

BAB II

LANDASAN TEORI

2.1 Penelitian Terdahulu

Subbab ini memaparkan landasan teoritis dan kajian penelitian terdahulu yang menjadi dasar konseptual serta metodologis dalam penelitian ini. Pembahasan diawali dengan telaah terhadap berbagai studi yang relevan dengan persepsi publik terhadap *artificial intelligence*, khususnya dalam konteks media sosial, *generative AI*, dan dunia kerja. Kajian ini juga mencakup penelitian yang menggunakan analisis sentimen kontekstual, *topic Modeling*, serta pendekatan klasifikasi teks untuk memahami opini publik pada platform digital seperti X/Twitter. Selanjutnya, penelitian terdahulu dianalisis secara komparatif untuk melihat perbedaan fokus kajian, data yang digunakan, metode atau model yang diterapkan, serta temuan yang dihasilkan. Melalui kajian tersebut, dapat dirumuskan posisi penelitian ini dalam peta riset yang telah ada, sekaligus menegaskan kontribusi penelitian, yaitu pemetaan sentimen kontekstual terhadap AI dalam konteks pekerjaan berbasis pengetahuan pada data X/Twitter berbahasa Indonesia berdasarkan kategori ancaman dan peluang, serta perbandingan performa model TF-IDF dengan SVM, IndoBERT, dan IndoBERTweet dalam proses klasifikasi.

Tabel 2. 1 Penelitian Terdahulu

Referensi ke-1	
Judul	“Whose AI? How different <i>publics</i> think about AI and its social impacts”
Nama Penulis	Bao et al. [7].
Sumber Jurnal	<i>Computers in Human Behavior</i> , Vol. 130, Article 107182
Tahun	2022
Permasalahan	Persepsi publik terhadap AI sering disederhanakan menjadi mendukung atau menolak, padahal persepsi terhadap AI dapat terbentuk dari kombinasi risiko dan manfaat.
Framework / Model	Survei nasional terhadap 2.700 responden di Amerika Serikat dengan metode <i>Latent Class Analysis</i> berdasarkan persepsi risiko dan manfaat AI.
Temuan	Ditemukan lima segmen publik: <i>negative</i> 33,3%, <i>ambivalent</i> 28,5%, <i>tepid</i> 24,0%, <i>ambiguous</i> 7,5%, dan <i>indifferent</i> 6,6%.

Pembahasan	Publik dapat melihat AI sebagai teknologi yang berisiko sekaligus bermanfaat, sehingga persepsi terhadap AI tidak cukup dipahami secara biner.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena mendukung dasar konseptual bahwa persepsi terhadap AI dapat dipahami melalui kombinasi risiko dan manfaat, sehingga sejalan dengan kategori ancaman dan peluang dalam penelitian ini. Gapnya, penelitian tersebut menggunakan pendekatan survei pada konteks publik Amerika Serikat, sehingga belum membahas persepsi publik berbasis data media sosial berbahasa Indonesia terhadap AI dalam konteks pekerjaan berbasis pengetahuan.
Referensi ke-2	
Judul	“Public perception of generative AI on Twitter: an empirical study <i>Based on occupation and usage</i> ”
Nama Penulis	Miyazaki et al. [19].
Sumber Jurnal	EPJ <i>Data science</i> , Vol. 13, <i>Article 2</i>
Tahun	2024
Permasalahan	Masih terbatas penelitian kuantitatif yang menganalisis persepsi publik terhadap <i>generative</i> AI di media sosial berdasarkan pekerjaan dan penggunaan.
Framework / Model	Analisis 3.065.972 <i>tweet</i> terkait <i>generative</i> AI dari Januari 2019 hingga Maret 2023, menggunakan analisis sentimen, ekstraksi pekerjaan, dan BERTopic.
Temuan	Sentimen terhadap <i>generative</i> AI secara umum positif. Namun, beberapa kelompok seperti ilustrator menunjukkan sentimen negatif karena isu penggunaan karya seni dalam pelatihan AI.
Pembahasan	Persepsi terhadap <i>generative</i> AI berbeda antar pekerjaan dan dipengaruhi oleh konteks penggunaan serta dampak yang dirasakan.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena sama-sama menggunakan data Twitter/X untuk menganalisis persepsi publik terhadap <i>generative</i> AI serta memperhatikan keterkaitan persepsi AI dengan pekerjaan atau <i>occupation</i> pengguna. Gapnya, penelitian tersebut berfokus pada <i>generative</i> AI secara umum dan data Twitter berbahasa Inggris, sedangkan penelitian ini berfokus pada data X/Twitter berbahasa Indonesia dalam konteks pekerjaan berbasis pengetahuan dengan kategori ancaman, peluang, dan netral.
Referensi ke-3	
Judul	“Public Attitudes Toward ChatGPT on Twitter: Sentiments, Topics, and Occupations”
Nama Penulis	Koonchanok et al. [20]
Sumber Jurnal	arXiv:2306.12951v2
Tahun	2024

Permasalahan	Penelitian sebelumnya tentang opini publik terhadap ChatGPT di Twitter masih banyak menggunakan rentang waktu pendek dan belum banyak melihat topik serta pekerjaan pengguna.
<i>Framework / Model</i>	Analisis 314.645 <i>tweet</i> #ChatGPT dari 5 Desember 2022 hingga 10 Juni 2023 menggunakan XLM-T untuk sentimen, BERTopic untuk topik, dan ekstraksi pekerjaan dari deskripsi pengguna.
Temuan	Sentimen publik terhadap ChatGPT cenderung netral hingga positif: 35,28% positif, 45,79% netral, dan 18,92% negatif. Topik dominan meliputi <i>Education</i> , <i>Bard</i> , <i>Search Engines</i> , <i>OpenAI</i> , <i>Marketing</i> , dan <i>Cybersecurity</i> .
Pembahasan	Diskusi publik tentang ChatGPT di Twitter bersifat dinamis, dengan variasi sentimen dan topik yang berubah dari waktu ke waktu.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena menganalisis sikap publik terhadap ChatGPT di Twitter melalui analisis sentimen, <i>topic Modeling</i> , dan <i>occupation analysis</i> , sehingga dapat menjadi pembandingan dalam memahami variasi sentimen dan topik percakapan AI di media sosial. Gapnya, penelitian tersebut berfokus pada ChatGPT dan menggunakan <i>topic Modeling</i> , sedangkan penelitian ini berfokus pada AI dalam pekerjaan berbasis pengetahuan serta menggunakan analisis frasa dominan pada setiap kategori sentimen kontekstual, bukan <i>topic Modeling</i> .
Referensi ke-4	
Judul	“Powerful tool or too powerful? Early <i>public</i> discourse about ChatGPT across 4 million tweets”
Nama Penulis	Ng dan Chow [31].
Sumber Jurnal	<i>PLOS ONE</i> , Vol. 19, No. 3, e0296882
Tahun	2024
Permasalahan	Masih terbatas studi <i>big data</i> yang membahas diskursus awal publik terhadap ChatGPT, khususnya terkait isu, tema, dan sentimen yang muncul setelah peluncurannya.
<i>Framework / Model</i>	Analisis 4.251.662 <i>tweet</i> berbahasa Inggris tentang ChatGPT pada periode 1 Desember 2022 hingga 1 Maret 2023 menggunakan <i>prominent peaks modelling</i> dan analisis sentimen berbasis kata kunci.
Temuan	Ditemukan enam tema utama, yaitu <i>hype</i> and <i>hesitance</i> , penggunaan dalam profesional dan akademik, bias demografis, eksperimen moral, serta AI sebagai cerminan pengetahuan manusia. Sentimen publik bersifat ganda, yaitu antusias terhadap manfaat ChatGPT tetapi juga waspada terhadap potensi penyalahgunaan.
Pembahasan	Diskursus awal ChatGPT menunjukkan pola <i>double-edged</i> , yaitu ChatGPT dipandang sebagai alat yang kuat dan

	bermanfaat, tetapi juga menimbulkan kekhawatiran terkait kredibilitas, bias, etika, dan dampak profesional.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena menunjukkan bahwa diskursus Twitter dapat dianalisis melalui lonjakan percakapan dan kata kunci berbasis sentimen untuk memahami isu yang dominan dalam percakapan publik. Gapnya, penelitian tersebut menganalisis diskursus awal ChatGPT pada <i>tweet</i> berbahasa Inggris menggunakan <i>prominent peaks Modeling</i> dan analisis kata kunci, sedangkan penelitian ini menganalisis data X/Twitter berbahasa Indonesia dalam konteks pekerjaan berbasis pengetahuan dengan kategori ancaman, peluang, dan netral serta menggunakan frasa dominan untuk menginterpretasikan isu pada tiap kategori sentimen.
Referensi ke-5	
Judul	“The ChatGPT bot is causing panic now – but it’ll soon be as mundane a tool as Excel”: analysing topics, sentiment and emotions relating to ChatGPT on Twitter
Nama Penulis	Heaton, Clos, Nichele, dan Fischer [22].
Sumber Jurnal	<i>Personal and Ubiquitous Computing</i> , Vol. 28, pp. 875–894
Tahun	2024
Permasalahan	Penelitian sebelumnya tentang ChatGPT di Twitter masih banyak berfokus pada reaksi awal dan belum banyak menggabungkan <i>topic Modeling</i> , <i>sentiment analysis</i> , dan <i>emotion detection</i> secara bersamaan.
Framework / Model	Analisis 88.058 <i>tweet</i> tentang ChatGPT dari November 2022 hingga Maret 2023 menggunakan LDA untuk <i>topic Modeling</i> , VADER untuk <i>sentiment analysis</i> , dan EmoLex untuk <i>emotion detection</i> .
Temuan	Hasil penelitian menemukan tujuh topik utama dalam diskursus ChatGPT, yaitu <i>text generation</i> , <i>chatbot development</i> , <i>writing assistance</i> , <i>data training</i> , efisiensi penggunaan, dampak terhadap bisnis, dan <i>cryptocurrency</i> . Sentimen secara umum cenderung positif dengan skor VADER berkisar antara 0,21 hingga 0,31, tetapi mengalami penurunan sekitar peluncuran ChatGPT Plus. Pada aspek emosi, pola yang dominan adalah <i>trust</i> dan <i>fear</i> , dengan penurunan <i>trust</i> yang berkaitan dengan kekhawatiran terhadap bias dan misinformasi.
Pembahasan	Hasil penelitian menunjukkan bahwa kekhawatiran terhadap ChatGPT memang muncul, terutama terkait bias, misinformasi, etika, dan dampak sosial, tetapi tidak mendominasi keseluruhan diskursus Twitter.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena menggabungkan analisis sentimen, <i>topic Modeling</i> , dan <i>emotion detection</i> untuk memahami diskursus ChatGPT di Twitter, termasuk adanya sentimen positif serta kekhawatiran terkait bias, misinformasi, etika, dan dampak sosial. Gapnya, penelitian tersebut berfokus

