dipertanyakan*”, pendekatan ABSA akan mengidentifikasi dua aspek berbeda: aspek pertama adalah “*tujuan kebijakan*” dan aspek kedua adalah “*transparansi pengelolaan*”. Setelah aspek diidentifikasi, ditentukan sentimen yang diekspresikan terhadap setiap aspek tersebut, yaitu *aspect-level sentiment classification*. Pada tahapan ini, model harus memahami konteks semantik untuk menentukan polaritas (positif, negatif, atau netral) yang paling tepat untuk setiap aspek. Terdapat beberapa pendekatan untuk klasifikasi sentimen berbasis aspek yang telah dikembangkan dalam penelitian sebelumnya [20], [21]. Pendekatan *lexicon-based* yang berisi kata-kata dan skornya untuk menghitung polaritas secara agregat. Pendekatan ini memiliki keuntungan dalam hal interpretabilitas dan tidak memerlukan data training yang besar, namun kurang mampu menangani konteks yang

kompleks dan kata-kata baru yang tidak ada dalam kamus. Pendekatan berbasis *machine learning* menggunakan algoritma klasifikasi tradisional, seperti *Support Vector Machine* (SVM), *Random Forest*, atau *Naïve Bayes* dengan fitur-fitur yang diekstrak secara manual dari teks. Pendekatan ini lebih fleksibel dibandingkan *lexicon-based*, namun masih memerlukan *feature engineering* untuk menangkap hubungan semantik yang kompleks. Pendekatan ketiga adalah berbasis *deep learning*, yang menggunakan arsitektur *neural network* (LSTM), CNN, atau transformer models (BERT atau IndoBERT) yang dapat belajar representasi fitur yang kaya dan *contextual* dari *data training*. Model-model ini terbukti mampu mencapai performa superior dalam menangkap dependensi jangka panjang dan ambiguitas semantik dalam teks.

Dengan hasil dari aspek dan sentimen (*aspect-sentiment pair*) akan menunjukkan aspek mana yang mendapat sentimen tertentu, memberikan insight yang jauh lebih granular dibandingkan analisis sentimen konvensional [14]. Untuk mengilustrasikan perbedaan ini, contoh teks sebelumnya "Danantara punya tujuan bagus, tapi transparansi pengelolaannya masih dipertanyakan" jika dianalisis menggunakan sentimen analysis konvensional hanya akan menghasilkan satu label, yaitu netral (karena ada sentimen positif dan negatif yang saling mengimbangi dan tidak jelas label manakah yang dominan), sehingga informasi bernilai tinggi tentang aspek spesifik yang menjadi kritik atau dukungan akan hilang. Sebaliknya, pendekatan ABSA akan menghasilkan output terstruktur berupa dua pasangan. Pertama ("tujuan kebijakan", positif) dan kedua ("transparansi pengelolaan", negatif) yang memberikan pemahaman mendalam tentang dimensi mana yang mendapat pujian dan dimensi mana yang mendapat kritik.

Dalam konteks kebijakan Danantara, opini masyarakat umumnya bersifat multidimensional, normatif, dan sering kali disampaikan secara implisit. Karakteristik ini membedakan data kebijakan publik dari domain ulasan produk atau layanan yang lebih eksplisit, sehingga menuntut pendekatan ABSA yang mampu memetakan sentimen pada level aspek secara sistematis dan kontekstual [14]. ABSA memiliki beberapa keunggulan signifikan dibandingkan analisis sentimen konvensional yang sangat relevan untuk mendukung pengambilan keputusan berbasis bukti. ABSA mampu mengidentifikasi sumber sentimen secara eksplisit, menunjukkan aspek mana yang mendapat dukungan dan aspek mana yang mendapat kritik, sehingga pembuat kebijakan dapat fokus pada perbaikan yang paling diperlukan daripada bereaksi terhadap sentimen global yang ambigu. ABSA memberikan pemahaman multi-dimensional yang lebih mendalam, karena kebijakan

Danantara memiliki banyak dimensi (transparansi, efektivitas, dampak ekonomi, kredibilitas lembaga, dan lainnya) dan ABSA memungkinkan analisis terpisah dan terukur untuk setiap dimensi. ABSA mendukung *evidence-based policy making* yang lebih kuat [4], karena dengan mengetahui aspek spesifik yang menjadi perhatian publik beserta polaritas sentimen terhadap setiap aspek, pemerintah dapat merancang strategi komunikasi dan perbaikan kebijakan yang lebih tepat sasaran dan berbasis data kuantitatif. Berdasarkan tantangan tersebut, ABSA tidak dapat dipandang sebagai proses tunggal yang berdiri sendiri, melainkan sebagai bagian dari ekosistem metodologi yang melibatkan identifikasi aspek, pemodelan konteks, dan klasifikasi sentimen yang saling berkaitan. Literatur ABSA modern menekankan bahwa keberhasilan pemetaan sentimen berbasis aspek sangat dipengaruhi oleh bagaimana setiap komponen metodologi tersebut diintegrasikan untuk memahami hubungan antara aspek dan opini dalam konteks bahasa alami [22]. Oleh karena itu, pendekatan ABSA modern menuntut integrasi antara pemodelan konteks global dan pemodelan urutan lokal, karena opini masyarakat sering menyampaikan evaluasi aspek secara tidak eksplisit dan terpisah dari ekspresi sentimennya.

