

BAB II

LANDASAN TEORI

2.1 Landasan Teori

2.1.1 Pengertian *Jabar Quick Response*

Pemerintah Provinsi Jawa Barat meluncurkan program *Jabar Quick Response* (JQR) yang kreatif untuk merespons dengan cepat berbagai permasalahan masyarakat. Meskipun JQR tidak memiliki kerangka hukum yang sangat jelas, penerapannya didasarkan pada sejumlah prinsip hukum yang lebih umum, seperti:

1. (Undang-Undang Nomor 25 Tahun 2009 tentang Pelayanan Publik, 2009)

Undang-undang ini mengatur penyampaian layanan publik yang berkualitas tinggi, tepat waktu, mudah digunakan, dan dengan harga yang wajar. Untuk melayani masyarakat dengan cepat dan efisien, JQR mematuhi nilai-nilai ini. Salah satu tanggung jawab utama *Jabar Quick Response*, yaitu menangani keluhan publik, juga diatur oleh regulasi ini.

2. (Peraturan Presiden Republik Indonesia Nomor 76 Tahun 2013 tentang Pengelolaan Pengaduan Pelayanan Publik, 2013)

Prinsip umum untuk menangani keluhan pelayanan publik, termasuk standar pelayanan minimum, prosedur pengaduan, dan penyelesaian keluhan, diatur dalam Peraturan Presiden ini.

3. Peraturan Daerah Provinsi Jawa Barat

Jabar Quick Response adalah cara untuk menerapkan strategi regional yang lebih komprehensif, yang dijelaskan dalam sejumlah peraturan daerah provinsi.

Jabar *Quick Response* didukung oleh beberapa landasan hukum yang lebih besar, meskipun kurang memiliki perlindungan hukum yang sangat tepat. Ini memungkinkan Jabar *Quick Response* untuk melayani masyarakat dengan sukses dan efisien (Anwar & Rustandi, 2023).

2.1.2 Analisis Sentimen

Salah satu aplikasi *Natural Language Processing* (NLP) yang paling terkenal adalah analisis sentimen. Subjek ilmiah NLP mempelajari bagaimana cara membuat komputer berfungsi dan berpikir seperti manusia. Di dalam kecerdasan buatan terdapat bidang pemrosesan bahasa alami, atau NLP. Dalam perkembangan data mining, *Artificial Intelligence* (AI) adalah salah satu dari empat subbidang penambangan data, bersama dengan pengambilan informasi, basis data, dan statistik. Dalam penerapannya, *AI* juga membutuhkan *machine learning* sebagai teknik untuk menyelesaikan. Penggunaan *machine learning* bertujuan untuk menggantikan peran manusia dalam pengambilan keputusan. Berbeda dengan manusia, *machine learning* tidak melibatkan emosi, sehingga setiap keputusan yang dihasilkan sepenuhnya didasarkan pada data yang telah diproses. (Irwansyah Saputra dan Dinar Ajeng Kristiyanti, 2022).

Analisis sentimen adalah prosedur komputasi yang mengelola, memahami, dan mengklasifikasikan emosi sebagai positif atau negatif dengan menggunakan teknik analisis teks pada data tekstual. Analisis sentimen adalah teknik yang umum digunakan karena meningkatnya permintaan dari individu atau kelompok untuk memahami sudut pandang orang lain. Dataset yang digunakan memiliki dampak pada analisis juga karena akan ditangani dengan cara yang berbeda (Widayat,

2021). Analisis sentimen menargetkan korporasi serta individu. Analisis sentimen atau yang dikenal sebagai penambangan opini adalah praktik menilai perasaan pengguna melalui pemeriksaan tulisan mereka. Analisis sentimen tidak terbatas pada karakter, tetapi juga dapat digunakan untuk menganalisis masalah dengan program, perusahaan, barang, aplikasi, dan topik lain yang terbuka untuk diskusi publik (Herlinawati dkk., 2020). Analisis sentimen memiliki manfaat dalam menghemat waktu dan usaha saat melakukan penelitian dengan banyak data. Berikut contoh penerapan analisis sentimen:

1. Di bidang bisnis, contohnya adalah untuk mengetahui bagaimana tanggapan masyarakat terhadap reputasi suatu merek produk baru, sehingga merek tersebut dapat ditingkatkan.
2. Di ranah politik, contohnya adalah untuk mengetahui seberapa populer seorang tokoh, sehingga dapat memberikan wawasan lebih mengenai tokoh yang bersangkutan.
3. Berbagai program seperti vaksinasi, penanganan COVID-19, pemilihan presiden, dan lainnya dapat dianalisis untuk memberikan masukan yang berguna dalam menyempurnakan kebijakan pemerintah terkait.

Secara garis besar, proses analisis sentimen terdiri dari lima tahapan utama, yaitu *crawling data*, *pre-processing*, *feature selection*, *classification*, dan *evaluation*. Data tidak terstruktur dapat diorganisir melalui analisis sentimen. Di antara keuntungan analisis sentimen adalah kemampuannya untuk mengevaluasi dan memberikan ide untuk berbagai bidang. Analisis sentimen merupakan sebuah alat yang bermanfaat untuk mengevaluasi pernyataan, peristiwa, maupun komentar

yang bersifat kontroversial. Selain itu, hasil dari analisis ini dapat memberikan wawasan bagi perusahaan, figur publik, maupun pemerintah dalam merumuskan langkah atau keputusan selanjutnya. (Natasuwarna, 2020).