	pada ChatGPT dan <i>tweet</i> berbahasa Inggris, sedangkan penelitian ini berfokus pada AI dalam pekerjaan berbasis pengetahuan pada data X/Twitter berbahasa Indonesia serta menginterpretasikan kecenderungan isu melalui frasa dominan pada masing-masing kategori sentimen.
Referensi ke-6	
Judul	“Sentiment analysis <i>Methods</i> , applications, and challenges: A Systematic Literature Review”
Nama Penulis	Mao, Liu, dan Zhang [23]
Sumber Jurnal	<i>Journal of King Saud University - Computer and Information Sciences</i> , Vol. 36, Article 102048
Tahun	2024
Permasalahan	Literatur analisis sentimen sebelumnya masih terbatas pada metode atau domain tertentu, sehingga diperlukan tinjauan yang lebih menyeluruh mengenai metode, aplikasi, dan tantangan analisis sentimen.
Framework / Model	<i>Systematic Literature Review</i> berbasis PRISMA dengan meninjau artikel dari <i>Scopus</i> , <i>Web of Science</i> , ScienceDirect, IEEE, ACM, Springer, dan <i>Google Scholar</i> .
Temuan	Analisis sentimen dapat dilakukan dengan pendekatan <i>Lexicon-Based</i> , <i>machine learning</i> , <i>deep learning</i> , <i>hybrid</i> , <i>transformer</i> , dan <i>transfer learning</i> . Studi ini juga menjelaskan bahwa metrik umum evaluasi analisis sentimen meliputi <i>accuracy</i> , <i>precision</i> , <i>recall</i> , dan F1-score.
Pembahasan	Analisis sentimen memiliki banyak aplikasi, termasuk media sosial, bisnis, pemerintahan, kesehatan, dan <i>Large Language Models</i> . Tantangan utama meliputi konteks bahasa, slang, sarkasme, data multilingual, dan kebutuhan model yang sesuai dengan domain.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena memberikan dasar metodologis mengenai pendekatan analisis sentimen, tantangan data media sosial, serta penggunaan metrik evaluasi seperti <i>accuracy</i> , <i>precision</i> , <i>recall</i> , dan F1-score. Gapnya, penelitian tersebut merupakan <i>Systematic Literature Review</i> yang bersifat umum, sehingga tidak secara empiris menganalisis sentimen kontekstual terhadap AI dalam pekerjaan berbasis pengetahuan pada data X/Twitter berbahasa Indonesia.
Referensi ke-7	
Judul	“The Implementation of TF-IDF and Word2Vec on Booster Vaccine Sentiment Analysis Using <i>Support Vector Machine Algorithm</i> ”
Nama Penulis	Agustina, Novita, Mustakim, dan Rozanda [28].
Sumber Jurnal	<i>Procedia Computer Science</i> , Vol. 234, pp. 156–163
Tahun	2024
Permasalahan	Pemilihan metode ekstraksi fitur berpengaruh terhadap performa klasifikasi sentimen, sehingga perlu

	dibandingkan efektivitas TF-IDF dan Word2Vec pada algoritma SVM.
<i>Framework / Model</i>	Analisis sentimen terhadap 13.297 <i>tweet</i> Indonesia bertema “vaksin <i>booster</i> ” menggunakan SVM dengan fitur TF-IDF dan Word2Vec, serta perbandingan <i>kernel</i> linear, RBF, dan <i>polynomial</i> .
Temuan	Kombinasi TF-IDF dengan SVM menggunakan <i>Data Split</i> 80:20 dan <i>kernel</i> RBF menghasilkan performa terbaik pada bagian hasil/kesimpulan, yaitu <i>precision</i> 84%, <i>recall</i> 85%, dan <i>F1-score</i> 83%.
Pembahasan	TF-IDF dengan SVM menunjukkan performa lebih baik dibandingkan Word2Vec dengan SVM. Hasil ini menunjukkan bahwa representasi fitur sederhana seperti TF-IDF masih efektif untuk klasifikasi sentimen berbasis Twitter.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena menunjukkan bahwa TF-IDF dengan SVM masih efektif digunakan dalam klasifikasi sentimen berbasis data Twitter berbahasa Indonesia. Gapnya, penelitian tersebut berfokus pada sentimen vaksin <i>booster</i> dan membandingkan fitur TF-IDF serta Word2Vec, sedangkan penelitian ini menggunakan TF-IDF dengan SVM sebagai <i>baseline</i> klasik untuk dibandingkan dengan model berbasis <i>transformer</i> pada konteks AI dan pekerjaan berbasis pengetahuan.
Referensi ke-8	
Judul	“Sentiment analysis classification system using hybrid BERT models”
Nama Penulis	Amira Samy Talaat [32].
Sumber Jurnal	<i>Journal of Big data</i> , Vol. 10, <i>Article</i> 110
Tahun	2023
Permasalahan	Model analisis sentimen berbasis BERT sudah efektif, tetapi akurasi masih dapat ditingkatkan, khususnya untuk data media sosial yang pendek dan mengandung emoji.
<i>Framework / Model</i>	Menggunakan model <i>hybrid</i> berbasis DistilBERT dan RoBERTa yang dikombinasikan dengan BiGRU dan BiLSTM. Model diuji pada tiga dataset: <i>Airlines</i> , <i>CrowdFlower</i> , dan <i>Apple</i> .
Temuan	Model <i>hybrid</i> berbasis BiGRU menunjukkan performa terbaik. RoBERTa secara umum menghasilkan akurasi lebih tinggi dibandingkan DistilBERT, dengan akurasi tertinggi 91,72% pada dataset <i>Apple</i> .
Pembahasan	Penggabungan BERT dengan BiGRU/BiLSTM dapat meningkatkan performa klasifikasi sentimen. Hasil juga menunjukkan bahwa keberadaan emoji dapat memengaruhi performa model.

Relevansi dan Gap terhadap Penelitian Ini	Relevan karena mendukung penggunaan model berbasis BERT atau <i>transformer</i> dalam analisis sentimen media sosial serta menunjukkan potensinya dibandingkan pendekatan konvensional. Gapnya, penelitian tersebut menggunakan model <i>hybrid</i> berbasis BERT pada dataset umum dan tidak berfokus pada data X/Twitter berbahasa Indonesia maupun konteks sentimen AI dalam pekerjaan berbasis pengetahuan.
Referensi ke-9	
Judul	“IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding”
Nama Penulis	Wilie et al. [29].
Sumber Jurnal	<i>Proceedings of AACL-IJCNLP</i> , pp. 843–857
Tahun	2020
Permasalahan	Penelitian NLP bahasa Indonesia masih terhambat karena keterbatasan dataset, benchmark, dan model pralatih yang dapat digunakan secara luas.
Framework / Model	Mengembangkan IndoNLU sebagai benchmark untuk 12 tugas NLU bahasa Indonesia dan memperkenalkan IndoBERT yang dilatih menggunakan dataset Indo4B.
Temuan	IndoBERT menunjukkan performa kuat pada berbagai tugas bahasa Indonesia. Pada tugas sentimen SmSA, IndoBERT- <i>base</i> memperoleh skor 87,73, sedangkan IndoBERT- <i>large</i> memperoleh skor 92,72.
Pembahasan	Model pralatih monolingual seperti IndoBERT mampu menangkap karakteristik bahasa Indonesia dengan baik, termasuk pada tugas klasifikasi teks dan analisis sentimen.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena menjadi dasar penggunaan IndoBERT sebagai model pralatih bahasa Indonesia yang telah dievaluasi pada berbagai tugas pemahaman bahasa, termasuk analisis sentimen. Gapnya, penelitian tersebut berfokus pada pengembangan benchmark dan sumber daya NLP bahasa Indonesia, bukan pada analisis sentimen kontekstual terhadap AI dalam pekerjaan berbasis pengetahuan pada data X/Twitter.
Referensi ke-10	
Judul	“IndoBERTweet: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization”
Nama Penulis	Koto, Lau, dan Baldwin [30].
Sumber Jurnal	<i>Proceedings of EMNLP</i>
Tahun	2021
Permasalahan	Model bahasa umum belum tentu optimal untuk data Twitter karena teks media sosial memiliki karakteristik khusus seperti bahasa informal, singkatan, <i>mention</i> , URL, dan kosakata domain tertentu.

<i>Framework / Model</i>	Mengembangkan IndoBERTweet sebagai model pralatih khusus Twitter Indonesia menggunakan 26 juta <i>tweet</i> atau sekitar 409 juta token, lalu dievaluasi pada tujuh dataset Twitter Indonesia.
Temuan	IndoBERTweet memperoleh rata-rata performa 86,1, sedangkan adaptasi kosakata dan <i>domain-adaptive pretraining</i> dengan metode <i>average of subwords</i> mencapai rata-rata terbaik sebesar 86,5.
Pembahasan	Model yang dilatih atau diadaptasi pada domain Twitter dapat meningkatkan performa pada tugas-tugas NLP berbasis Twitter dibandingkan model umum.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena IndoBERTweet dikembangkan khusus untuk data Twitter berbahasa Indonesia, sehingga sesuai digunakan sebagai model pembanding pada data media sosial. Gapnya, penelitian tersebut berfokus pada pengembangan model bahasa untuk Twitter Indonesia, sedangkan penelitian ini menggunakan IndoBERTweet untuk mengklasifikasikan sentimen kontekstual ancaman, peluang, dan netral terhadap AI dalam konteks pekerjaan berbasis pengetahuan.

Berdasarkan sepuluh penelitian terdahulu yang telah dipaparkan, kajian yang relevan dengan penelitian ini dapat dipetakan ke dalam tiga kelompok utama. Kelompok pertama berkaitan dengan kajian persepsi publik terhadap *artificial intelligence* yang menunjukkan bahwa AI tidak hanya dipandang sebagai teknologi yang membawa manfaat, tetapi juga dapat dipersepsikan sebagai risiko [7]. Penelitian dalam kelompok ini menjadi dasar konseptual bahwa persepsi terhadap AI dapat dibaca melalui dimensi ancaman dan peluang. Kelompok kedua berkaitan dengan penelitian berbasis X/Twitter yang menganalisis diskursus publik mengenai *generative AI* dan ChatGPT melalui analisis sentimen, *topic Modeling*, analisis kata kunci, maupun analisis pekerjaan pengguna [19], [20], [31], [22]. Penelitian-penelitian tersebut menunjukkan bahwa media sosial dapat digunakan untuk memahami dinamika opini publik terhadap AI, tetapi sebagian besar masih berfokus pada data berbahasa Inggris, ChatGPT atau *generative AI* secara umum, serta belum secara khusus menganalisis percakapan X/Twitter berbahasa Indonesia dalam konteks pekerjaan berbasis pengetahuan. Kelompok ketiga berkaitan dengan pendekatan metodologis dalam klasifikasi sentimen, mulai dari metode klasik seperti TF-IDF dengan SVM hingga model berbasis *transformer* seperti BERT, IndoBERT, dan IndoBERTweet [23], [28], [29], [30], [32]. Kajian dalam kelompok

ini menjadi dasar pemilihan model dalam penelitian, sekaligus menunjukkan perlunya perbandingan antara pendekatan *machine learning* klasik dan model berbasis *transformer* pada data media sosial berbahasa Indonesia.

Berdasarkan pemetaan tersebut, posisi penelitian ini terletak pada penggabungan konteks, kategori sentimen, metode interpretasi, dan perbandingan model. Dari sisi konteks, penelitian ini berfokus pada percakapan X/Twitter berbahasa Indonesia mengenai AI dalam pekerjaan berbasis pengetahuan. Dari sisi kategori, penelitian ini tidak menggunakan polaritas umum positif, negatif, dan netral, melainkan kategori sentimen kontekstual ancaman dan peluang. Dari sisi interpretasi, penelitian ini tidak menggunakan *topic Modeling*, tetapi menggunakan frasa dominan untuk mengidentifikasi kecenderungan isu pada masing-masing kategori sentimen. Dari sisi model, penelitian ini membandingkan TF-IDF dengan SVM, IndoBERT, dan IndoBERTweet untuk melihat performa pendekatan klasik dan *transformer* dalam klasifikasi sentimen kontekstual. Dengan demikian, penelitian ini diposisikan untuk memberikan pemetaan sentimen kontekstual terhadap *artificial intelligence* dalam pekerjaan berbasis pengetahuan pada data X/Twitter berbahasa Indonesia periode 2023–2025, sekaligus mengevaluasi model klasifikasi yang digunakan dalam konteks tersebut.