2.1.2 LDA (*Latent Dirichlet Allocation*)

LDA merupakan salah satu metode *topic modeling* yang berfungsi untuk menemukan pola atau tema tersembunyi (*latent topics*) dari kumpulan dokumen teks tanpa memerlukan data berlabel. Model ini bekerja dengan mengasumsikan bahwa setiap dokumen terdiri atas kombinasi dari beberapa topik, dan setiap topik memiliki distribusi tertentu atas kata-kata. Maka, adanya LDA akan membantu kita memahami apa yang sebenarnya dibicarakan dalam kumpulan teks. Terdapat penelitian yang dilakukan oleh Hidayati [23] mengungkap bahwa pola percakapan tersembunyi dalam ribuan ulasan pelanggan atau postingan media sosial bisa menangkap tema-tema penting yang mungkin tidak langsung terlihat, khususnya ketika berhadapan dengan volume data teks yang besar dan beragam.

Dalam konteks media sosial di Indonesia, Wahyudi et al. [24] membuktikan bahwa LDA efektif untuk mengidentifikasi pola diskusi publik dalam kumpulan data besar. Penerapan LDA pada judul-judul berita daring selama masa darurat COVID-19 menghasilkan pengelompokan topik dominan seperti respons kebijakan pemerintah, penanganan kesehatan, dan partisipasi masyarakat. Hal ini menunjukkan bahwa pendekatan LDA relevan digunakan untuk menelusuri kecenderungan sentimen terhadap isu kebijakan yang kompleks dan dinamis. Sementara itu, Santoso et al. [25] menegaskan bahwa penggunaan LDA pada data media sosial, seperti Instagram dapat membantu memetakan

aspek-aspek utama dalam wacana publik sebelum dilakukan ABSA. Dengan mengekstraksi topik-topik laten terlebih dahulu, peneliti dapat mengidentifikasi fokus perhatian publik terhadap isu tertentu, misalnya transparansi, tata kelola, atau keuangan. Oleh karena itu, dalam penelitian ini LDA digunakan sebagai tahap awal untuk menentukan aspek-aspek utama yang akan dianalisis lebih lanjut menggunakan model ABSA.

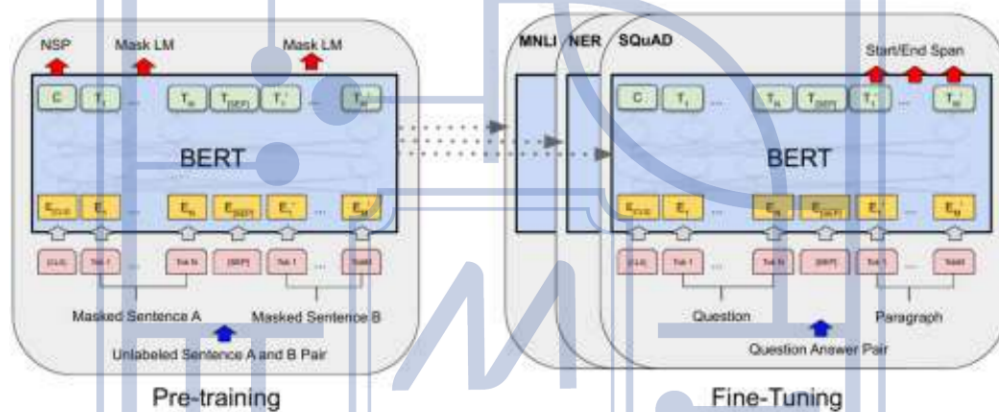
LDA dikategorikan sebagai *generative probabilistic* model yang menggunakan distribusi *Dirichlet* dengan parameter α dan β untuk membentuk distribusi topik dalam dokumen serta distribusi kata dalam topik. Proses inferensi seperti *Gibbs sampling* digunakan untuk mengestimasi variabel laten dari dokumen yang diamati [23], [24]. Pendekatan ini memungkinkan pemetaan tema secara *probabilistik*, sehingga setiap dokumen dapat direpresentasikan sebagai campuran dari berbagai topik dengan proporsi tertentu. Dengan begitu, LDA menjadi dasar penting dalam proses ekstraksi aspek untuk mendukung analisis sentimen yang lebih terarah.