Terdapat beberapa jenis analisis sentimen yaitu *emotion detection*, *aspectbased sentiment analysis*, dan *fine grand sentiment analysis*. *Fine sentiment analysis* merupakan jenis analisis sentimen yang memberikan penilaian secara lebih terperinci dan biasanya diterapkan dalam bidang *e-commerce*. *Emoticon detection* adalah jenis analisis yang dimaksudkan untuk memahami perasaan seperti kebahagiaan, kesedihan, kemarahan, dan lainnya yang terdapat dalam sebuah pesan. *Aspect-based sentiment analysis*, Analisis seperti ini dimanfaatkan untuk mengenali berbagai faktor yang memengaruhi opini serta penilaian yang diberikan oleh klien (Geofanni Nerissa Arviana, 2021a)

Analisis sentimen sebenarnya adalah proses multi-tahap yang dimulai dengan pengumpulan data dari sumber yang telah ditentukan (sosial media, website, dll), menggunakan metode seperti *web crawler* atau *web scraper* untuk menarik data. Selanjutnya, data tersebut dapat diolah menggunakan algoritma seperti *Naive Bayes* atau SVM untuk membangun model analisis sentimen. Langkah selanjutnya adalah tahap evaluasi, yang menentukan apakah model tersebut memadai atau masih perlu perbaikan. Penelitian sentimen oleh karena itu merupakan alat yang krusial untuk meramalkan tren pasar, mengukur opini publik, dan mendukung perusahaan serta organisasi dalam mengambil keputusan strategis (Geofanni Nerissa Arviana, 2021a).

2.1.3 *Natural Language Processing (NLP)*

Natural Language Processing (NLP) merupakan cabang penting dari ilmu komputer yang terletak di antara linguistik dan pembelajaran mesin. Proses ini memerlukan sejumlah besar perhitungan agar komputer dapat memahami pernyataan yang ditulis dalam bahasa manusia (Aditya Jain, 2018). Tujuan dari *Natural Language Processing (NLP)* adalah untuk mengurangi beban kerja pengguna dan memenuhi kebutuhan akan komunikasi bahasa yang murni antara manusia dan komputer. NLP memberikan manfaat bagi pengguna yang tidak memiliki waktu untuk menguasai bahasa baru karena tidak semua pengguna mungkin merupakan penutur asli dari bahasa yang spesifik untuk mesin (Khurana dkk., 2023).

Natural Language Processing (NLP) secara umum terdiri dari dua komponen utama, yaitu *Natural Language Understanding* dan *Natural Language Generation*, yang masing-masing berfungsi untuk memahami serta menghasilkan teks. NLP berperan sebagai fondasi utama dalam sistem pengenalan bahasa, seperti yang digunakan oleh *Siri (Apple)* maupun Google. Sistem NLP dapat dimulai dengan menentukan struktur dan sifat morfologis kata, seperti makna atau *part-of-speech*, pada tingkat kata setelah itu, dapat memeriksa urutan kata, tata bahasa, dan makna seluruh kalimat di tingkat kata dan tingkat kalimat. Dalam konteks tertentu, satu kata atau kalimat mungkin memiliki arti atau konotasi berbeda yang terhubung ke banyak kata atau kalimat lain dalam konteks (Khurana dkk., 2023). ada beberapa area utama pada NLP, diantaranya:

1. *Question Answering System (QAS)*, adalah kemampuan komputer untuk merespons pertanyaan dari pengguna. Dalam hal ini, pengguna mengajukan pertanyaan langsung kepada komputer daripada memasukkan kata kunci yang terkait dengan jawaban yang diinginkan.
2. *Summarization*, Ini memungkinkan pengguna untuk mengubah dokumen teks yang panjang menjadi slide presentasi dengan merangkum serangkaian email atau makalah
3. *Machine Translation*, salah satu alat yang telah menerapkan bidang studi ini adalah *Google Translate*, yang dapat menerjemahkan teks dari satu bahasa ke bahasa lain dan memahami bahasa manusia sebagai hasilnya.
4. *Speech Recognition*, adalah bidang NLP yang sangat menantang. Ini disebabkan oleh fakta bahwa komputer perlu memahami bahasa manusia, yang sering kali disampaikan di bidang ini melalui pertanyaan dan instruksi.
5. *Document Classification* merupakan salah satu bidang dalam NLP yang paling berhasil. Tugas utamanya adalah menentukan kategori atau tempat yang paling sesuai untuk dokumen baru yang dimasukkan ke dalam sistem. Pendekatan ini sangat bermanfaat dalam berbagai aplikasi seperti penyaringan spam, klasifikasi artikel berita, dan ulasan film.

Natural Language Processing (NLP) merupakan pendekatan yang sangat relevan untuk menganalisis data dari media sosial seperti *Twitter*. Hal ini karena *Twitter* menyediakan data teks dalam jumlah besar yang mencerminkan opini dan respons masyarakat secara *real-time*. Dengan NLP, teknologi ini dapat

mengekstraksi makna dari teks singkat yang umumnya memiliki batasan karakter, seperti *tweet*.