2.2 Teori yang berkaitan

2.2.1 *Artificial Intelligence* dan Transformasi Dunia Kerja

Artificial intelligence (AI) merupakan cabang ilmu komputer yang berfokus pada pengembangan sistem yang mampu meniru kemampuan kognitif manusia, seperti belajar, memahami bahasa, mengenali pola, serta mengambil keputusan [3]. Dalam perkembangannya, AI tidak lagi terbatas pada sistem berbasis aturan (*rule-Based systems*), tetapi telah berevolusi menjadi sistem berbasis pembelajaran mesin (*Machine Learning*) dan *deep learning* yang mampu memproses data dalam jumlah besar serta menghasilkan prediksi atau rekomendasi secara otomatis [33]. Perkembangan ini semakin pesat sejak munculnya teknologi *generative AI* yang mampu menghasilkan teks, gambar, kode program, dan berbagai bentuk konten digital lainnya secara otonom.

Transformasi yang dibawa oleh AI mencakup aspek teknologis dan struktural dalam dunia kerja, berbagai sektor industri mulai mengintegrasikan AI untuk meningkatkan efisiensi operasional, otomatisasi proses, serta pengambilan keputusan berbasis data. Dalam konteks ini, AI sering dipandang sebagai teknologi disruptif karena kemampuannya menggantikan sebagian tugas yang sebelumnya dilakukan oleh manusia, terutama pekerjaan yang bersifat repetitif dan berbasis aturan [34]. Namun demikian, di sisi lain, AI juga membuka peluang baru melalui penciptaan jenis pekerjaan baru, peningkatan produktivitas, serta kolaborasi antara manusia dan mesin.

Secara konseptual, terdapat dua perspektif utama dalam melihat dampak AI terhadap pekerjaan, yaitu *automation* dan *augmentation*. Perspektif *automation* menekankan bahwa AI berpotensi menggantikan tenaga kerja manusia pada tugas-tugas tertentu, sehingga menimbulkan kekhawatiran terhadap hilangnya lapangan pekerjaan. Sebaliknya, perspektif *augmentation* melihat AI sebagai alat yang memperkuat kapabilitas manusia, membantu proses pengambilan keputusan, serta mendorong peningkatan keterampilan melalui *reskilling* dan *upskilling* [35]. Dalam kerangka teori *Skill-Biased Technological Change*, adopsi teknologi cenderung meningkatkan permintaan terhadap tenaga kerja dengan keterampilan tinggi, sekaligus menekan pekerjaan yang bersifat rutin dan kurang terampil [35]. Dalam konteks Indonesia, transformasi digital yang dipercepat oleh adopsi AI menghadirkan tantangan dan peluang yang tidak merata di berbagai sektor. Perbedaan tingkat literasi digital, kesiapan infrastruktur, serta kapasitas sumber daya manusia dapat memengaruhi bagaimana AI diterima dan dimaknai oleh masyarakat [7]. Oleh karena itu, persepsi publik terhadap AI dalam dunia kerja menjadi aspek yang penting untuk dipahami. Persepsi tersebut tidak hanya dipengaruhi oleh realitas perubahan teknologi, tetapi juga oleh narasi yang berkembang di ruang digital, termasuk media sosial, yang membentuk opini kolektif mengenai apakah AI dipandang sebagai ancaman terhadap pekerjaan, peluang dalam mendukung aktivitas kerja, atau isu yang dibicarakan secara deskriptif tanpa kecenderungan yang kuat. Dengan demikian, pemahaman terhadap *artificial intelligence* dan

transformasi dunia kerja tidak dapat dilepaskan dari dinamika sosial dan persepsi publik yang menyertainya. Landasan konseptual ini menjadi penting sebagai dasar untuk menganalisis bagaimana diskursus mengenai AI dalam dunia kerja direpresentasikan dalam bentuk sentimen kontekstual pada media sosial.

2.2.2 *Knowledge Work* dan *Knowledge Worker*

Knowledge work merupakan bentuk pekerjaan yang berpusat pada penggunaan pengetahuan, informasi, dan kemampuan kognitif sebagai sumber utama dalam menyelesaikan tugas. Berbeda dengan pekerjaan manual yang lebih banyak mengandalkan aktivitas fisik, *Knowledge work* menekankan kemampuan berpikir, menganalisis informasi, memecahkan masalah, berkomunikasi, serta mengambil keputusan. Dalam konteks ini, *Knowledge worker* dapat dipahami sebagai pekerja yang memiliki keahlian, pendidikan, atau pengalaman tertentu, kemudian menggunakan kemampuan tersebut untuk memperoleh, menciptakan, membagikan, maupun menerapkan pengetahuan dalam pekerjaannya [36].

Dalam literatur, istilah *Knowledge worker* digunakan untuk menggambarkan berbagai profesi yang memiliki karakteristik pekerjaan berbasis pengetahuan. Beberapa profesi yang umum digunakan sebagai contoh atau sampel *Knowledge worker* antara lain manajer, insinyur, profesional teknologi informasi, akuntan, pengacara, analis keuangan, dan profesor [14]. Selain itu, penelitian terdahulu juga menunjukkan bahwa profesi seperti *manager*, *engineer*, *researcher/scientist*, *consultant*, *professor*, *lawyer*, dan *physician* sering muncul dalam studi-studi yang membahas komunitas *Knowledge worker* [14]. Namun, penggunaan contoh profesi tersebut perlu dipahami sebagai gambaran umum, karena batasan istilah *Knowledge worker* dalam literatur masih dapat bervariasi tergantung konteks penelitian.

Pekerjaan berbasis pengetahuan menjadi relevan dalam konteks perkembangan *artificial intelligence*, khususnya *generative AI*, karena banyak tugas utama *Knowledge worker* berkaitan langsung dengan kemampuan yang

mulai dapat dibantu oleh AI. Aktivitas seperti penyusunan teks, pencarian dan analisis informasi, pembuatan rekomendasi, pengolahan data, komunikasi profesional, serta dukungan pengambilan keputusan merupakan contoh aktivitas kognitif dalam pekerjaan berbasis pengetahuan yang semakin banyak bersinggungan dengan teknologi AI. Penelitian sebelumnya menunjukkan bahwa *generative* AI dalam berbagai industri berbasis pengetahuan sering dipersepsikan sebagai alat yang dapat membantu pekerjaan rutin di bawah pengawasan manusia, meskipun tetap memunculkan isu seperti *deskilling*, dehumanisasi, disinformasi, dan perubahan struktur pekerjaan [37]. Oleh karena itu, *Knowledge work* menjadi konteks yang penting dalam penelitian ini untuk menganalisis bagaimana publik memandang AI sebagai ancaman, peluang, atau netral dalam dunia kerja berbasis pengetahuan.

2.2.3 Persepsi Ancaman dan Peluang terhadap *Artificial Intelligence*

Persepsi terhadap *artificial intelligence* dalam dunia kerja tidak selalu dapat dipahami secara sederhana sebagai sikap positif atau negatif. Dalam konteks teknologi yang berdampak luas seperti AI, individu dapat memandang teknologi tersebut melalui dua sisi utama, yaitu sebagai sumber risiko sekaligus sebagai sumber manfaat. Persepsi risiko muncul ketika AI dipahami sebagai teknologi yang dapat mengganggu stabilitas pekerjaan, menggantikan sebagian tugas manusia, mengurangi kebutuhan terhadap tenaga kerja tertentu, atau menimbulkan ketidakpastian terhadap masa depan profesi. Sebaliknya, persepsi manfaat muncul ketika AI dipandang sebagai teknologi yang dapat membantu pekerjaan, meningkatkan produktivitas, mempercepat proses pengambilan keputusan, serta membuka peluang baru dalam dunia kerja [7].

Dalam konteks dunia kerja, persepsi ancaman terhadap AI berkaitan erat dengan kekhawatiran bahwa kemampuan teknologi dapat mengambil alih sebagian fungsi yang sebelumnya dilakukan oleh manusia. Ancaman tersebut tidak terbatas pada aspek ekonomi, seperti risiko kehilangan pekerjaan, tetapi juga dapat menyentuh aspek psikologis dan profesional pekerja. Kehadiran AI dapat memunculkan *identity threat*, yaitu kondisi ketika pekerja merasa bahwa

identitas profesional, status, atau makna dari pekerjaan yang dilakukan mulai terancam oleh kemampuan teknologi baru [17]. Oleh karena itu, AI dapat dipersepsikan sebagai ancaman ketika dikaitkan dengan penggantian pekerjaan, penurunan peran manusia, ketidakpastian karier, atau melemahnya nilai keahlian profesional.

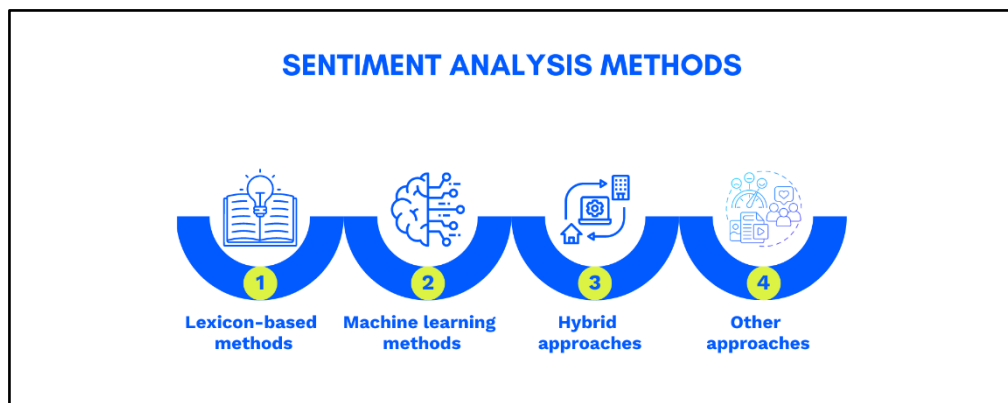
Di sisi lain, AI juga dapat dipersepsikan sebagai peluang ketika teknologi tersebut dipahami sebagai alat yang mendukung dan memperkuat kemampuan manusia. Dalam perspektif *augmentation*, AI diposisikan sebagai teknologi pendukung yang memperkuat kemampuan manusia, bukan semata-mata sebagai pengganti pekerjaan. Peran ini terlihat dari kemampuannya dalam membantu penyelesaian tugas, meningkatkan efisiensi, memperluas akses informasi, serta mendukung pengambilan keputusan [35]. Dengan demikian, persepsi peluang terhadap AI muncul ketika masyarakat melihat AI sebagai alat bantu kerja, sarana peningkatan produktivitas, pendukung inovasi, atau teknologi yang membuka kebutuhan terhadap keterampilan dan profesi baru.