2.1.3 BERT (*Bidirectional Encoder Representations from Transformers*)

Transformer merupakan arsitektur *deep learning* yang dirancang untuk memproses data sekuensial dengan memanfaatkan mekanisme *self-attention*, sehingga mampu memodelkan hubungan antar-token secara paralel dan menangkap dependensi jarak jauh secara lebih efektif dibandingkan model berbasis *recurrent neural network* (RNN). Salah satu implementasi transformer yang paling berpengaruh adalah BERT, yaitu model pre-training yang dirancang untuk mempelajari representasi kontekstual dua arah dari teks. Dalam ABSA, BERT berperan sebagai *feature extractor* kontekstual yang mampu menangkap makna aspek dan opini secara implisit, terutama pada teks pendek media sosial yang sarat ambiguitas semantik. Representasi vektor yang dihasilkan BERT memungkinkan pemisahan konteks aspek dan sentimen secara lebih presisi dibandingkan *embedding* statis, sehingga menjadi fondasi penting dalam pengembangan ABSA modern [26].

Zevana dan Riana [27] telah melakukan penelitian *fine-tuning* pada IndoBERT dan menggabungkannya dengan arsitektur CNN serta Bi-LSTM. Hasilnya menunjukkan peningkatan *accuracy* signifikan dalam tugas klasifikasi teks berbahasa Indonesia, terutama pada data yang mengandung slang dan struktur kalimat tidak baku. Temuan ini mengindikasikan bahwa meskipun BERT mampu menghasilkan representasi semantik kontekstual yang kuat, penambahan lapisan sekuensial, seperti Bi-LSTM diperlukan untuk memodelkan dependensi urutan token secara eksplisit, yang sangat relevan dalam tugas

ABSA. Dengan demikian, BERT tidak hanya diposisikan sebagai model *end-to-end*, tetapi sebagai komponen representasi awal dalam sistem *hybrid* yang lebih komprehensif. Penelitian ini membuktikan bahwa BERT dapat berperan sebagai representasi semantik yang kuat dalam sistem *hybrid* yang lapisan tambahan membantu menangkap fitur urutan atau spasial dari teks. Selain itu, Ahmadian et al. [9] menjelaskan bahwa model berbasis transformer, seperti BERT dan RoBERTa mampu memberikan performa stabil pada tugas klasifikasi sentimen dan emosi berbahasa Indonesia, bahkan pada dataset terbatas. Mereka menekankan pentingnya *pre-training* berbasis bahasa lokal dan transfer learning untuk mengatasi keterbatasan data berlabel.



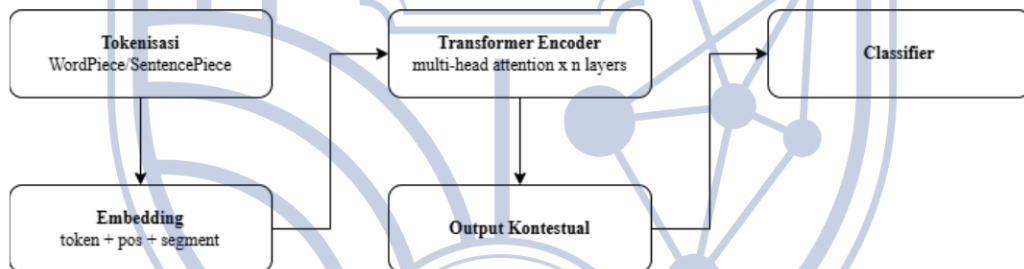
Gambar 2.1 Arsitektur BERT [3]

Gambar 2.1 menunjukkan bahwa arsitektur BERT terdiri dari dua tahap utama, yaitu *pre-training* dan *fine-tuning*. Tahapan *pre-training*, BERT akan menggunakan dua metode pembelajaran, yaitu MLM dan NSP dengan memanfaatkan pasangan kalimat tidak berlabel (*Unlabeled Sentence A dan B Pair*) untuk mempelajari representasi kontekstual dua arah dari teks. Arsitektur bidireksional pada BERT memberikan kapabilitas untuk memprediksi konteks sebelum dan sesudah kata secara paralel, sehingga model mampu memahami konteks kalimat secara mendalam dan mempelajari hubungan semantik antar kata dengan baik. Pada tahapan *fine-tuning*, model BERT yang telah dilatih sebelumnya kemudian disesuaikan untuk tugas spesifik, seperti *Machine Reading Comprehension* (MNLI/NER/SQuAD) dengan menambahkan satu output layer tanpa memodifikasi arsitektur dasar, yang ditunjukkan melalui pasangan pertanyaan-jawaban (*question answer pair*). Proses *fine-tuning* ini membuat modifikasi model menjadi lebih efisien baik dari segi kebutuhan data training untuk tugas spesifik maupun komputasi, sehingga dapat menghasilkan model dengan kebutuhan resource yang rendah namun tetap memiliki performa yang baik. Representasi input pada BERT memiliki tiga komponen yaitu *token*

embeddings (C , T_1 , T_2 , dll), *segment embeddings* (E_{Sub} , E_1 , E_2 , dll), dan *position embeddings* (Tok 1, Tok N, dll) yang memungkinkan model memproses hingga 512 token maksimum dengan menggunakan token special, seperti CLS untuk klasifikasi dan SEP untuk memisahkan kalimat [28]. Oleh karena itu, teori BERT menjadi dasar utama bagi penelitian ini, sebelum mengarah pada adaptasi model khusus bahasa Indonesia (IndoBERT) yang akan dijelaskan pada sub-bagian berikutnya.