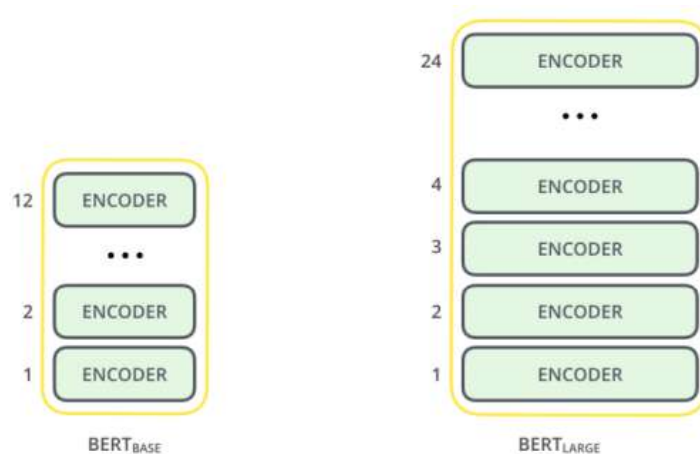
Namun, penerapan NLP pada teks Twitter berbahasa Indonesia memiliki tantangan tersendiri. Karakteristik bahasa Indonesia di media sosial cenderung informal, sering menggunakan bahasa gaul, singkatan, atau bahkan campuran bahasa. Selain itu, sumber daya NLP untuk Bahasa Indonesia masih terbatas dibandingkan dengan Bahasa Inggris, sehingga diperlukan model dan *preprocessing* yang disesuaikan secara lokal.

2.1.4 *Bidirectional Encoder Representations from Transformers (BERT)*

Bidirectional Encoder Representations from Transformers (BERT) merupakan metode *machine learning* berbasis arsitektur *Transformer* yang digunakan untuk *pre-training* dalam pemrosesan bahasa alami (NLP), dan dikembangkan oleh Google. BERT pertama kali diperkenalkan oleh Devlin dan rekan-rekannya pada tahun 2019 (Devlin dkk., 2019), Google sendiri telah menerapkan BERT untuk memahami maksud dari penelusuran pengguna secara lebih akurat. Sebelumnya, teknik representasi kata seperti *GloVe*, *word2vec*, dan *FastText* digunakan, namun masih memiliki keterbatasan. Sebagai solusi atas permasalahan tersebut, hadir *ELMo (bidirectional LSTM)* yang memanfaatkan konteks kata sebelum dan sesudah untuk membentuk *word embedding*. Selanjutnya, BERT diperkenalkan dengan performa yang lebih unggul dibanding model-model sebelumnya. BERT dirancang untuk mempelajari representasi dua arah dari teks yang tidak diberi label, dengan memperhatikan konteks di sebelah kiri dan kanan setiap kata pada semua lapisan model. Model BERT yang telah melalui tahap *pre-*

training ini kemudian dapat disesuaikan (*fine-tuning*) untuk berbagai tugas NLP seperti klasifikasi teks, penjawaban pertanyaan, dan terjemahan, hanya dengan menambahkan satu lapisan output tambahan (Devlin dkk., 2019).

Menurut (Devlin dkk., 2019) BERT model *architecture* dibentuk dari *multilayer bi-directional Transformer encoder*. terdapat 2 bentuk model:



Gambar 2.1 Arsitektur BERT (Devlin dkk., 2019)

1. BERT BASE: model ini dibangun dari 12 *Transformer block*, 12 *Attention layer* dan 768 *hidden layers*.
2. BERT LARGE: model ini memiliki *layer* dan *attention layer* yang lebih banyak dari BERT BASE untuk mendapatkan hasil yang lebih baik yaitu 24 *transformer block*, 16 *attention head* dan 1024 *hidden layer*.

BERT dapat dimanfaatkan dalam dua tahap, yaitu *pre-training* dan *fine-tuning*. Pada tahap *pre-training*, model dilatih menggunakan data tanpa label untuk menyelesaikan berbagai tugas awal. Sedangkan pada tahap *fine-tuning*, BERT memulai dengan parameter yang sudah diperoleh dari proses pelatihan sebelumnya,

lalu seluruh parameter tersebut disesuaikan kembali menggunakan data berlabel. Proses ini memungkinkan BERT untuk menyelesaikan berbagai tugas NLP seperti klasifikasi teks, penjawaban pertanyaan, dan pengenalan entitas bernama (*Named Entity Recognition*) (Devlin dkk., 2019).

a. *Pre-Training*

Pre-training BERT dapat digunakan untuk menyelesaikan tugas unsupervised seperti *Masked model language* dan *next sentence prediction* (Devlin dkk., 2019).

1. *Masked language Model*

Masked Language Model digunakan untuk menyembunyikan atau menutupi kata-kata acak yang tidak mungkin dalam kalimat. [MASK] token digunakan untuk menggantikan 15% kata dalam setiap urutan kata sebelum dimasukkan ke BERT. Model kemudian akan mencoba memprediksi kata tersebut nilai asli menggunakan konteks yang disediakan oleh kata lain yang tidak ditutup dengan [MASK] dalam urutan kata.