Berdasarkan pemahaman tersebut, penelitian ini menggunakan kategori ancaman dan peluang sebagai fokus utama untuk mengklasifikasikan persepsi publik terhadap AI dalam konteks pekerjaan berbasis pengetahuan. Kategori ancaman digunakan untuk teks yang menunjukkan kekhawatiran terhadap dampak AI pada pekerjaan, seperti penggantian tenaga kerja, pemutusan hubungan kerja, penurunan nilai profesi, atau ketidakpastian masa depan pekerjaan. Kategori peluang digunakan untuk teks yang menunjukkan pandangan bahwa AI dapat membantu pekerjaan, meningkatkan efisiensi, membuka peluang baru, atau mendukung produktivitas. Sementara itu, kategori netral digunakan pada tahap pelabelan sebagai kategori penyaring untuk teks yang bersifat informatif, deskriptif, atau tidak menunjukkan kecenderungan yang jelas sebagai ancaman maupun peluang. Dengan demikian, penggunaan kategori ancaman dan peluang diharapkan dapat menangkap dua arah utama persepsi publik terhadap AI secara lebih kontekstual dibandingkan klasifikasi sentimen umum berbasis positif, negatif, dan netral.

2.2.4 *Sentiment Analysis* dan Kategorisasi Sentimen Kontekstual

Sentiment analysis merupakan teknik dalam *Natural Language Processing* (NLP), *text mining*, dan *data mining* yang digunakan untuk mengidentifikasi, mengekstraksi, serta mengklasifikasikan opini, sentimen, sikap, dan persepsi yang terkandung dalam data teks. Pendekatan ini banyak digunakan untuk menganalisis opini publik pada media digital karena mampu membantu proses pengolahan komentar, ulasan, dan percakapan dalam jumlah besar secara otomatis dan sistematis [38]. Dalam penelitian ini, *sentiment analysis* digunakan untuk memetakan bagaimana *artificial intelligence* dipersepsikan dalam konteks dunia kerja, khususnya apakah AI dipandang sebagai ancaman atau peluang.

Secara umum, analisis sentimen sering menggunakan kategori polaritas seperti positif, negatif, dan netral. Namun, dalam konteks tertentu, kategori polaritas umum belum selalu mampu menggambarkan makna opini secara spesifik [23]. Pada diskursus mengenai AI dan dunia kerja, suatu pernyataan tidak terbatas untuk menunjukkan emosi positif atau negatif, tetapi juga dapat mencerminkan persepsi terhadap AI sebagai risiko pekerjaan, alat bantu produktivitas, atau informasi yang bersifat deskriptif. Oleh karena itu, penelitian ini menggunakan kategorisasi sentimen kontekstual yang disesuaikan dengan permasalahan AI dalam dunia kerja. Kategori ancaman digunakan untuk merepresentasikan persepsi yang melihat AI sebagai risiko terhadap pekerjaan, sedangkan kategori peluang digunakan untuk merepresentasikan persepsi yang melihat AI sebagai teknologi pendukung aktivitas kerja. Sementara itu, kategori netral digunakan pada tahap pelabelan sebagai kategori penyaring terhadap teks yang tidak menunjukkan kecenderungan kuat sebagai ancaman maupun peluang. Penyesuaian kategori ini dilakukan agar hasil analisis mampu merepresentasikan dua arah utama persepsi publik terhadap AI dalam pekerjaan berbasis pengetahuan secara lebih relevan dibandingkan klasifikasi berbasis polaritas umum.



Gambar 2. 1 Klasifikasi Metode *Sentiment Analysis*

Sumber: [23]

Berdasarkan Gambar 2.1, metode analisis sentimen dapat dibagi menjadi empat pendekatan utama, yaitu:

1) *Lexicon-Based Methods*

Pendekatan ini menggunakan kamus atau leksikon yang telah diberi label sentimen tertentu, seperti positif, negatif, atau netral. Setiap kata dalam teks akan dicocokkan dengan kamus tersebut, kemudian skor sentimen dihitung berdasarkan akumulasi nilai kata yang ditemukan. Metode ini tidak memerlukan proses pelatihan model karena berbasis aturan, sehingga relatif sederhana dan mudah diimplementasikan. Namun, pendekatan ini memiliki keterbatasan dalam memahami konteks kalimat, makna ganda, ironi, serta bahasa informal yang sering muncul di media sosial. Akibatnya, akurasi metode leksikal cenderung lebih rendah pada teks yang kompleks dan dinamis seperti *tweet* [38].

2) *Machine Learning Methods*

Pendekatan *machine learning* menggunakan data berlabel untuk melatih model agar mampu mengenali pola sentimen dalam teks. Model akan belajar dari fitur yang diekstraksi, seperti frekuensi kata atau representasi numerik tertentu, untuk kemudian melakukan klasifikasi pada data baru. Metode ini mencakup algoritma klasik seperti *support vector machine* (SVM), *Logistic Regression*, dan *Naïve Bayes*, serta berkembang menuju *deep learning* dan

transformer. Keunggulan utama pendekatan ini adalah kemampuannya menyesuaikan diri terhadap pola data dan menghasilkan akurasi yang lebih tinggi dibanding metode berbasis leksikon, khususnya pada dataset besar [39].

3) *Hybrid Approaches*

Pendekatan *hybrid* mengombinasikan dua atau lebih metode untuk meningkatkan performa klasifikasi. Misalnya, penggunaan *embedding* berbasis leksikon yang dikombinasikan dengan algoritma *machine learning*, atau penggabungan model *transformer* dengan arsitektur *deep learning* seperti BiLSTM atau CNN. *Hybrid* bertujuan untuk memanfaatkan keunggulan masing-masing pendekatan, seperti kekuatan representasi kontekstual dari *transformer* dan kemampuan pemodelan sekuensial dari jaringan saraf. Pendekatan ini sering digunakan untuk meningkatkan stabilitas dan generalisasi model [38].

4) *Other Approaches*

Kategori *other Approaches* mencakup pendekatan lain yang tidak sepenuhnya masuk ke dalam kategori leksikon, *machine learning*, maupun *hybrid*. Pendekatan ini dapat meliputi *rule-Based systems*, *aspect-Based sentiment analysis*, pendekatan multimodal, *transfer learning*, serta model bahasa pralatih. Kategori ini berkembang seiring meningkatnya kebutuhan analisis sentimen pada domain yang lebih spesifik dan kompleks, terutama ketika penelitian bukan hanya ingin mengetahui polaritas umum, tetapi juga aspek, *target* opini, atau konteks tertentu dalam teks [39].

Dalam penelitian ini, tahap pemodelan klasifikasi dilakukan dengan pendekatan *supervised learning* karena model dilatih menggunakan data yang telah memiliki label akhir. Label akhir yang digunakan dalam pemodelan difokuskan pada dua kategori, yaitu ancaman dan peluang, sedangkan kategori netral digunakan pada tahap pelabelan sebagai kategori penyaring. Model yang digunakan meliputi TF-IDF dengan *support vector machine* (SVM) sebagai *baseline* klasik, serta IndoBERT dan IndoBERTweet sebagai model berbasis *transformer* untuk menangkap konteks bahasa secara lebih mendalam. Dengan

demikian, *sentiment analysis* dalam penelitian ini digunakan untuk memetakan persepsi publik terhadap AI secara lebih kontekstual dalam dunia kerja berbasis pengetahuan di Indonesia.

2.2.5 Media Sosial X/Twitter sebagai Sumber Data Opini Publik

Media sosial merupakan ruang digital yang memungkinkan pengguna untuk menyampaikan opini, pengalaman, respons, dan persepsi terhadap berbagai isu secara terbuka. Dalam konteks penelitian opini publik, media sosial dapat digunakan sebagai sumber data karena memuat percakapan masyarakat dalam jumlah besar, bersifat aktual, dan mencerminkan respons pengguna terhadap isu tertentu. Berbeda dengan survei tradisional yang membutuhkan proses pengumpulan data secara langsung, data media sosial dapat memberikan gambaran opini publik secara lebih cepat dan dalam skala yang lebih luas. Penelitian tentang *social media-Based public opinion analysis* menunjukkan bahwa media sosial dapat menjadi cara baru untuk merepresentasikan dan mengukur opini publik terhadap isu tertentu, meskipun tetap memiliki tantangan terkait kualitas data, pengumpulan data, dan aspek etis [40].

Salah satu platform media sosial yang banyak digunakan dalam penelitian opini publik adalah X/Twitter. Platform ini memiliki karakteristik percakapan yang singkat, terbuka, *real-time*, serta berbasis *user-generated content*. Karakter tersebut membuat X/Twitter sering digunakan untuk menangkap opini, sentimen, dan reaksi publik terhadap peristiwa, kebijakan, produk, maupun perkembangan teknologi. Twitter juga sering dipandang sebagai sumber data yang potensial untuk analisis berbasis *big data* karena jutaan pengguna dapat mengunggah pesan, berinteraksi, dan menyebarkan informasi secara cepat dalam berbagai konteks sosial [41].

Dalam penelitian analisis sentimen, X/Twitter menjadi sumber data yang relevan karena teks yang diunggah pengguna umumnya memuat opini dan sentimen terhadap isu tertentu. Twitter *sentiment analysis* digunakan untuk memproses sifat subjektif dari data Twitter, termasuk opini dan sentimen yang terkandung di dalamnya [42]. Selain itu, fitur seperti *hashtag*, *mention*, *retweet*,

dan pencarian berbasis kata kunci memudahkan peneliti dalam mengidentifikasi percakapan yang berkaitan dengan topik tertentu. Dalam konteks penelitian ini, X/Twitter digunakan untuk mengumpulkan opini publik mengenai *artificial intelligence* dan dunia kerja di Indonesia, sehingga data yang diperoleh dapat digunakan untuk memetakan persepsi publik ke dalam kategori ancaman, peluang, dan netral.

Dalam penelitian berbasis media sosial, penentuan rentang waktu pengambilan data menjadi aspek penting karena opini publik bersifat dinamis dan dapat berubah mengikuti perkembangan isu tertentu. Percakapan di X/Twitter mengenai teknologi baru seperti *generative AI* dapat mengalami perubahan seiring meningkatnya perhatian publik, munculnya penggunaan baru, serta berkembangnya pembahasan mengenai dampaknya terhadap masyarakat dan pekerjaan. Penelitian terdahulu menunjukkan bahwa setelah kemunculan ChatGPT, minat publik terhadap AI meningkat secara signifikan dan diskursus mengenai *generative AI* mulai melibatkan berbagai kelompok pekerjaan [20]. Selain itu, penelitian mengenai sikap publik terhadap ChatGPT di Twitter juga menunjukkan bahwa sentimen dan topik pembahasan dapat berubah dari waktu ke waktu [19]. Oleh karena itu, pemilihan periode data perlu disesuaikan dengan konteks perkembangan isu yang diteliti. Dalam penelitian ini, rentang waktu 2023–2025 digunakan untuk menangkap dinamika percakapan publik mengenai AI setelah meningkatnya diskursus *generative AI*, sekaligus mencakup periode ketika isu AI semakin relevan dalam pembahasan transformasi dunia kerja dan kebutuhan keterampilan masa depan [1].

Meskipun demikian, penggunaan X/Twitter sebagai sumber data juga memiliki keterbatasan. Data media sosial tidak selalu merepresentasikan seluruh populasi karena hanya mencerminkan pengguna yang aktif di platform tersebut. Selain itu, data yang diperoleh dapat mengandung *noise*, seperti duplikasi, spam, akun bot, bahasa informal, sarkasme, serta opini yang ambigu. Oleh karena itu, penelitian berbasis media sosial perlu melalui proses pembersihan data, seleksi kata kunci, *preprocessing*, dan pelabelan yang sesuai agar data yang dianalisis tetap relevan dengan tujuan penelitian. Dengan

mempertimbangkan kelebihan dan keterbatasan tersebut, X/Twitter tetap relevan digunakan sebagai sumber data untuk memahami dinamika opini publik terhadap AI dalam dunia kerja, terutama karena diskursus mengenai teknologi baru sering berkembang secara cepat di ruang digital.