2.1.4 IndoBERT

Model IndoBERT merupakan adaptasi dari arsitektur BERT yang dirancang secara khusus untuk memahami bahasa Indonesia. Model ini dikembangkan dengan melatih algoritma transformer menggunakan korpus besar berbahasa Indonesia yang mencakup berbagai domain teks, seperti berita, artikel, dan unggahan media sosial. Menurut Saadah et al. [29], keunggulan utama IndoBERT terletak pada kemampuannya dalam menangkap konteks dua arah (*bidirectional context*) yang makna suatu kata tidak hanya dipengaruhi oleh kata sebelumnya, tetapi juga oleh kata setelahnya. Karakteristik ini menjadi krusial dalam analisis teks kebijakan publik, di mana opini masyarakat sering disampaikan secara implisit, normatif, dan tidak selalu menyebutkan aspek secara eksplisit. Oleh karena itu, IndoBERT sangat relevan digunakan dalam ABSA untuk domain kebijakan publik berbahasa Indonesia



Gambar 2.2 Alur Proses Pembentukan Respresentasi Kontekstual IndoBERT [28]

Gambar 2.2 tahapan pemrosesan teks IndoBERT dimulai dengan tokenisasi *sub-word*, yaitu pemecahan kata menjadi unit-unit kecil (sub-tokens) menggunakan metode *WordPiece tokenization*. Teknik ini memungkinkan model mengenali kata-kata baru atau slang yang tidak terdapat dalam kosakata awalnya. Misalnya, kata “pemerintahannya” akan diuraikan menjadi “pemerintah” dan “##annya”. Apriliani et al. [28] menjelaskan bahwa strategi tokenisasi *sub-word* ini berperan penting dalam menangani variasi morfologis serta istilah campuran (*code-mixing*) yang umum ditemukan dalam teks media sosial Indonesia. Setelah proses tokenisasi, setiap token dikonversi menjadi representasi numerik melalui *embedding layer*, yang terdiri dari token *embeddings*, *segment embeddings*, dan *position*

embeddings sama dengan bagian BERT. Lapisan ini membantu model memahami perbedaan konteks antar-token dalam satu urutan kalimat.

Kemudian teks yang sudah direpresentasikan dalam bentuk *embedding* akan diproses melalui 12 lapisan *encoder transformer*, yang masing-masing terdiri atas dua komponen utama, yaitu *multi-head self-attention* dan *feed-forward neural network*. Ahmadian et al. [9] menjelaskan bahwa mekanisme *self-attention* memungkinkan model menghitung hubungan semantik antar-kata dalam satu kalimat, sehingga IndoBERT mampu memahami konteks yang lebih kompleks, termasuk hubungan antar-frasa yang berjauhan. Lapisan-lapisan *encoder* ini bekerja secara bertumpuk untuk memperdalam pemahaman model terhadap konteks semantik teks hingga menghasilkan representasi kontekstual (*contextual representation*), yaitu vektor makna yang menggambarkan hubungan setiap kata terhadap keseluruhan kalimat. Representasi kontekstual yang dihasilkan oleh IndoBERT ini kemudian dapat dimanfaatkan sebagai input pada lapisan lanjutan dalam arsitektur ABSA, seperti Bi-LSTM, untuk memodelkan hubungan sekuensial antar-token yang berkaitan dengan aspek dan polaritas sentimen. Dengan demikian, IndoBERT berfungsi sebagai tahap awal dalam ekosistem metodologi ABSA yang terintegrasi.

Kemudian Aryanti et al. [3] menambahkan bahwa representasi kontekstual hasil IndoBERT dapat digunakan sebagai masukan (input) untuk berbagai tugas lanjutan seperti ABSA, klasifikasi respons publik, dan deteksi emosi. Dengan melakukan *fine-tuning*, model dapat disesuaikan dengan domain tertentu, misalnya ulasan produk, kebijakan publik, atau aplikasi kesehatan mental. Hasil penelitian menunjukkan bahwa penggunaan IndoBERT secara *fine-tuned* mampu meningkatkan *accuracy* secara signifikan dibandingkan model berbasis LSTM atau CNN tradisional, karena setiap token telah membawa pemahaman kontekstual yang kaya dari keseluruhan teks. IndoBERT diposisikan sebagai *feature extractor* kontekstual utama yang menghasilkan representasi semantik token-level dari teks media sosial berbahasa Indonesia. Representasi ini selanjutnya diproses oleh lapisan BiLSTM untuk menangkap dependensi urutan dan hubungan aspek-sentimen secara lebih eksplisit. Pendekatan ini sejalan dengan tren ABSA modern yang menekankan integrasi *transformer* dan model sekuensial untuk meningkatkan performa ABSA pada domain kompleks, seperti kebijakan publik.