2. *Next Sentence Prediction*

Next Sentence Prediction (NSP) merupakan salah satu metode yang digunakan untuk memprediksi apakah suatu kalimat merupakan kelanjutan dari kalimat sebelumnya. Dalam proses pelatihan BERT, model diberikan pasangan kalimat dan dilatih untuk menentukan apakah kalimat kedua benar-benar merupakan kelanjutan dari kalimat pertama dalam dokumen asli. Selama pelatihan, 50% dari pasangan kalimat tersebut memang berasal

dari urutan sebenarnya dalam dokumen, sementara 50% sisanya terdiri dari kalimat kedua yang diambil secara acak dari korpus. Asumsinya, kalimat yang dipilih secara acak tidak memiliki kesinambungan makna dengan kalimat pertama.

Dalam tahap *pre-training*, BERT memanfaatkan arsitektur *transformer* yang sangat andal. Arsitektur ini memungkinkan model untuk mengatasi tantangan dalam memahami hubungan kata-kata yang berjauhan dalam suatu teks melalui penggunaan *attention mechanism* yang efisien. Pada struktur *transformer*, terdapat beberapa lapisan yang disebut *encoder*, yang berfungsi untuk menerima input berupa teks dan menghasilkan representasi kontekstual yang mendalam dan bermakna.

Setelah melalui tahap *pre-training*, BERT dapat disesuaikan lebih lanjut melalui proses *fine-tuning* untuk menangani tugas-tugas spesifik. *Fine-tuning* ini melibatkan pelatihan ulang model menggunakan data yang telah diberi label, sesuai dengan kebutuhan tugas tertentu. Sebagai contoh, model BERT yang sebelumnya telah dilatih secara umum dapat diadaptasi untuk berbagai keperluan seperti analisis sentimen, pemrosesan bahasa alami, pengenalan entitas bernama (*Named Entity Recognition*), dan tugas-tugas NLP lainnya.

Tahap *pre-training* pada BERT memungkinkan model untuk mempelajari representasi teks yang lebih mendalam serta memahami struktur dan konteks bahasa secara lebih efektif. Kemampuan ini telah

memberikan kemajuan yang signifikan dalam berbagai jenis tugas pemrosesan bahasa alami.

b. *Fine-Tuning*

Model *pre-training* BERT yang telah dilatih dapat digunakan untuk menyelesaikan tugas NLP lainnya dengan menambahkan *layer* yang sesuai dengan tugas yang ingin di selesaikan pada model yang ada. BERT dapat digunakan untuk klasifikasi teks, *sentiment analysis*, *question-answering*, *named entity recognition* dan tugas *natural language* lainnya (Devlin dkk., 2019).

Model BERT yang telah melalui pelatihan umum kemudian dimuat kembali, dan bagian lapisan akhirnya disesuaikan agar cocok dengan tugas tertentu. Biasanya, lapisan akhir BERT yang dikenal sebagai *classifier layer* akan diganti dengan lapisan klasifikasi baru yang sesuai dengan jumlah kelas pada tugas yang ingin diselesaikan. Setelah itu, model dilatih menggunakan metode pembelajaran terawasi, seperti *backpropagation* dan *stochastic gradient descent*, untuk mengurangi tingkat kesalahan dalam menyelesaikan tugas yang ditargetkan (Devlin dkk., 2019).

2.1.4.1 IndoBERT

IndoBERT merupakan pre-trained model yang dilatih berdasarkan arsitektur BERT khusus dalam dataset Bahasa Indonesia yang disebut INDOLEM. Dataset INDOLEM sendiri terdiri dari beberapa tugas untuk Bahasa Indonesia, seperti morfo-sintaks yang mengarah pada tugas grammatical atau aturan bahasa (terdiri dari NER dan POS Tagging), semantik, dan wacana. Model IndoBERT dilatih dan dievaluasi menggunakan dataset dari INDOLEM dan mengungguli hasil

percobaan dengan algoritma lainnya yaitu MBERT, dan Bi-LSTM-CRF (Koto dkk., 2021).

2.1.4.2 Multilingual BERT

Multilingual BERT (MBERT) merupakan pengembangan dari model bahasa tunggal BERT yang telah dilatih menggunakan *corpora monolingual* dalam 104 bahasa. MBERT disempurnakan menggunakan data pelatihan khusus dari satu bahasa dan dilakukan proses evaluasi dalam bahasa yang berbeda, sehingga memungkinkan untuk digunakan dalam lintas bahasa, dan bahkan MBERT mampu melakukan tugas generalisasi lintas bahasa dengan baik. Model MBERT juga sudah dilatih menggunakan kamus dalam Bahasa Indonesia sehingga dapat digunakan untuk tugas dalam Bahasa Indonesia (Koto dkk., 2021).

2.1.4.3 IndoROBERTa

IndoROBERTa merupakan model bahasa berbasis arsitektur RoBERTa yang telah disesuaikan untuk bahasa Indonesia. RoBERTa sendiri adalah singkatan dari "*Robustly Optimized BERT Pretraining Approach*", yang merupakan pengembangan dari model BERT dengan optimasi pada proses *pretraining*, seperti penggunaan data yang lebih besar, penghapusan token segmentasi, dan pelatihan dengan *batch size* yang lebih besar.