2.2.6 Natural Language Processing (NLP)

Natural Language Processing (NLP) merupakan cabang dari *artificial intelligence* yang berfokus pada interaksi antara komputer dan bahasa manusia. NLP bertujuan untuk memungkinkan sistem komputer memahami, memproses, dan menganalisis teks atau bahasa alami secara otomatis. Dalam konteks analisis sentimen, NLP berperan penting dalam mengubah data teks yang tidak terstruktur menjadi representasi yang dapat diolah oleh model komputasi. Data yang berasal dari media sosial seperti X (Twitter) umumnya bersifat tidak terstruktur dan mengandung berbagai variasi bahasa, termasuk singkatan, slang, emoji, serta kesalahan penulisan [43]. Oleh karena itu, sebelum dilakukan proses klasifikasi sentimen, teks perlu melalui serangkaian tahapan *preprocessing* untuk meningkatkan kualitas data dan mengurangi *noise*. Tahapan ini bertujuan untuk menyederhanakan struktur teks tanpa menghilangkan makna utama yang terkandung di dalamnya. Secara umum, tahapan dalam NLP untuk analisis sentimen meliputi:

1) Case folding

Proses mengubah seluruh huruf dalam teks menjadi huruf kecil untuk menghindari perbedaan representasi antara huruf kapital dan huruf kecil.

2) Cleaning

Penghapusan elemen yang tidak relevan seperti URL, *mention*, *hashtag*, tanda baca tertentu, angka, atau karakter khusus yang tidak berkontribusi terhadap analisis sentimen.

3) Tokenization

Proses memecah teks menjadi unit-unit kecil yang disebut token, biasanya berupa kata atau sub-kata.

4) *Stopword Removal*

Penghapusan kata-kata umum yang sering muncul tetapi tidak memiliki makna signifikan dalam menentukan sentimen, seperti “dan”, “yang”, atau “di”.

5) *Stemming* atau *Lemmatization*

Proses mengubah kata ke bentuk dasarnya agar variasi morfologis dari kata yang sama dapat diperlakukan sebagai satu representasi. Setelah *preprocessing* dilakukan, teks perlu direpresentasikan dalam bentuk numerik agar dapat diproses oleh algoritma *machine learning*. Beberapa metode representasi teks yang umum digunakan antara lain:

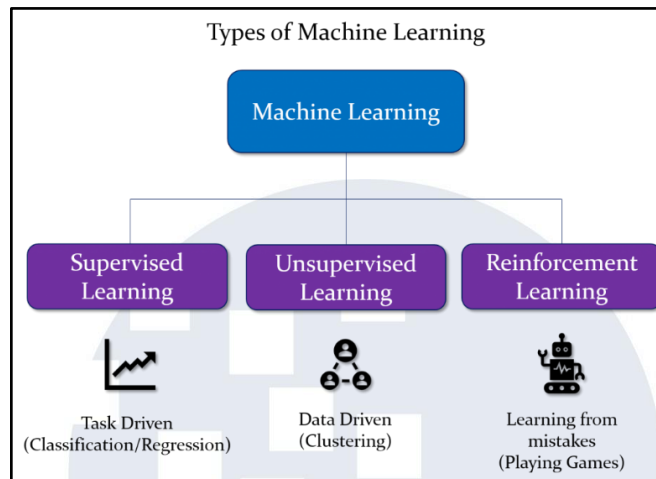
- *Bag of Words* (BoW), yang merepresentasikan teks berdasarkan frekuensi kemunculan kata tanpa mempertimbangkan urutan kata.
- TF-IDF (*Term Frequency–Inverse Document Frequency*), yang memberikan bobot lebih tinggi pada kata-kata yang dianggap penting dalam suatu dokumen dibandingkan kata yang sering muncul di seluruh dokumen.
- *Word Embedding*, yang memetakan kata ke dalam ruang vektor berdimensi tertentu untuk menangkap hubungan semantik antar kata.

Dalam penelitian ini, representasi teks menjadi komponen penting karena akan memengaruhi performa model klasifikasi. Pendekatan berbasis fitur seperti TF-IDF akan digunakan pada model klasik (SVM), sementara model *transformer* seperti IndoBERT dan IndoBERTweet menggunakan representasi kontekstual yang dihasilkan melalui mekanisme *embedding* internalnya. Dengan demikian, NLP berfungsi sebagai fondasi utama yang menjembatani data teks mentah dengan proses pembelajaran model klasifikasi.

2.2.7 *Machine Learning* dalam Klasifikasi Teks

Machine learning merupakan cabang dari *artificial intelligence* yang berfokus pada pengembangan sistem yang mampu belajar dari data dan meningkatkan performanya tanpa diprogram secara eksplisit untuk setiap tugas tertentu. Dalam konteks analisis sentimen, *machine learning* digunakan untuk membangun model klasifikasi yang dapat mengenali pola dalam teks dan

mengelompokkan data ke dalam kategori sentimen yang telah ditentukan [44]. Secara umum, *machine learning* dapat diklasifikasikan ke dalam tiga pendekatan utama sebagaimana ditunjukkan pada Gambar 2.2.



Gambar 2. 2 Klasifikasi Pendekatan *Machine Learning*
Sumber: [45]

Berdasarkan Gambar 2.2, pendekatan *machine learning* terdiri dari:

1. *Supervised learning*

Supervised learning merupakan pendekatan pembelajaran yang menggunakan data berlabel sebagai dasar pelatihan model. Dalam pendekatan ini, setiap data memiliki pasangan input dan output yang telah diketahui sebelumnya. Model akan belajar memetakan hubungan antara fitur input dan label *target*, sehingga dapat memprediksi label pada data baru. Pada klasifikasi teks, *supervised learning* digunakan untuk mengelompokkan dokumen ke dalam kategori tertentu, seperti positif, negatif, atau netral. Algoritma seperti *support vector machine* (SVM), *logistic regression*, dan model *transformer* termasuk dalam kategori ini. Pendekatan ini sangat umum digunakan dalam analisis sentimen karena tersedia data yang telah dilabeli.

2. *Unsupervised learning*

Unsupervised learning tidak menggunakan data berlabel dalam proses pelatihannya. Model berusaha menemukan pola atau struktur tersembunyi dalam data secara mandiri. Contoh penerapannya dalam analisis teks adalah

clustering atau *topic Modeling*. Meskipun bermanfaat untuk eksplorasi data, pendekatan ini tidak secara langsung digunakan untuk klasifikasi sentimen berbasis label seperti pada penelitian ini.

3. *Reinforcement Learning*

Reinforcement learning merupakan pendekatan di mana model belajar melalui interaksi dengan lingkungan dan menerima umpan balik dalam bentuk *reward* atau *penalty*. Pendekatan ini umumnya digunakan pada sistem pengambilan keputusan dinamis, seperti game atau sistem rekomendasi adaptif, dan jarang digunakan secara langsung untuk tugas klasifikasi sentimen tradisional.

Dalam penelitian ini, tahap pemodelan klasifikasi termasuk dalam pendekatan *supervised learning* karena model dilatih menggunakan data yang telah memiliki label akhir. Label akhir yang digunakan pada tahap pemodelan difokuskan pada dua kategori, yaitu ancaman dan peluang, sedangkan kategori netral digunakan sebagai penyaring pada tahap pelabelan. Proses pembentukan label dilakukan melalui kombinasi pelabelan manual dan *semi-supervised labeling*, kemudian hasil label akhir digunakan untuk melatih model klasifikasi. Model yang digunakan meliputi TF-IDF dengan *support vector machine* (SVM) sebagai *baseline* klasik, serta IndoBERT dan IndoBERTweet sebagai model berbasis *transformer*. Dengan memahami klasifikasi pendekatan *machine learning* ini, posisi metodologis penelitian menjadi lebih jelas sebelum memasuki pembahasan yang lebih teknis mengenai metode dan algoritma yang digunakan pada subbab berikutnya.

2.3 *Framework, Metode dan Algoritma yang digunakan*

2.3.1 CRISP-DM

CRISP-DM (*Cross Industry Standard Process for Data mining*) merupakan salah satu *framework* yang banyak digunakan dalam penelitian berbasis *data mining* dan *machine learning*. *Framework* ini menyediakan alur kerja sistematis yang terdiri dari enam tahapan utama, mulai dari pemahaman

permasalahan hingga implementasi model [46]. Penggunaan CRISP-DM dalam penelitian ini bertujuan untuk memastikan bahwa proses analisis sentimen dilakukan secara terstruktur, terarah, dan sesuai dengan tujuan penelitian. Secara umum, CRISP-DM terdiri dari enam tahapan, yaitu *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling*, *Evaluation*, dan *Deployment*. Keenam tahapan tersebut diadaptasi dalam penelitian ini sebagai berikut:

1. *Business Understanding*

Tahap ini berfokus pada pemahaman permasalahan penelitian dan tujuan yang ingin dicapai. Dalam konteks penelitian ini, permasalahan utama adalah bagaimana sentimen kontekstual terhadap *artificial intelligence* dalam pekerjaan berbasis pengetahuan muncul pada data X/Twitter berbahasa Indonesia yang diperoleh berdasarkan kata kunci penelitian pada periode 2023–2025. Tujuan yang ingin dicapai adalah membangun model klasifikasi sentimen yang mampu mengelompokkan opini publik ke dalam dua kategori, yaitu ancaman dan peluang, serta membandingkan performa antara pendekatan klasik dan *transformer*.

2. *Data Understanding*

Pada tahap ini dilakukan pengumpulan dan eksplorasi data awal. Data penelitian diperoleh melalui proses *scraping* pada platform X (Twitter) dengan rentang waktu 1 Januari 2023 hingga 31 Desember 2025 menggunakan kata kunci yang berkaitan dengan AI dan dunia kerja. Data yang terkumpul kemudian dianalisis secara deskriptif untuk memahami karakteristiknya, seperti jumlah data, distribusi teks, serta potensi ketidakseimbangan kelas. Pada tahap ini dilakukan pengumpulan dan eksplorasi data awal. Data penelitian diperoleh melalui proses *scraping* pada platform X (Twitter) dengan rentang waktu 1 Januari 2023 hingga 31 Desember 2025 menggunakan kata kunci yang berkaitan dengan AI dan dunia kerja. Data yang terkumpul kemudian dianalisis secara deskriptif untuk memahami karakteristiknya, seperti jumlah data, distribusi teks, serta potensi ketidakseimbangan kelas.

3. *Data Preparation*

Tahap *Data Preparation* mencakup proses pembersihan, pelabelan, dan transformasi data agar siap digunakan dalam tahap *Modeling*. Proses ini meliputi penghapusan duplikasi, pembersihan karakter khusus, *Case folding*, serta penyesuaian teks sesuai kebutuhan model. Pada penelitian ini, pelabelan dilakukan melalui kombinasi pelabelan manual dan *semi-supervised labeling*. Kategori netral digunakan sebagai penyaring, sedangkan kategori ancaman dan peluang digunakan sebagai label akhir pada tahap pemodelan. Data yang telah siap kemudian dibagi menjadi data latih, data validasi, dan data uji untuk mendukung proses pelatihan dan evaluasi model.

