

BAB II

TINJAUAN PUSTAKA

2.1. Landasan Teori

2.1.1. Bias Media

2.1.1.1. Definisi dan Karakteristik

Bias media memiliki tiga karakteristik utama menurut literatur: (1) bias cenderung condong, dengan tujuan mempengaruhi opini publik ke arah tertentu; (2) bias sering terjadi secara konsisten pada setiap media atau jurnalis; dan (3) jurnalis mungkin menambahkan bias untuk membuat cerita lebih berkesan. Bias ini juga selalu disampaikan dalam format jurnalistik. Bias media dapat didefinisikan sebagai situasi di mana jurnalis memiringkan informasi dalam berita untuk mempengaruhi opini penerima, yang terjadi secara sering dan konsisten pada media atau jurnalis tertentu. Selain itu, berbagai definisi bias media juga menawarkan klasifikasi berdasarkan dua faktor utama: (1) maksud penulis, dan (2) konteks terjadinya bias (Rodrigo-Ginés dkk., 2024).

2.1.1.2. Jenis Bias Media Menurut Maksud Penulis

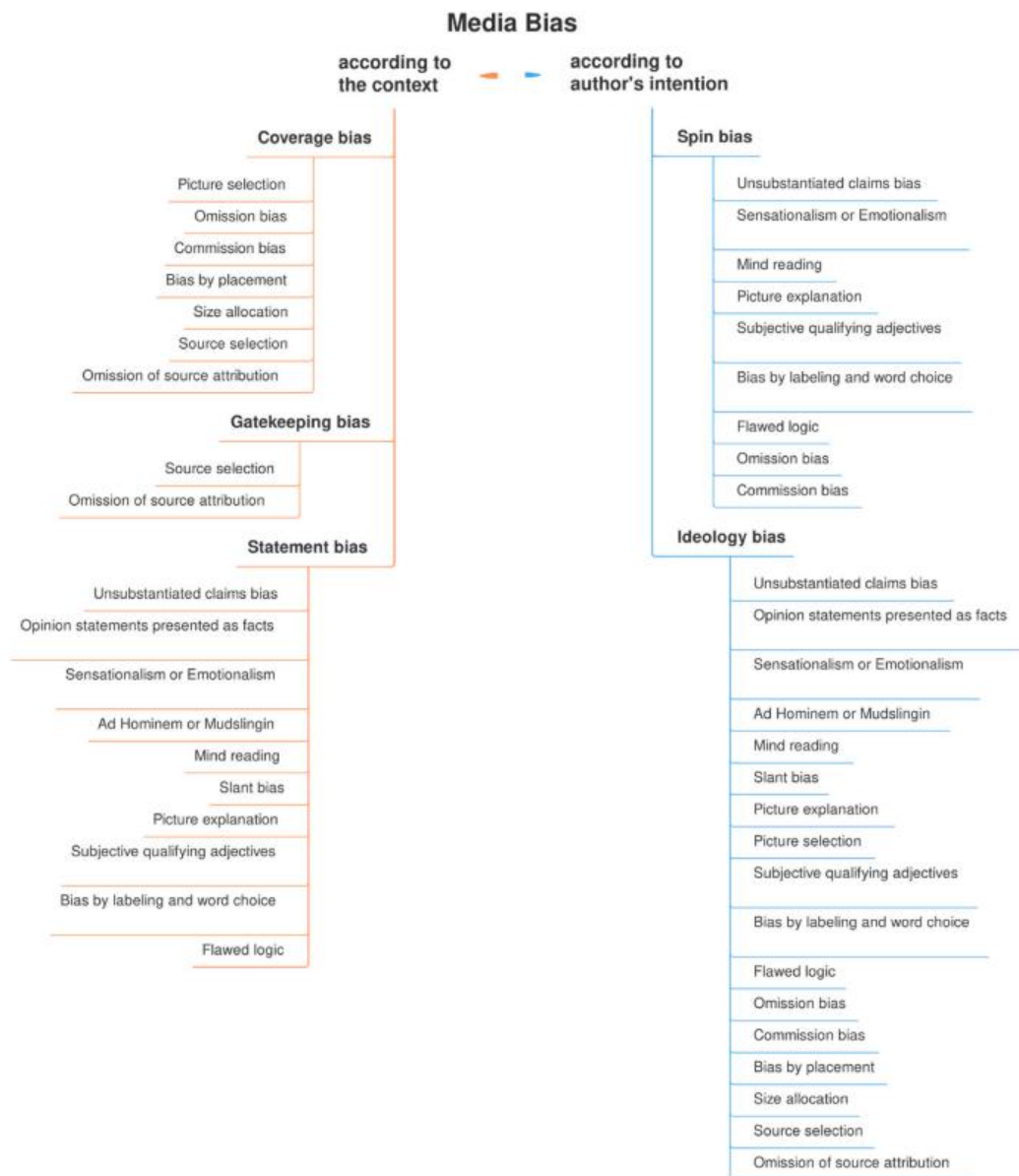
Bias media terbagi dalam dua kategori berdasarkan maksud penulis: bias putaran (retorika) dan bias ideologi. Bias putaran terjadi ketika jurnalis berusaha menciptakan cerita yang berkesan, sementara bias ideologi, atau *framing*, terjadi ketika informasi disajikan secara parsial. Bias ideologi sering ditemukan dalam isu politik, namun tidak terbatas pada spektrum politik tradisional (kiri vs. kanan) (Rodrigo-Ginés dkk., 2024).