2.1.5 BiLSTM (Bidirectional Long Short-Term Memory)

BiLSTM merupakan varian jaringan saraf rekuren (RNN) yang menggunakan sel LSTM untuk memproses urutan data dari dua arah, yaitu maju (*forward*) dan mundur (*backward*), sehingga mampu menangkap konteks kata sebelum dan sesudah suatu token dalam kalimat. Hal ini penting karena dalam teks sentimen (termasuk media sosial) sering terdapat negasi, ironi, atau rujukan ke kata-sebelumnya maupun setelahnya yang mempengaruhi makna. Studi yang dilakukan oleh Cahyaningtyas et al. pada ulasan hotel berbahasa Indonesia menunjukkan bahwa BiLSTM termasuk salah satu topologi yang dipertimbangkan dalam tugas ABSA karena kemampuannya memodelkan urutan dua arah [17]. Dengan ABSA, kemampuan BiLSTM dalam memodelkan dependensi sekuensial dua arah menjadi krusial untuk memahami relasi antara aspek dan ekspresi sentimen yang dapat muncul pada posisi berbeda dalam satu kalimat

Dalam kerangka ABSA, tugas utama mencakup identifikasi aspek yang dibahas dalam teks serta klasifikasi sentimen terhadap setiap aspek tersebut. BiLSTM berperan penting pada kedua tahap ini karena pemrosesan dua arah memungkinkan model memahami kemunculan aspek sekaligus ekspresi sentimen yang dapat muncul sebelum maupun setelah aspek dalam kalimat [18]. Pada pendekatan ABSA modern, Bi-LSTM umumnya tidak digunakan secara terpisah, melainkan dikombinasikan dengan *embedding* kontekstual yang dihasilkan oleh model transformer, seperti IndoBERT. Representasi token-level dari IndoBERT kemudian diproses ulang oleh Bi-LSTM untuk memodelkan dependensi urutan secara eksplisit, sehingga hubungan aspek–sentimen dapat ditangkap secara lebih presisi, khususnya pada teks bahasa Indonesia yang bersifat informal dan tidak baku. Turangan et al. [30] menunjukkan bahwa integrasi Bi-LSTM dalam pipeline hybrid berbasis BERT mampu meningkatkan *accuracy* ekstraksi aspek dan klasifikasi sentimen dibandingkan pendekatan berbasis embedding statis. Dibandingkan lapisan Dense atau CNN yang cenderung menangkap pola lokal tanpa mempertimbangkan urutan dua arah secara penuh, Bi-LSTM memiliki keunggulan dalam memodelkan dependensi kontekstual jangka panjang antara aspek dan ekspresi sentimen [19]. Keunggulan ini menjadi krusial dalam ABSA kebijakan publik, di mana aspek dan opini sering muncul pada posisi yang berjauhan dalam satu kalimat.

2.1.6 Evaluation Metrics

Dalam penelitian ABSA, evaluasi model bertujuan untuk menilai seberapa baik sistem mengenali aspek dalam teks dan menentukan polaritas sentimen (positif, negatif, atau

netral). Beberapa metrik evaluasi yang akan digunakan pada penelitian ini adalah *accuracy*, *presisi*, *recall*, dan *F1-score*.

Accuracy (*accuracy*) akan menunjukkan proporsi prediksi yang benar dari total data sebagai berikut.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \dots\dots\dots(1)$$

TP (*True Positive*) adalah jumlah prediksi positif yang benar, TN (*True Negative*) jumlah prediksi negatif yang benar, FP (*False Positive*) jumlah prediksi positif yang salah, dan FN (*False Negative*) jumlah prediksi negatif yang salah [12]. *Accuracy* cocok untuk dataset yang kelasnya relatif seimbang, namun bisa menyesatkan jika kelas tertentu jauh lebih dominan [17].

Presisi (*Precision*) mengukur seberapa tepat model saat memprediksi sebuah kelas positif, misalnya sentimen positif sebagai berikut.

$$Precision = \frac{TP}{TP+FP} \dots\dots\dots(2)$$

Sementara itu, *Recall* mengukur kemampuan model untuk menangkap semua kasus positif yang sebenarnya sebagai berikut.

$$Recall = \frac{TP}{TP+FN} \dots\dots\dots(3)$$

Kedua metrik ini saling melengkapi. Agar penilaian lebih seimbang, digunakan *F1-score*, yang merupakan rata-rata harmonis antara presisi dan recall sebagai berikut.

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \dots\dots\dots(4)$$

F1-score membantu memastikan model tidak hanya tepat (presisi tinggi) tetapi juga tidak melewatkan banyak kasus positif (recall tinggi) [18].