4. *Modeling*

Pada tahap *Modeling*, dilakukan pembangunan model klasifikasi menggunakan tiga pendekatan utama, yaitu TF-IDF dengan *support vector machine* (SVM), IndoBERT, dan IndoBERTweet. Model dilatih menggunakan data latih dan dilakukan proses penyesuaian parameter untuk memperoleh performa terbaik. Tahap ini bertujuan untuk membandingkan efektivitas representasi fitur tradisional dengan representasi kontekstual berbasis *transformer*.

5. *Evaluation*

Tahap evaluasi bertujuan untuk menilai performa model berdasarkan metrik klasifikasi seperti *accuracy*, *precision*, *recall*, dan *F1-score*. Evaluasi dilakukan untuk memastikan bahwa model memiliki akurasi tinggi, serta mampu mengklasifikasikan setiap kategori sentimen secara seimbang. Hasil evaluasi akan digunakan sebagai dasar dalam menentukan model terbaik.

6. *Deployment*

Tahap terakhir dalam CRISP-DM adalah *Deployment*, yaitu implementasi model ke dalam lingkungan yang dapat digunakan secara praktis. Dalam penelitian ini, model terbaik akan diimplementasikan dalam bentuk aplikasi berbasis web menggunakan Streamlit. Implementasi ini

memungkinkan pengguna untuk memasukkan teks dan memperoleh hasil klasifikasi sentimen secara langsung, sehingga hasil penelitian bersifat teoretis dan juga aplikatif.

Dengan menerapkan *framework* CRISP-DM, penelitian ini dilakukan secara sistematis mulai dari identifikasi masalah hingga implementasi model. Pendekatan ini memastikan bahwa setiap tahapan analisis memiliki keterkaitan yang jelas dan mendukung tujuan utama penelitian dalam memetakan sentimen publik terhadap AI dalam dunia kerja.

2.3.2 Uji Kesepakatan antar *Annotator* Menggunakan Cohen's *Kappa*

Cohen's *Kappa* merupakan metode statistik yang digunakan untuk mengukur tingkat kesepakatan antara dua *annotator* atau penilai dalam mengklasifikasikan data ke dalam kategori tertentu. Metode ini banyak digunakan dalam penelitian yang melibatkan proses pelabelan manual, terutama ketika label yang diberikan bersifat kategorikal atau nominal. Berbeda dengan persentase kesepakatan biasa, Cohen's *Kappa* menghitung jumlah label yang sama antara dua *annotator*, serta mempertimbangkan kemungkinan kesepakatan yang terjadi secara kebetulan [47].

Dalam penelitian analisis sentimen, Cohen's *Kappa* digunakan untuk mengetahui sejauh mana dua *annotator* memberikan label yang konsisten terhadap data teks. Hal ini penting karena proses pelabelan data media sosial memiliki potensi subjektivitas, terutama pada teks yang bersifat ambigu atau memiliki makna yang berdekatan antar kategori. Pada penelitian ini, Cohen's *Kappa* digunakan untuk mengukur kesepakatan antara *annotator* pertama dan *annotator* kedua dalam memberi label sentimen ke dalam tiga kategori, yaitu ancaman, peluang, dan netral. Secara umum, rumus Cohen's *Kappa* adalah sebagai berikut:

$$\kappa = \frac{P_o - P_e}{1 - P_e} \quad (2.1)$$

di mana:

- κ = nilai Cohen's *Kappa*,
- P_o = proporsi kesepakatan aktual antara *annotator*, dan
- P_e = proporsi kesepakatan yang diperkirakan terjadi secara kebetulan.

Nilai P_o menunjukkan tingkat kesepakatan yang benar-benar terjadi antara dua *annotator* berdasarkan hasil pelabelan. Sementara itu, nilai P_e menunjukkan peluang kesepakatan yang mungkin terjadi secara acak berdasarkan distribusi label dari masing-masing *annotator*. Dengan demikian, Cohen's *Kappa* memberikan ukuran reliabilitas yang lebih kuat dibandingkan persentase *agreement* biasa karena telah memperhitungkan faktor kebetulan dalam proses pelabelan.

Nilai Cohen's *Kappa* berada pada rentang -1 hingga 1. Nilai yang semakin mendekati 1 menunjukkan tingkat kesepakatan yang semakin tinggi, sedangkan nilai yang mendekati 0 menunjukkan bahwa kesepakatan antar *annotator* cenderung terjadi secara kebetulan. Nilai negatif menunjukkan bahwa tingkat kesepakatan antar *annotator* lebih rendah dibandingkan kesepakatan yang terjadi secara acak. Interpretasi nilai Cohen's *Kappa* dapat dilihat pada Tabel 2.2.

Tabel 2. 2 Interpretasi Nilai Cohen's *Kappa*

No	Nilai <i>Kappa</i>	Tingkat Kesepakatan
1	< 0,00	Poor <i>agreement</i>
2	0,00–0,20	Slight <i>agreement</i>
3	0,21–0,40	Fair <i>agreement</i>
4	0,41–0,60	Moderate <i>agreement</i>
5	0,61–0,80	Substantial <i>agreement</i>
6	0,81–1,00	Almost perfect <i>agreement</i>

Berdasarkan interpretasi tersebut, semakin tinggi nilai Cohen's *Kappa* maka semakin baik tingkat konsistensi pelabelan antar *annotator*. Dalam penelitian ini, Cohen's *Kappa* digunakan sebagai dasar untuk menilai reliabilitas proses pelabelan manual sebelum label final digunakan sebagai ground truth dalam proses pemodelan. Dengan adanya pengujian ini, hasil pelabelan diharapkan lebih konsisten dan dapat dipertanggungjawabkan secara ilmiah.

2.3.3 Semi-Supervised Labeling

Semi-supervised labeling merupakan pendekatan pelabelan yang memanfaatkan kombinasi data berlabel dan data belum berlabel untuk membantu proses pembentukan label pada dataset. Pendekatan ini digunakan ketika seluruh data sulit dilabeli secara manual karena jumlah data yang besar, keterbatasan waktu, serta kebutuhan menjaga konsistensi pelabelan. Dalam konteks klasifikasi teks, pendekatan ini dapat membantu memperluas data berlabel dengan memanfaatkan model awal yang dilatih dari sejumlah data berlabel sebagai dasar untuk memberikan label semu pada data lainnya [48].

Salah satu metode yang umum digunakan dalam *semi-supervised labeling* adalah *self-training*. Pada metode ini, model awal dilatih menggunakan data berlabel yang disebut *seed data*. Setelah model awal terbentuk, model tersebut digunakan untuk memprediksi label pada data yang belum berlabel. Label hasil prediksi tersebut disebut *pseudo-label*. Konsep *pseudo-labeling* pada dasarnya memilih kelas dengan probabilitas prediksi tertinggi sebagai label semu untuk data yang belum berlabel [49]. Dengan demikian, data yang awalnya belum memiliki label dapat dimanfaatkan untuk memperluas cakupan data berlabel dalam proses penelitian. Secara umum, proses pembentukan *pseudo-label* dapat dirumuskan pada Rumus 2.2.

$$\hat{y}_i = \arg \max_{y \in Y} P_{\theta}(y|x_i) \quad (2.2)$$

di mana:

- $(\hat{y}_i)$ = *pseudo-label* untuk data ke-(i)
- (x_i) = proporsi kesepakatan aktual antara *annotator*, dan
- (Y) = proporsi kesepakatan yang diperkirakan terjadi secara kebetulan.
- $(P_{\theta}(y|x_i))$ = proporsi kesepakatan yang diperkirakan terjadi secara kebetulan.
- (θ) = parameter model

Selain label prediksi, proses *pseudo-labeling* juga dapat mempertimbangkan nilai *confidence*. Nilai *confidence* diperoleh dari probabilitas prediksi tertinggi yang diberikan model terhadap suatu kelas. Rumus nilai *confidence* ditunjukkan pada Rumus 2.3.

$$c_i = \max_{y \in Y} P_{\theta}(y|x_i) \quad (2.3)$$

di mana:

- (c_i) = nilai *confidence* pada data ke-(i)
- $(P_{\theta}(y|x_i))$ = probabilitas prediksi model terhadap kelas (y) berdasarkan data (x_i)

Selanjutnya, *pseudo-label* dapat diterima apabila nilai *confidence* memenuhi ambang batas tertentu. Proses seleksi *pseudo-label* berdasarkan ambang *confidence* ditunjukkan pada Rumus 2.4.

$$\tilde{D} = \{(x_i, \hat{y}_i) \mid c_i \geq \tau\} \quad (2.4)$$

di mana:

- $(\tilde{D})$ = himpunan data dengan *pseudo-label* yang diterima
- (x_i) = data teks ke-(i)
- $(\hat{y}_i)$ = *pseudo-label* hasil prediksi model
- (c_i) = nilai *confidence*
- (τ) = nilai ambang batas *confidence*

Pemilihan nilai ambang *confidence* dapat berbeda pada setiap penelitian karena bergantung pada karakteristik data, model yang digunakan, serta tujuan seleksi *pseudo-label*. Nilai ambang yang terlalu rendah dapat meningkatkan risiko diterimanya label semu yang kurang tepat, sedangkan nilai ambang yang terlalu tinggi dapat membuat jumlah data yang diterima menjadi terlalu sedikit. Pada penelitian ini, nilai ambang *confidence* ditetapkan sebesar 0,75 [50]. Dalam penelitian ini, nilai tersebut digunakan sebagai batas seleksi pada skenario pelabelan, sedangkan data dengan nilai *confidence*

rendah tetap ditinjau kembali melalui proses *manual review* untuk menjaga kualitas label akhir.

Dalam penelitian ini, pendekatan *semi-supervised labeling* digunakan setelah proses pelabelan manual dan uji kesepakatan antar-annotator dilakukan. Data berlabel hasil pelabelan manual digunakan sebagai dasar untuk membentuk *seed* data. Selanjutnya, model awal dilatih menggunakan *seed* data tersebut dan digunakan untuk menghasilkan *pseudo-label* pada data lainnya. Data dengan nilai *confidence* yang memenuhi ambang batas diterima sebagai label semu, sedangkan data dengan nilai *confidence* rendah ditinjau kembali melalui proses *manual review*. Dengan pendekatan ini, proses pelabelan dapat dilakukan secara lebih efisien, tetapi tetap mempertahankan kontrol kualitas melalui pedoman pelabelan dan peninjauan manual terhadap data dengan tingkat keyakinan rendah

2.3.4 Representasi Teks Menggunakan TF-IDF

Dalam proses klasifikasi teks berbasis *machine learning* klasik, data teks perlu diubah menjadi representasi numerik agar dapat diproses oleh algoritma komputasi. Salah satu metode representasi teks yang umum digunakan adalah TF-IDF (*Term Frequency–Inverse Document Frequency*). Metode ini digunakan untuk mengukur tingkat kepentingan suatu kata dalam sebuah dokumen relatif terhadap seluruh dokumen dalam kumpulan data. TF-IDF merupakan pengembangan dari metode *Bag of Words* yang menghitung frekuensi kemunculan kata, serta mempertimbangkan seberapa umum kata tersebut muncul dalam keseluruhan korpus [51]. Dengan demikian, kata yang sering muncul pada satu dokumen tetapi jarang muncul pada dokumen lain akan memiliki bobot yang lebih tinggi. Secara matematis, TF-IDF terdiri dari dua komponen utama, yaitu *term frequency* (TF) dan *inverse document frequency* (IDF).