Dalam konteks bahasa Indonesia, IndoROBERTa dikembangkan untuk mengatasi keterbatasan model *multilingual* yang seringkali kurang optimal dalam memahami nuansa bahasa lokal. Model ini dilatih menggunakan korpus besar berbahasa Indonesia, seperti Indo4B, yang mencakup berbagai domain teks untuk

memastikan representasi bahasa yang kaya dan kontekstual (Richardson & Wicaksana, 2022).

2.1.5 Pengujian

Pengujian dengan *Confusion Matrix* adalah langkah krusial dalam membangun model analisis sentimen handal. Dengan memahami konsep ini, bisa memastikan bahwa model *Confusion Matrix* dapat memberikan hasil yang akurat dan dapat diandalkan.

Confusion Matrix adalah sebuah metode yang digunakan dalam evaluasi kinerja model klasifikasi untuk menganalisis seberapa baik model tersebut melakukan klasifikasi pada *dataset* pengujian. *Confusion Matrix* memberikan gambaran tentang seberapa baik atau buruk model dalam mnengklasifikasikan data (Hemmatian, 2019).

Tabel 2.1 Confusion Matrix

<i>Predicted Values</i>	<i>Actual Values</i>	
	<i>True Negative (TN)</i>	<i>False Positif (FP)</i>
	<i>False Negative (FN)</i>	<i>True Positif (TP)</i>

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$$

$$Precision = \frac{TP}{TP+FP}$$

$$Recall = \frac{TP}{TP+FN}$$

Keterangan:

- a. True Positive (TP)

Merupakan jumlah sampel positif yang diklasifikasikan dengan benar sebagai positif oleh model.

b. True Negative (TN)

Merupakan jumlah sampel negatif yang diklasifikasikan dengan benar sebagai negatif oleh model.

c. False Positive (FP)

Merupakan jumlah sampel negatif yang salah diklasifikasikan sebagai positif oleh model.

d. False Negative (FN)

Merupakan jumlah sampel positif yang salah diklasifikasikan sebagai negatif oleh model

2.2 Penelitian Terkait

2.2.1 *State of The Art*

Terdapat beberapa penelitian terdahulu yang dilakukan dalam topik Analisis Sentimen Terhadap Program Jabar *Quick Response* Menggunakan *Algoritma Bidirectional Encoder Representations from Transformers* (BERT). Namun, setiap penelitian menggunakan metode yang unik dan memberikan hasil yang berbeda-beda. Berikut ini penelitian yang terkait dengan penelitian yang dilakukan disajikan pada Tabel 2.2.

Tabel 2.2 *State of The Art*

No	Peneliti	Judul	Permasalahan	Algoritma	Hasil Penelitian
1.	(Sihombing & Situmorang, 2024)	Prediksi Sentimen Pada Teks Media Sosial Corporate University Menggunakan RoBERTa	Menilai bagaimana performa model IndoRoBERTa dibandingkan varian IndoBERT dalam	IndoRoBERTa, IndoBERT	IndoRoBERTa memiliki F1-Score 95.2% dan akurasi 96.2%, lebih tinggi dibandingkan semua model

			klasifikasi sentimen media sosial corporate university.		IndoBERT yang diuji (maks F1-Score IndoBERT = 90.2%)
2.	(Manoppo dkk., 2025)	Analisis Sentimen Publik di Media Sosial Terhadap Kenaikan PPN 12% di Indonesia Menggunakan Indobert	Menganalisis persepsi publik di Indonesia terhadap kebijakan kenaikan Pajak Pertambahan Nilai (PPN) menjadi 12%, dengan fokus pada pengguna platform X (sebelumnya Twitter)	IndoBERT	Secara kuantitatif, model mencapai <i>accuracy</i> 84,94%, <i>precision</i> 85,60%, <i>recall</i> 84,94%, dan <i>F1-score (weighted)</i> 84,37%. Analisis distribusi sentimen lebih lanjut menunjukkan bahwa sentimen publik yang dominan adalah negatif

3.	(Purwati, 2023)	Analisis Sentimen Berita Vaksin Covid-19 Dengan Robustly Optimized BERT Pre-Training APPROACH(ROBERTa)	Pandemi Covid-19 mendorong banyak pihak agar mampu beradaptasi dengan kondisi terkini. Salah satu program yang diluncurkan pemerintah agar dapat mengatasi penyebaran Covid-19 adalah dengan menjalankan program Vaksinasi. Hal ini menunjukkan perlunya analisis mengenai	<i>RoBERTa</i>	RoBERTa menunjukkan performa yang optimal setelah proses pre-training dan fine-tuning dengan pengaturan hyperparameter seperti epoch, batch size, dan learning rate
----	-----------------	--	--	----------------	---