2.2 Penelitian Terdahulu

Penelitian mengenai ABSA telah mengalami perkembangan signifikan dalam beberapa tahun terakhir, terutama dengan kemajuan model berbasis transformer dan pendekatan hibrida untuk memahami opini publik secara mendalam. Menurut Zhu et al. [19], ABSA memungkinkan analisis sentimen tidak hanya pada tingkat kalimat atau dokumen, tetapi pada tingkat aspek spesifik, sehingga menghasilkan pemahaman yang lebih granular

terhadap opini masyarakat. Pendekatan ini sangat berguna untuk menganalisis kebijakan publik karena mampu mengungkap dimensi mana yang mendapat dukungan dan mana yang menuai kritik. Hua et al. [14] melakukan tinjauan sistematis terhadap berbagai domain penerapan ABSA dan menunjukkan bahwa metode ini banyak digunakan pada ranah pelayanan publik, kebijakan sosial, dan transportasi, karena kemampuannya mendeteksi isu spesifik berdasarkan aspek. Namun, sebagian besar penelitian tersebut masih terbatas pada bahasa Inggris dan belum banyak diadaptasi untuk konteks linguistik Indonesia. Penelitian yang dilakukan Perikos et al. [31] memperkenalkan *explainable* ABSA dengan model transformer yang memberikan interpretabilitas tinggi pada setiap pasangan aspek dan sentimen. Model ini memungkinkan pembuat kebijakan memahami alasan di balik prediksi sentimen, namun penelitian tersebut masih bersifat konseptual dan belum diterapkan pada domain kebijakan nyata. Dalam konteks pengambilan keputusan kebijakan, Nath et al. [4] mengusulkan model ABSA berbasis *hybrid deep learning* untuk mengevaluasi respons publik terhadap kebijakan pemerintah India. Hasilnya menunjukkan bahwa ABSA efektif dalam mengungkap dimensi-dimensi spesifik (efisiensi, transparansi, dan keadilan) yang menjadi perhatian publik, sekaligus mendukung *evidence-based policy making*. Terdapat penelitian yang dilakukan oleh Sejati et al. [7] yang menganalisis opini publik terhadap Kementerian Keuangan Republik Indonesia (Kemenkeu RI) melalui media sosial menggunakan ABSA berbasis BERT. Hasilnya menunjukkan bahwa ABSA dapat memberikan wawasan mendalam tentang bagaimana masyarakat menilai aspek transparansi, pelayanan publik, dan kebijakan fiskal lembaga tersebut. Namun, penelitian tersebut masih berfokus pada satu institusi dan belum mengkaji kebijakan spesifik yang berdampak langsung pada masyarakat luas.

Dengan demikian, penelitian ini mengisi celah tersebut dengan menerapkan model ensemble IndoBERT–BiLSTM untuk melakukan analisis sentimen berbasis aspek terhadap kebijakan Danantara, dengan tujuan memberikan pemahaman yang lebih spesifik dan kuantitatif terhadap dimensi kebijakan yang mendapat dukungan maupun kritik publik. Berikut adalah Tabel 2.2 yang menunjukkan ringkasan penelitian terdahulu yang menjadi rujukan utama dalam penelitian ini, mencakup peneliti, judul penelitian (tahun), metode penelitian, dan hasil penelitian.