1. *Term Frequency* (TF)

Term frequency mengukur seberapa sering suatu kata muncul dalam sebuah dokumen. Secara umum, TF dirumuskan sebagai:

$$TF(t, d) = \frac{f_{t,d}}{\sum_{t' \in d} f_{t',d}} \quad (2.5)$$

di mana:

- $f_{t,d}$ adalah jumlah kemunculan term t dalam dokumen d
- Penyebut adalah total seluruh kata dalam dokumen d

TF merepresentasikan kepentingan kata dalam satu dokumen tertentu.

2. Inverse Document Frequency (IDF)

Inverse document frequency mengukur seberapa penting suatu kata dalam keseluruhan korpus dokumen. Kata yang sering muncul di banyak dokumen akan memiliki bobot lebih rendah dibandingkan kata yang jarang muncul.

Rumus IDF adalah:

$$IDF(t) = \log \left(\frac{N}{df_t} \right) \quad (2.6)$$

di mana:

- N adalah jumlah total dokumen
- df_t adalah jumlah dokumen yang mengandung term t

IDF bertujuan untuk mengurangi bobot kata-kata umum yang kurang informatif.

3. Perhitungan TF-IDF

Nilai akhir TF-IDF diperoleh dari hasil perkalian TF dan IDF:

Rumus TF-IDF adalah:

$$TF - IDF(t, d) = \frac{f_{t,d}}{\sum_{t' \in d} f_{t',d}} \times \log \left(\frac{N}{df_t} \right) \quad (2.7)$$

di mana:

- $TF\text{-}IDF(t,d)$ = bobot term t dalam dokumen d
- $TF(t,d)$ = nilai frekuensi term dalam dokumen
- $IDF(t)$ = nilai kelangkaan term dalam korpus

Nilai ini merepresentasikan tingkat kepentingan suatu kata dalam dokumen tertentu relatif terhadap seluruh dokumen dalam dataset. Dalam penelitian ini, TF-IDF digunakan sebagai metode representasi teks untuk model *support vector machine* (SVM). Pendekatan ini dipilih karena TF-IDF mampu menghasilkan fitur numerik yang stabil dan efektif untuk klasifikasi teks berbasis *machine learning* klasik. Selain itu, kombinasi TF-IDF dan SVM telah banyak digunakan dalam penelitian analisis sentimen dan terbukti memiliki performa yang kompetitif, terutama pada dataset berbahasa Indonesia dengan ukuran menengah.

Meskipun demikian, TF-IDF memiliki keterbatasan karena tidak mempertimbangkan urutan kata maupun konteks semantik secara mendalam. Setiap kata diperlakukan sebagai entitas independen tanpa memahami hubungan antar kata dalam kalimat. Oleh karena itu, dalam penelitian ini TF-IDF dengan SVM digunakan sebagai *baseline* klasik yang kemudian dibandingkan dengan model berbasis *transformer* yang mampu menghasilkan representasi kontekstual.

2.3.5 Algoritma *Support Vector Machine* (SVM)

Support vector machine (SVM) merupakan salah satu algoritma *supervised learning* yang banyak digunakan untuk tugas klasifikasi dan regresi. Dalam konteks klasifikasi teks, SVM dikenal memiliki performa yang baik terutama pada data berdimensi tinggi seperti hasil representasi TF-IDF. Algoritma ini bekerja dengan mencari *hyperplane* optimal yang mampu memisahkan data dari dua kelas atau lebih dengan margin terbesar. Secara geometris, SVM berupaya menemukan sebuah *hyperplane* yang memisahkan data dari kelas yang berbeda dengan jarak maksimum terhadap titik data terdekat dari masing-masing kelas. Titik-titik data yang berada paling dekat dengan *hyperplane* disebut sebagai *Support vectors*, karena titik-titik inilah

yang menentukan posisi dan orientasi *hyperplane*. Misalkan terdapat dataset dengan pasangan data (x_i, y_i) , di mana x_i adalah vektor fitur dan $y_i \in \{-1, +1\}$ adalah label kelas. *Hyperplane* dalam ruang fitur dapat dinyatakan sebagai:

$$w^T x + b = 0 \quad (2.8)$$

di mana:

- w : adalah vektor bobot (*weight vector*),
- x : adalah vektor fitur,
- $w^T x$: adalah hasil perkalian antara vektor bobot dan vektor fitur,
- b : adalah bias.

Tujuan utama SVM adalah mencari nilai w dan b yang memaksimalkan margin, yaitu jarak antara *hyperplane* dengan *Support vectors* dari masing-masing kelas. Secara matematis, permasalahan optimasi SVM dapat dirumuskan sebagai:

$$\min_{w,b} \frac{1}{2} \|w\|^2 \quad (2.9)$$

dengan kendala:

$$y_i(w^T x_i + b) \geq 1, \quad \forall i \quad (2.10)$$

Fungsi objektif tersebut bertujuan untuk meminimalkan norma vektor bobot, yang secara tidak langsung memaksimalkan margin pemisah antar kelas. Dalam praktiknya, SVM juga dapat menggunakan fungsi *kernel* untuk menangani data yang tidak dapat dipisahkan secara linear. Fungsi *kernel* memungkinkan data diproyeksikan ke ruang berdimensi lebih tinggi sehingga pemisahan menjadi memungkinkan. Beberapa jenis *kernel* yang umum digunakan antara lain linear, *polynomial*, dan radial basis function (RBF). Pada klasifikasi teks berbasis TF-IDF, *kernel* linear sering digunakan karena data teks umumnya sudah berada pada ruang berdimensi tinggi [52].

Dalam penelitian ini, SVM digunakan sebagai model *baseline* klasik dengan representasi fitur TF-IDF. Pemilihan SVM didasarkan pada beberapa pertimbangan, yaitu kemampuannya menangani data berdimensi tinggi, ketahanan terhadap *overfitting* pada dataset dengan fitur besar, serta performanya yang stabil dalam berbagai penelitian analisis sentimen berbahasa Indonesia. Model SVM dalam penelitian ini akan dibandingkan dengan model berbasis *transformer* untuk mengevaluasi perbedaan performa antara pendekatan berbasis fitur dan pendekatan berbasis representasi kontekstual.

2.3.6 Transformer, IndoBERT, dan IndoBERTweet

Perkembangan model klasifikasi teks mengalami lompatan signifikan dengan hadirnya arsitektur *transformer*. Berbeda dengan pendekatan tradisional yang bergantung pada representasi berbasis frekuensi kata seperti TF-IDF, model *transformer* mampu memahami hubungan kontekstual antar kata dalam suatu kalimat melalui mekanisme *attention*. Kemampuan ini memungkinkan model untuk menangkap makna yang lebih kompleks, termasuk dependensi jarak jauh antar kata.

1. Arsitektur Transformer

Transformer diperkenalkan sebagai arsitektur berbasis *self-attention* yang tidak menggunakan mekanisme sekuensial seperti RNN atau LSTM. Komponen utama dalam *transformer* adalah mekanisme *attention* yang memungkinkan model memberi bobot berbeda pada setiap kata dalam kalimat saat memproses representasi suatu kata.

Mekanisme dasar *attention* dapat dirumuskan sebagai:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (2.11)$$

di mana:

- $Q = \text{Query}$
- $K = \text{Key}$
- $V = \text{Value}$

- d_k = dimensi vektor key

Rumus tersebut menunjukkan bahwa skor perhatian dihitung berdasarkan kesamaan antara *query* dan *key*, kemudian dinormalisasi menggunakan fungsi *softmax*. *Self-attention* memungkinkan setiap kata dalam kalimat untuk “melihat” seluruh kata lain dalam kalimat tersebut, sehingga konteks dapat dipahami secara dua arah (*bidirectional*).

2. BERT (*Bidirectional Encoder Representations from Transformers*)

BERT merupakan model berbasis *transformer* yang dilatih menggunakan pendekatan *bidirectional training*, sehingga dapat memahami konteks dari arah kiri dan kanan secara bersamaan. Berbeda dengan model sekuensial tradisional, BERT tidak memproses kata secara satu arah, melainkan mempertimbangkan keseluruhan konteks dalam satu waktu. Dalam penerapan klasifikasi sentimen, BERT digunakan melalui proses *fine-tuning*. Pada tahap ini, representasi akhir dari token khusus (CLS) digunakan sebagai input untuk layer klasifikasi yang umumnya menggunakan fungsi *softmax*:

$$\hat{y} = \text{softmax}(Wx + b) \quad (2.12)$$

di mana:

- x adalah representasi vektor dari [CLS],
- W adalah matriks bobot,
- b adalah bias.

Hasil *softmax* berupa probabilitas untuk masing-masing kelas sentimen.

3. IndoBERT

IndoBERT merupakan model BERT yang telah dilatih menggunakan korpus besar berbahasa Indonesia. Model ini dirancang untuk memahami karakteristik morfologi dan struktur sintaksis bahasa Indonesia secara lebih akurat dibandingkan model multilingual umum [29]. Dalam penelitian ini, IndoBERT digunakan untuk melakukan *fine-tuning* pada dataset sentimen yang berkaitan dengan diskursus *artificial intelligence* dan dunia kerja. Representasi kontekstual yang dihasilkan IndoBERT diharapkan mampu

menangkap makna implisit dalam teks yang tidak dapat direpresentasikan melalui pendekatan berbasis frekuensi kata.

4. IndoBERTweet

IndoBERTweet merupakan model *transformer* yang secara khusus dilatih menggunakan data media sosial Indonesia, khususnya Twitter. Model ini lebih adaptif terhadap bahasa informal, singkatan, slang, serta struktur kalimat tidak baku yang umum ditemukan dalam percakapan daring [30]. Karena penelitian ini menggunakan data dari platform X (Twitter), IndoBERTweet menjadi model yang relevan untuk dibandingkan dengan IndoBERT dan pendekatan klasik TF-IDF dengan SVM. Perbandingan ini bertujuan untuk mengevaluasi apakah model yang dilatih pada domain media sosial memiliki performa yang lebih baik dalam memahami opini publik digital.

2.3.7 Evaluasi Model

Evaluasi model dilakukan untuk mengukur sejauh mana model klasifikasi mampu memprediksi kategori sentimen secara akurat dan konsisten. Dalam penelitian ini, evaluasi digunakan untuk membandingkan performa tiga pendekatan model, yaitu TF-IDF dengan SVM, IndoBERT, dan IndoBERTweet, dalam mengklasifikasikan sentimen publik ke dalam dua kelas yaitu ancaman dan peluang. Karena penelitian ini menggunakan klasifikasi dua-kelas, evaluasi tidak terbatas untuk melihat ketepatan secara keseluruhan, tetapi juga mempertimbangkan performa model pada masing-masing kelas sentimen.

1. *Confusion Matrix*

Confusion matrix merupakan representasi tabel yang menunjukkan perbandingan antara label aktual dan label hasil prediksi model [53]. Pada klasifikasi dua-kelas, *confusion matrix* akan berbentuk matriks berukuran $n \times n$, di mana n adalah jumlah kelas.

Melalui *confusion matrix*, dapat diketahui:

- a. Jumlah data ancaman yang benar diklasifikasikan sebagai ancaman dan jumlah data peluang yang benar diklasifikasikan sebagai peluang.
- b. Jumlah data ancaman yang salah diklasifikasikan sebagai peluang serta jumlah data peluang yang salah diklasifikasikan sebagai ancaman.
- c. Distribusi kesalahan prediksi antar kelas yang menjadi dasar perhitungan metrik evaluasi lainnya, seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

2. *Accuracy*

Accuracy mengukur proporsi keseluruhan prediksi yang benar dibandingkan dengan seluruh data yang diuji.