			berita terhadap Vaksinasi Covid-19, untuk mengukur bagaimana animo masyarakat terkait program Vaksinasi Covid-19.		
4.	(Putri & Ardiansyah, 2023)	Analisis Sentimen Terhadap Kemajuan Kecerdasan Buatan di Indonesia Menggunakan BERT dan RoBERTa	Perbedaan persepsi masyarakat terhadap kemajuan AI yang muncul di berbagai platform media sosial, sehingga diperlukan pemetaan sentimen	IndoBERT, IndoRoBERTa	Model indobenchmark/indobert- base-p1 menghasilkan akurasi validasi 84% dan testing 83% pada skema 1, serta validasi 83% dan testing 84% pada skema 2. Sementara RoBERTa

			publik agar dapat dijadikan acuan bagi pemangku kebijakan dalam membuat regulasi pemanfaatan AI.		lebih rendah dengan akurasi maksimal hanya 81% validasi dan 79% testing. Sentimen negatif menjadi yang paling dominan.
5.	(Jaya, 2023)	Analisis Sentimen Pandangan Publik terhadap Profesi PNS dari Twitter Menerapkan Indonesian RoBERTa Base Sentiment Classifier	Banyaknya opini publik terkait profesi PNS di Twitter yang belum dianalisis secara mendalam untuk mengetahui persepsi masyarakat.	Indonesia RoBERTa Base Sentiment Classifier	Sentimen netral mendominasi (86,4%), diikuti negatif (8,6%) dan positif (4,8%). Model mencapai akurasi tinggi hingga 94,36% pada data pelatihan.

6.	(Dwiyansaputra dkk., 2025)	Analisis Sentimen Pengguna Platform Media Sosial X pada Topik Pemilihan Presiden 2024 Menggunakan Perbandingan Model Monolingual dan Multilingual Bert	IndoBERT dan mBERT dipilih sebagai model analisis sentimen karena supremasi mereka dalam memahami konteks linguistik secara menyeluruh melalui penggunaan arsitektur Transformer. IndoBERT disesuaikan untuk bahasa Indonesia dan dapat	IndoBERT MultilingualBERT	Model IndoBERT memiliki akurasi tertinggi yaitu 84% dengan presisi 75%, <i>recall</i> 80%, dan <i>F1-Score</i> sebesar 78%. Sedangkan Mbert mencatatkan akurasi 81%, presisi 69%, <i>recall</i> 78%, dan <i>F1-Score</i> 73%.
----	-------------------------------	---	--	------------------------------	---

			menghasilkan temuan analisis yang lebih tepat dalam memahami ciri khas bahasa tersebut, sebagai model multilingual mBERT memberikan fleksibilitas untuk mengevaluasi teks dalam banyak bahasa.		
7.	(Adinda Nur Ashifa dkk., 2025)	Analisis Sentimen Tanggapan Masyarakat terhadap Layanan	Banyaknya keluhan masyarakat terhadap layanan kesehatan di	RoBERTa	Sentimen netral tertinggi pada antrean (0.53), positif tertinggi pada fasilitas kesehatan (0.45),

		Kesehatan di Kota Surabaya dengan RoBERTa	Surabaya yang disampaikan lewat media sosial, namun belum dianalisis secara sistematis.		dan negatif tertinggi pada tenaga kesehatan (0.33). Akurasi model: 87,2%.
8.	(Zain dkk., 2021)	Analisis Sentimen Pendapat Masyarakat Mengenai Vaksin Covid-19 Pada Media Sosial Twitter dengan Robustly Optimized BERT Pretraining Approach	Perlunya pemahaman terhadap opini masyarakat terkait vaksin Covid-19 agar program vaksinasi nasional dapat berjalan lebih efektif, mengingat banyaknya informasi	Indonesian RoBERTa Base Sentiment Classifier	Dari 109.202 tweet, 69,6% bersentimen netral, 24,7% negatif, dan hanya 5,7% positif. Akurasi prediksi untuk label positif 84%, netral 97%, dan negatif 93%. Analisis menunjukkan bahwa sentimen negatif dominan dan terkait dengan ketidakpercayaan

			dan pendapat di media sosial yang beragam.		terhadap vaksin dan efek sampingnya.
9.	(Fatmasari dkk., 2024)	Implementasi Algoritma BERT Pada Komentar Layanan Akademik dan Non Akademik Universitas Terbuka di Media Sosial	Media sosial TikTok dan Twitter (X) seringkali menjadi tempat untuk menyampaikan pendapat atau komentar terhadap suatu hal. Universitas Terbuka merupakan suatu kampus yang memiliki media sosial TikTok dan Twitter (X) dengan	IndoBERT	Penelitian ini dilakukan untuk mengetahui tingkat layanan bidang akademik dan non-akademik Universitas Terbuka. Data yang dianalisis sebanyak 685 data komentar pada media sosial TikTok dan Twitter (X) dengan kata kunci Universitas Terbuka. Metode yang digunakan adalah analisis menggunakan model pre-trained BERT. Pada model