Tabel 2.2 Ringkasan Penelitian Terdahulu

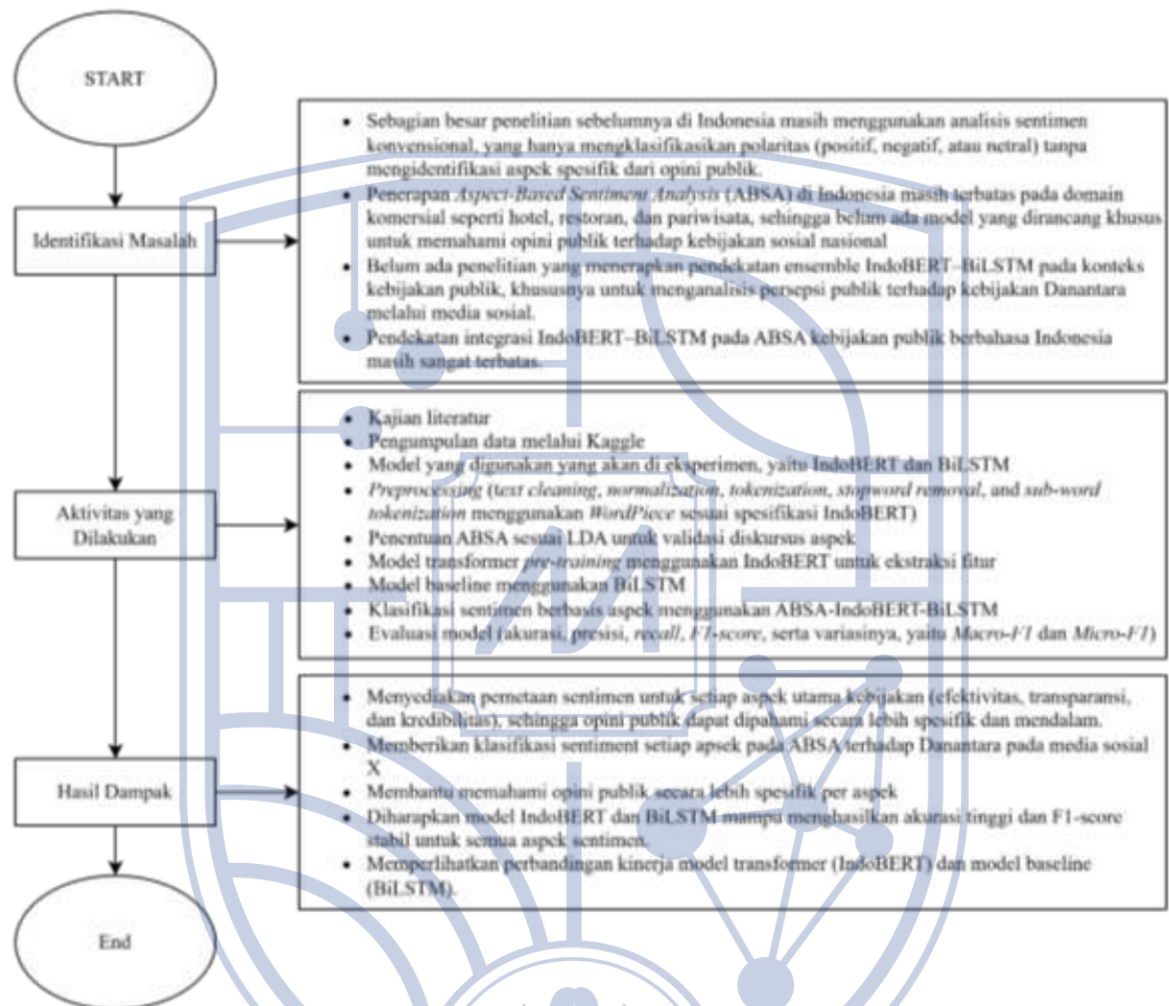
Penulis (Tahun)	Judul Penelitian	Metode Penelitian	Hasil Penelitian
Liang Zhu, Ming Xu, Yong Bao, Yan Xu, dan Xiangliang Kong, (2022)	<i>Deep Learning for Aspect-Based Sentiment Analysis: A Review</i>	Studi tinjauan literatur (review) terhadap model deep learning seperti LSTM, CNN, dan Transformer untuk ABSA.	ABSA mampu menganalisis sentimen pada tingkat aspek, memberikan pemahaman yang lebih detail terhadap opini publik, dan menunjukkan bahwa model berbasis Transformer memberikan performa terbaik.
Yucheng Hua, Paul Denny, Katya Taskova, dan Jack Wicker, (2023)	<i>A Systematic Review of Aspect-Based Sentiment Analysis: Domains, Methods, and Trends</i>	Systematic Literature Review (SLR) terhadap penerapan ABSA di berbagai domain (pelayanan publik, kebijakan sosial, transportasi).	ABSA banyak digunakan di bidang pelayanan publik dan kebijakan sosial, namun, mayoritas penelitian masih terbatas pada teks berbahasa Inggris dan belum banyak diterapkan di konteks linguistik Indonesia.
Isidoros Perikos dan Athanasios Diamantopoulos, (2024)	<i>Explainable Aspect-Based Sentiment Analysis Using Transformer Models</i>	Model Transformer dengan penjelasan berbasis attention (<i>Explainable ABSA</i>).	ABSA berbasis Transformer dapat memberikan interpretabilitas tinggi terhadap pasangan aspek-sentimen, namun masih bersifat konseptual dan belum diterapkan pada kasus kebijakan nyata.
Dipak Nath, Shiv Kumar	<i>Unveiling the Impact of</i>	Hybrid Deep Learning	ABSA efektif dalam mengidentifikasi dimensi

Dwivedi, dan Amit Prakash Dwivedi, (2023)	<i>Indian Government Policies Using ABSA with Multi-Criteria Decision- Making and Hybrid Deep Learning</i>	(menggabungkan model neural network dengan multi-criteria decision-making).	kebijakan publik (efisiensi, transparansi, keadilan), mendukung pengambilan keputusan berbasis bukti (<i>evidence-based policy making</i>).
Priska Trisna Sejati, Farrikh Alzami, Aris Marjuni, Heni Indrayani, Ika Dewi Puspitarini, (2024)	<i>Aspect-Based Sentiment Analysis for Enhanced Understanding of 'Kemenkeu' Tweets</i>	ABSA berbasis BERT diterapkan pada data Twitter tentang Kementerian Keuangan RI.	ABSA mampu memberikan wawasan mendalam tentang opini publik terhadap aspek transparansi, pelayanan publik, dan kebijakan fiskal, namun fokus hanya pada satu lembaga tanpa menyoroti kebijakan spesifik berskala nasional.