Rumus *accuracy*:

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}} \quad (2.13)$$

di mana:

- TP (*True Positive*) = jumlah data pada kelas positif yang berhasil diprediksi dengan benar sebagai kelas positif.
- TN (*True Negative*) = jumlah data pada kelas negatif yang berhasil diprediksi dengan benar sebagai kelas negatif.
- FP (*False positive*) = jumlah data pada kelas negatif yang salah diprediksi sebagai kelas positif.
- FN (*False Negative*) = jumlah data pada kelas positif yang salah diprediksi sebagai kelas negatif.

Dalam konteks klasifikasi dua-kelas, *accuracy* dihitung berdasarkan jumlah seluruh prediksi benar dibagi total data. Metrik ini memberikan gambaran umum mengenai performa model. Namun, *accuracy* memiliki keterbatasan ketika distribusi data tidak seimbang. Misalnya, jika kelas ancaman jauh lebih banyak dibandingkan kelas peluang, model dapat memperoleh *accuracy* tinggi hanya dengan lebih sering memprediksi ancaman,

meskipun performa pada kelas lain rendah. Oleh karena itu, *accuracy* perlu dikombinasikan dengan metrik lain untuk memberikan evaluasi yang lebih adil.

3. *Precision*

Precision mengukur tingkat ketepatan model dalam memprediksi suatu kelas tertentu. *Precision* menunjukkan berapa banyak data yang diprediksi sebagai suatu kelas yang benar-benar termasuk dalam kelas tersebut.

Rumus *precision*:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2.14)$$

Dalam penelitian ini, *precision* penting untuk melihat seberapa tepat model ketika memprediksi sentimen ancaman dan peluang. Misalnya, *precision* tinggi pada kelas ancaman berarti sebagian besar *tweet* yang diprediksi sebagai ancaman memang benar mengandung narasi ancaman. *Precision* menjadi krusial ketika kesalahan prediksi positif (*false positive*) dapat menyebabkan interpretasi yang keliru terhadap arah opini publik.

4. *Recall*

Recall mengukur kemampuan model dalam menangkap seluruh data yang benar-benar termasuk dalam suatu kelas.

Rumus *recall*:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2.15)$$

Recall tinggi menunjukkan bahwa model mampu menemukan sebagian besar data aktual dari suatu kelas. Dalam konteks penelitian ini, *recall* pada kelas ancaman penting untuk memastikan bahwa opini publik yang benar-benar mengandung kekhawatiran terhadap AI tidak terlewatkan oleh model.

Jika *recall* rendah, berarti banyak data yang seharusnya diklasifikasikan sebagai ancaman atau peluang justru salah diprediksi sebagai kelas lain.

5. F1-score

F1-score merupakan rata-rata harmonis antara *precision* dan *recall*, yang bertujuan untuk menyeimbangkan kedua metrik tersebut.

Rumus F1-score:

$$\text{F1-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.16)$$

F1-score digunakan ketika diperlukan keseimbangan antara ketepatan dan kelengkapan prediksi. Dalam kasus klasifikasi sentimen dua-kelas seperti penelitian ini, F1-score memberikan gambaran yang lebih representatif dibandingkan *accuracy* saja, terutama ketika distribusi data antar kelas tidak seimbang.

Karena penelitian ini menggunakan dua kelas sentimen, yaitu ancaman, dan peluang, maka nilai *precision*, *recall*, dan F1-score dihitung untuk masing-masing kelas dan kemudian dirata-ratakan menggunakan pendekatan *Macro average* dan *weighted average*. *Macro average* menghitung rata-rata sederhana tanpa mempertimbangkan jumlah data pada setiap kelas, sehingga memberikan gambaran performa model yang setara untuk seluruh kelas. Sementara itu, *weighted average* memperhitungkan proporsi jumlah data pada masing-masing kelas, sehingga lebih merepresentasikan kondisi distribusi data yang sebenarnya. Penggunaan kedua pendekatan ini memungkinkan evaluasi yang lebih komprehensif terhadap performa model, terutama dalam situasi di mana distribusi data antar kelas tidak seimbang [54].

Dalam konteks penelitian ini, metrik evaluasi bukan hanya digunakan untuk mengidentifikasi model dengan performa terbaik secara keseluruhan, tetapi juga untuk menganalisis keseimbangan performa antar kelas sentimen serta menentukan model yang paling sesuai untuk diimplementasikan dalam

aplikasi berbasis Streamlit. Dengan pendekatan evaluasi yang menyeluruh, hasil klasifikasi tidak semata-mata dinilai berdasarkan tingkat akurasi umum, melainkan juga berdasarkan kemampuan model dalam mengklasifikasikan dinamika opini publik secara lebih adil dan kontekstual terhadap isu *artificial intelligence* dalam dunia kerja.

2.4 Tools/Software yang digunakan

Dalam penelitian ini, proses pengolahan data, pembangunan model, hingga implementasi sistem dilakukan dengan memanfaatkan beberapa *tools* utama yang saling terintegrasi dalam satu ekosistem pengembangan. Penggunaan *tools* yang tepat menjadi faktor penting untuk memastikan bahwa setiap tahapan penelitian, mulai dari pengumpulan data, *preprocessing* teks, ekstraksi fitur, pelatihan model, evaluasi performa, hingga tahap *Deployment*, dapat berjalan secara sistematis dan efisien. *Tools* yang digunakan dalam penelitian ini meliputi bahasa pemrograman Python sebagai fondasi utama pengembangan, lingkungan kerja Jupyter Notebook untuk proses eksplorasi dan eksperimen model secara interaktif, serta *framework* Streamlit yang digunakan untuk mengimplementasikan model terbaik ke dalam bentuk aplikasi berbasis web. Integrasi ketiga *tools* tersebut memungkinkan proses penelitian berfokus pada aspek analitis, serta menghasilkan sistem yang dapat digunakan secara praktis.

2.4.1 Python

Python merupakan bahasa pemrograman tingkat tinggi yang banyak digunakan dalam bidang *data science*, *machine learning*, dan pengolahan bahasa alami (*Natural Language Processing*). Dalam penelitian ini, Python digunakan sebagai bahasa pemrograman utama untuk mengintegrasikan seluruh tahapan analisis, mulai dari proses pengumpulan data, *preprocessing* teks, ekstraksi fitur, pembangunan model klasifikasi, hingga evaluasi performa. Bahasa pemrograman Python memiliki karakteristik sintaks yang sederhana serta dukungan pustaka yang luas sehingga banyak digunakan dalam penelitian berbasis analisis data [55]. Identitas visual Python ditunjukkan pada Gambar 2.3.



Gambar 2. 3 Bahasa Pemrograman Python
Sumber: [56]

Dalam penelitian ini, Python dimanfaatkan melalui berbagai pustaka seperti Pandas untuk manipulasi dan pengolahan data, NumPy untuk komputasi numerik, Scikit-learn untuk implementasi TF-IDF dan *support vector machine*, serta *transformers* dari Hugging Face untuk *fine-tuning* model IndoBERT dan IndoBERTweet. Dengan dukungan ekosistem tersebut, Python memungkinkan seluruh proses analisis dilakukan secara terintegrasi dan efisien dalam satu lingkungan pengembangan.

2.4.2 Jupyter Notebook

Jupyter Notebook merupakan lingkungan pengembangan interaktif yang banyak digunakan dalam penelitian berbasis *data science* dan *machine learning*. Platform ini memungkinkan peneliti untuk menuliskan dan menjalankan kode secara bertahap, sekaligus menampilkan output dalam bentuk tabel, grafik, maupun visualisasi lainnya dalam satu dokumen yang terstruktur [57]. Dalam penelitian ini, Jupyter Notebook digunakan sebagai media utama untuk melakukan eksplorasi data, *preprocessing* teks, pelatihan model, serta evaluasi performa klasifikasi. Lingkungan ini mendukung proses eksperimen yang iteratif, sehingga peneliti dapat melakukan pengujian berbagai parameter dan konfigurasi model secara fleksibel. Tampilan identitas visual Jupyter Notebook ditunjukkan pada Gambar 2.4.



Gambar 2. 4 Jupyter Notebook
Sumber: [57]

Selain itu, Jupyter Notebook juga mendukung integrasi dengan berbagai pustaka Python yang digunakan dalam penelitian ini, seperti Scikit-learn dan *transformers*. Melalui dukungan komputasi tambahan seperti GPU pada platform berbasis *cloud*, proses *fine-tuning* model *transformer* dapat dilakukan dengan lebih efisien. Dengan demikian, Jupyter Notebook berperan sebagai lingkungan kerja utama yang memfasilitasi seluruh proses eksperimen dan pengembangan model secara sistematis.

2.4.3 TweetHarvest

TweetHarvest merupakan alat berbasis *command-line* yang digunakan untuk melakukan pengumpulan data dari hasil pencarian X/Twitter. Alat ini bekerja dengan memanfaatkan *Playwright* untuk membuka halaman pencarian X/Twitter dan mengambil data *tweet* berdasarkan kata kunci serta rentang waktu tertentu. Hasil pengambilan data kemudian disimpan dalam format CSV sehingga dapat digunakan untuk proses analisis data pada tahap berikutnya [58].

Dalam penelitian ini, TweetHarvest digunakan pada tahap pengumpulan data karena mampu mendukung proses *scraping* berdasarkan kata kunci yang telah ditentukan. Kata kunci tersebut disusun sesuai dengan ruang lingkup penelitian, yaitu percakapan mengenai *artificial intelligence* dalam konteks dunia kerja dan pekerjaan berbasis pengetahuan. Data yang diperoleh melalui

TweetHarvest kemudian digabungkan, dibersihkan, disaring, dan diproses lebih lanjut sebelum digunakan dalam tahap pelabelan dan pemodelan sentimen.

2.4.4 Streamlit

Streamlit merupakan *framework* berbasis Python yang digunakan untuk membangun aplikasi web interaktif secara sederhana dan cepat. *Framework* ini dirancang khusus untuk kebutuhan *data science* dan *machine learning*, sehingga memungkinkan model yang telah dilatih untuk diimplementasikan dalam bentuk aplikasi yang dapat digunakan secara langsung oleh pengguna [59]. Dalam penelitian ini, Streamlit digunakan pada tahap *Deployment* untuk mengimplementasikan model klasifikasi sentimen terbaik hasil evaluasi. Melalui aplikasi yang dibangun, pengguna dapat memasukkan teks terkait diskursus *artificial intelligence* dan dunia kerja, kemudian sistem akan memproses teks tersebut dan menampilkan hasil klasifikasi ke dalam kategori ancaman, peluang, atau netral. Identitas visual Streamlit ditunjukkan pada Gambar 2.5.



Gambar 2. 5 Streamlit
Sumber: [60]

Penggunaan Streamlit memungkinkan hasil penelitian tidak terbatas pada tahap analisis dan evaluasi model, tetapi juga dapat direpresentasikan dalam bentuk sistem yang aplikatif dan interaktif. Dengan demikian, implementasi ini menunjukkan bahwa model yang dikembangkan dalam

penelitian mampu digunakan secara praktis untuk mendukung analisis sentimen terhadap opini publik.



UMN

UNIVERSITAS
MULTIMEDIA
NUSANTARA