			ribuan pengikut. Penelitian ini dilakukan untuk mengetahui tingkat layanan bidang akademik dan non-akademik Universitas Terbuka.		ini diperoleh nilai akurasi sebesar 90% dengan proporsi data latih dan data uji 80:20.
10.	(Putu dkk., 2023)	Analisis Sentimen terhadap Perundungan Siber pada Twitter menggunakan Algoritma <i>Bidirectional Encoder Representations from Transformer</i> (BERT)	Perundungan siber adalah penggunaan teknologi informasi dan komunikasi untuk pelecehan yang disengaja/terencana dan aktivitas	IndoBERT	Tingkat akurasi pada klasifikasi perundungan siber dengan menggunakan data uji dan data validasi dalam Bahasa Indonesia yang dihasilkan dengan menggunakan algoritma BERT

			permusuhan, yang dilakukan berulang dan terstruktur. Perundungan siber termasuk intimidasi seseorang melalui media sosial, pelecehan, sexting, penipuan, meniru, dan mengirim pesan jahat melalui ruang obrolan dan pesan instan. Hal ini menunjukkan perlunanya analisis		sebesar 0,81 yang dimana dalam persentase akurasiya adalah 81%.
--	--	--	---	--	--

			terhadap Perundungan Siber pada Twitter.		
11.	(Nelly Sofi dkk., 2023)	Analisis Sentimen Masyarakat Pengguna Media Sosial Twitter Terhadap Motogp Mandalika Lombok Menggunakan Metode <i>Bidirectional Encoder Representation from Transformers</i> (BERT)	Balapan MotoGP Satu di Nusa Tenggara Barat Lombok, Mandalika yang diselenggarakan pada 18 Maret 2022, mendapatkan banyak tanggapan ataupun teaksi dari masyarakat di media social terutama Twitter. Tanggapan-tanggapan tersebut ada yang setuju	IndoBERT	Hasil evaluasi dari model tersebut mendapatkan akurasi sebesar 55%. Precision untuk positif sebesar 56%, netral sebesar 59%, dan negatif sebesar 44%. Recall untuk positif sebesar 74%, netral sebesar 29%, dan negatif 54%. F1-score untuk positif sebesar 64%, netral sebesar 38%, dan negatif sebesar 48%.

			dan tidak mengenai penyelenggaraan MotoGP di Mandalika ini, untuk mengetahui tanggapan masyarakat yang setuju atau tidak diperlukan sistem yang dapat mengolah data tweets dengan metode analisis sentimen.		
12.	(Adrinta Abdurrazzaq & Lesmana Tjiong, 2022)	Analisis Sentimen KUHP Baru Pada Data Twitter Menggunakan BERT	RUU KUHP dianggap perlu disahkan untuk mengganti undang – undang yang lama yang	IndoBERT, SVM	Hasil analisis dengan Model BERT mampu mencapai akurasi sebesar 81%. Hasil ini unggul 6% dibandingkan

			masih berbasis hukum kolonial, tetapi pengesahan undang – undang ini menuai berbagai sentimen dan kritik. Masyarakat menilai banyak pasal – pasal kontroversial di dalamnya, padahal dalam perjalanan perancangannya masyarakat telah berulang kali mengkritik dan		dengan model analisis sentimen menggunakan Support Vector Machine.
--	--	--	---	--	---

			mempertanyakan pasal – pasal tersebut.		
13.	(Kusuma & Mogi, 2023)	Implementasi BERT pada Analisis Sentimen Ulasan Destinasi Wisata Bali	Dalam beberapa tahun terakhir, kontribusi sektor pariwisata di Bali meningkat signifikan. Opini publik dan ulasan tentang destinasi wisata ini dapat digunakan untuk mengidentifikasi destinasi wisata baru yang sedang naik daun dan diminati. Itulah	IndoBERT	Hasilnya menunjukkan bahwa analisis sentimen menggunakan model IndoBERT dengan optimizer AdamW mencapai akurasi 97% dan AdaFactor mencapai akurasi 98,2%.

			sebabnya penting untuk memanfaatkan opini positif atau negatif ini untuk memperoleh informasi menarik dan penting tentang destinasi wisata ini untuk digunakan lebih lanjut.		
14.	(Akhmad, 2023)	Analisis Sentimen Ulasan Aplikasi DLU Ferry Pada Google Play Store Menggunakan Bidirectional Encoder	DLU Ferry merupakan aplikasi yang dikeluarkan oleh PT. Dharma Lautan Utama untuk memudahkan	IndoBERT	Hasil pengujian pada penelitian memperoleh akurasi sebesar 86% dengan pemilihan hyperparameter, yaitu

		Representations from Transformers	pelanggan dalam melakukan pemesanan tiket. Aplikasi DLU Ferry masih memiliki kelemahan, oleh karena itu perlu perbaikan untuk meningkatkan kualitas aplikasi ini. Untuk mengetahui kelemahan aplikasi ini dapat diperoleh dari ulasan yang ditulis konsumen pada Google Play Store.		batch size 32, learning rate $3e-6$, dan epoch 5.
--	--	-----------------------------------	--	--	--