Berdasarkan tinjauan terhadap lima penelitian terdahulu pada Tabel 2.1, dilakukan analisis pemilihan model untuk menentukan arsitektur yang rujukan utama bagi penelitian ini. Meskipun model IndoBERT murni telah menunjukkan performa kuat untuk sentimen umum [29], penelitian ini menerapkan arsitektur hibrida IndoBERT-BiLSTM untuk menangkap ketergantungan urutan dua arah yang lebih kompleks pada teks media sosial X. Pemilihan model ini didasarkan pada analisis bahwa integrasi *pre-trained* model bahasa Indonesia dengan lapisan sekuensial BiLSTM mampu menutupi kelemahan model dasar dalam memahami konteks kalimat yang tidak beraturan. Selain itu, berbeda dengan penelitian Turangan et al. [30] yang berfokus pada ekstraksi fitur aspek secara umum, penelitian ini mengisi *research gap* dengan mengombinasikan identifikasi aspek berbasis LDA dan validasi statistik melalui *McNemar's Test*, sehingga hasil evaluasi model tidak hanya memiliki akurasi tinggi tetapi juga terbukti signifikan secara statistika dan bukan merupakan faktor kebetulan.

2.3 Kerangka Konseptual

Kerangka ini dirancang untuk menunjukkan alur penyelesaian masalah penelitian melalui solusi yang diusulkan dan hasil yang diperoleh pada penelitian ini adalah sebagai berikut.



Gambar 2.3 Kerangka Konseptual

Gambar 2.3 menggambarkan alur sistematis dalam menganalisis sentimen publik terhadap kebijakan Danantara menggunakan pendekatan ABSA. Alur ini terdiri dari 3 bagian utama, yaitu identifikasi masalah, aktivitas yang dilakukan dan hasil beserta dampak.

Penelitian ini diawali dengan identifikasi permasalahan terkait volume dan kompleksitas opini masyarakat terhadap kebijakan Danantara di media sosial X. Informasi yang diunggah masyarakat pada platform tersebut sangat beragam, mulai dari dukungan, kritik, hingga opini netral, namun sulit untuk dianalisis secara manual mengingat tingginya jumlah unggahan yang diproduksi setiap hari. Oleh karena itu, penelitian ini menggunakan dataset yang berisi kumpulan teks berbahasa Indonesia yang relevan dengan isu sosial dan

kebijakan publik. Dataset tersebut menjadi fondasi untuk proses analisis yang bertujuan memahami kecenderungan opini publik secara lebih terstruktur dan berbasis aspek.

Setelah dataset diperoleh, tahap berikutnya adalah melakukan pra-pemrosesan data agar teks siap digunakan dalam proses pelatihan model. Tahap ini meliputi pembersihan teks dari karakter yang tidak relevan, normalisasi bahasa, tokenisasi, serta pelabelan sentimen jika diperlukan. Proses ini penting karena kualitas data sangat menentukan performa model analisis sentimen. Selanjutnya, data dipersiapkan untuk mengembangkan dua pendekatan pemodelan, yaitu *fine-tuned* IndoBERT dan BiLSTM. IndoBERT dipilih karena merupakan model berbasis transformer yang telah dilatih secara khusus pada teks bahasa Indonesia, sementara BiLSTM digunakan sebagai pembanding berbasis deep learning klasik yang mampu menangkap konteks sekuensial dalam kalimat. Kedua model dilatih dan diuji menggunakan data yang sama untuk memperoleh hasil yang objektif dan terukur.

Dan tahap terakhir menyediakan pemetaan sentimen untuk setiap aspek utama kebijakan, sehingga opini publik dapat dipahami secara lebih spesifik dan mendalam. Dan juga akan mengevaluasi performa kedua model dengan menggunakan metrik *accuracy*, *precision*, *recall*, dan F1-score. Evaluasi ini bertujuan untuk mengetahui model mana yang memiliki kemampuan terbaik dalam mengidentifikasi sentimen positif, negatif, dan netral dalam konteks diskusi mengenai Danantara di media sosial. Hasil evaluasi tersebut kemudian dianalisis untuk melihat bagaimana pola sentimen publik terbentuk dan apa implikasinya terhadap persepsi masyarakat. Dengan demikian, penelitian tidak hanya membandingkan performa model, tetapi juga menggambarkan bagaimana sentimen yang berkembang di media sosial dapat memberikan dampak terhadap pemahaman publik mengenai kebijakan Danantara. Temuan ini diharapkan dapat menjadi dasar bagi pembuat kebijakan atau pihak terkait dalam menilai opini publik secara lebih komprehensif.