15.	(Farida & Rochmawati, 2024)	Analisis Sentimen Masyarakat terhadap Fenomena Childfree Menggunakan Metode <i>Long Short-Term Memory dan Bidirectional Encoder Representations from Transformers</i> di Twitter	Twitter menjadi tempat yang ramai ketika terdapat isu terkini di negara ini bahkan dunia, beberapa isu terkini menjadi perhatian khusus oleh masyarakat, salah satunya tentang isu fenomena childfree, sehingga penelitian ini dilakukan dengan tujuan untuk mengetahui analisis	LSTM, IndoBERT	Pada pengujian model dilakukan dengan pengujian beberapa parameter untuk mendapatkan model yang optimal diantaranya ukuran batch size 128, dropout 0.5, dense layer 32 dan lstm layer 64. Hyperparameter tersebut dilakukan pelatihan model yang menghasilkan performa terbaik dengan akurasi sebesar 0.9585, f1-score 0.9589, loss 0.1001, kemudian model dapat memprediksi
-----	-----------------------------	--	--	----------------	--

			sentimen terhadap fenomena tersebut.		tweets childfree dan menghasilkan precision sebesar 0.7839, recall 0.77, dan f1-score 0.7697.
--	--	--	---	--	--

2.2.2 Matriks Penelitian

Tabel 2.3 Matriks Penelitian

No	Nama Penulis dan Tahun Penelitian	Judul Penelitian	Algoritma		
			INDORO BERTA	INDO BERT	MULTILINGUAL BERT
1.	(Sihombing & Situmorang, 2024)	Prediksi Sentimen Pada Teks Media Sosial Corporate University Menggunakan RoBERTa	✓	✓	
2.	(Manoppo dkk., 2025)	Analisis Sentimen Publik di Media Sosial Terhadap Kenaikan PPN 12% di Indonesia Menggunakan Indobert		✓	
3.	(Purwati, 2023)	Analisis Sentimen Berita Vaksin Covid-19 Dengan Robustly Optimized BERT Pre-Training APPROACH(ROBERTa)	✓		

4.	(Putri & Ardiansyah, 2023)	Analisis Sentimen Terhadap Kemajuan Kecerdasan Buatan di Indonesia Menggunakan BERT dan RoBERTa	✓	✓	
5.	(Jaya, 2023)	Analisis Sentimen Pandangan Publik terhadap Profesi PNS dari Twitter Menerapkan Indonesian RoBERTa Base Sentiment Classifier	✓		
6.	(Dwiyansaputra dkk., 2025)	Analisis Sentimen Pengguna Platform Media Sosial X pada Topik Pemilihan Presiden 2024 Menggunakan Perbandingan Model Monolingual dan Multilingual Bert		✓	✓
7.	(Adinda Nur Ashifa dkk., 2025)	Analisis Sentimen Tanggapan Masyarakat terhadap Layanan Kesehatan di Kota Surabaya dengan RoBERTa	✓		

8.	(Zain dkk., 2021)	Analisis Sentimen Pendapat Masyarakat Mengenai Vaksin Covid-19 Pada Media Sosial Twitter dengan Robustly Optimized BERT Pretraining Approach	✓		
9.	(Fatmasari dkk., 2024)	Implementasi Algoritma BERT Pada Komentar Layanan Akademik dan Non Akademik Universitas Terbuka di Media Sosial		✓	
10.	(Putu dkk., 2023)	Analisis Sentimen terhadap Perundungan Siber pada Twitter menggunakan Algoritma <i>Bidirectional Encoder Representations from Transformer</i> (BERT)		✓	
11.	(Nelly Sofi dkk., 2023)	Analisis Sentimen Masyarakat Pengguna Media Sosial Twitter Terhadap Motogp Mandalika Lombok Menggunakan Metode <i>Bidirectional Encoder Representation from Transformers</i> (BERT)		✓	

12.	(Adrinta Abdurrazzaq & Lesmana Tjiong, 2022)	Analisis Sentimen KUHP Baru Pada Data Twitter Menggunakan BERT		✓	
13.	(Kusuma & Mogi, 2023)	Implementasi BERT pada Analisis Sentimen Ulasan Destinasi Wisata Bali		✓	
14.	(Akhmad, 2023)	Analisis Sentimen Ulasan Aplikasi DLU Ferry Pada Google Play Store Menggunakan Bidirectional Encoder Representations from Transformers		✓	
15.	(Farida & Rochmawati, 2024)	Analisis Sentimen Masyarakat terhadap Fenomena Childfree Menggunakan Metode <i>Long Short-Term Memory</i> dan <i>Bidirectional Encoder Representations from Transformers</i> di Twitter		✓	

16.	(Gumelar, 2024)	Perbandingan Model IndoBERT, MultilingualBERT, dan IndoROBERTa dalam Analisis Sentimen Terhadap Program Jabar <i>Quick Response</i>	✓	✓	✓
-----	-----------------	---	---	---	---

Tabel 2.3 Memberikan perbandingan dengan penelitian terdahulu dan yang akan dilakukan. Dari hasil studi literatur mengenai penelitian sebelumnya, dapat disimpulkan hasil dari pembahasan penelitian tersebut hanya melakukan dengan pengujian 1 model. Maka dari itu penelitian yang akan dilakukan membandingkan 3 model yaitu IndoBERT, MultilingualBERT, dan IndoROBERTa dengan mengambil data tentang program Jabar *Quick Response* di media sosial sebagai sumber datanya sehingga dapat diukur performa dan hasil tingkat akurasi yang terbaik dari masing masing 3 model tersebut.