2.1.1.3. Jenis Bias Media Menurut Konteksnya

Bias media dapat dibagi menjadi tiga kategori berdasarkan konteks: bias liputan, bias penjaga gerbang, dan bias pernyataan (Saez-Trumper dkk., 2013). Bias liputan muncul ketika media lebih memilih meliput berita tertentu, terutama yang negatif, sehingga menciptakan persepsi yang salah tentang realitas (Stafford, 2014). Bias penjaga gerbang terjadi ketika media memilih atau menolak berita untuk diliput, menghasilkan penggambaran peristiwa yang tidak seimbang (Saez-Trumper dkk., 2013). Bias pernyataan muncul ketika media menyajikan informasi dengan cara yang memengaruhi persepsi melalui pilihan kata atau bahasa yang sarat muatan (Hamborg dkk., 2019).

2.1.1.4. Bentuk Bias Media

Bentuk bias media diklasifikasikan berdasarkan dua jenis bias menghasilkan 17 bentuk bias media, yaitu bias klaim yang tidak berdasar, pernyataan opini disajikan sebagai bias fakta, bias sensasionalisme/emosionalisme, bias pengacau, bias membaca pikiran, bias miring, bias kelalaian, bias komisi, bias berdasarkan pelabelan dan pilihan kata, bias logika yang salah, bias berdasarkan penempatan, bias kata sifat yang memenuhi syarat subjektif, bias alokasi ukuran, bias pemilihan sumber, penghilangan bias atribut sumber, bias pemilihan gambar, bias penjelasan gambar (Rodrigo-Ginés dkk., 2024).



Gambar 2.1 Bentuk Bias Media (Rodrigo-Ginés dkk., 2024)

2.1.2. *Natural Language Processing (NLP)*

NLP adalah bidang Kecerdasan Buatan dan Linguistik, yang ditujukan untuk membuat komputer memahami pernyataan atau kata-kata yang ditulis dalam bahasa manusia. *NLP* dapat diklasifikasikan menjadi dua bagian, yaitu *Natural Language Understanding* (Linguistic) (NLU) dan *Natural Language Generation*

(NLG). *Natural Language Understanding* (NLU) memungkinkan mesin untuk memahami bahasa alami dan menganalisisnya dengan mengekstraksi konsep, entitas, emosi, kata kunci, dan lain-lain. Sementara itu, *Natural Language Generation* (NLG) adalah proses menghasilkan frasa, kalimat, dan paragraf yang bermakna dari representasi internal. Bidang *NLP* terkait dengan berbagai teori dan teknik yang berhubungan dengan masalah bahasa alami dalam berkomunikasi dengan komputer (Khurana dkk., 2023). *NLP* juga mencakup berbagai aplikasi seperti analisis sentimen, *machine translation*, *speech recognition*, dan *question answering* (Otter dkk., 2021).

2.1.3. Deep Learning

Deep Learning merupakan cabang dari *Machine Learning* dan Kecerdasan Buatan (AI) yang berfokus pada penggunaan jaringan saraf tiruan (Artificial Neural Networks, ANN) dengan struktur berlapis-lapis untuk menyelesaikan permasalahan komputasi kompleks. Berbeda dengan pendekatan tradisional *Machine Learning* yang umumnya menggunakan algoritma pembelajaran dengan satu atau sedikit lapisan tersembunyi, *Deep Learning* memanfaatkan *deep neural networks* (jaringan saraf dalam) yang terdiri dari banyak lapisan untuk memodelkan data yang lebih kompleks dan melakukan tugas-tugas pembelajaran dengan performa tinggi. *Deep Learning* telah menjadi inti dari perkembangan teknologi dalam Revolusi Industri Keempat (Industry 4.0), di mana teknologi otomatisasi dan sistem cerdas menjadi bagian penting dalam berbagai aplikasi industri dan penelitian. Salah satu keunggulan *Deep Learning* adalah kemampuannya untuk mempelajari pola-pola dari data yang sangat besar (big data), serta menghasilkan

model yang akurat tanpa perlu proses ekstraksi fitur manual, yang sering kali diperlukan dalam pendekatan *Machine Learning* tradisional (Sarker, 2021).

2.1.4. RoBERTa

RoBERTa adalah model bahasa berbasis Transformer yang dioptimalkan untuk meningkatkan performa *BERT* melalui modifikasi strategi pra-pelatihan, peningkatan skala data, dan penyesuaian hiperparameter. Model ini dikembangkan oleh tim peneliti Facebook AI dan Universitas Washington untuk mengatasi kelemahan *BERT*, seperti penggunaan data yang tidak maksimal dan ketidakefisienan pelatihan. *RoBERTa* menggantikan *static masking BERT* dengan *dynamic masking*, di mana pola masking dihasilkan secara acak setiap kali data diproses. Teknik ini meningkatkan kemampuan model untuk mempelajari representasi konteks yang lebih beragam. Eksperimen menunjukkan bahwa tugas *Next Sentence Prediction* (NSP) pada *BERT* tidak memberikan manfaat signifikan. *RoBERTa* menghilangkan *NSP* dan fokus pada pelatihan *Masked Language Modeling* (MLM) dengan input teks panjang (hingga 512 token) untuk meningkatkan pemahaman intra-kalimat (Liu dkk., 2019).

2.1.5. Data Augmentation

Augmentasi data adalah teknik yang digunakan untuk meningkatkan jumlah dan variasi data pelatihan dengan memodifikasi data yang ada tanpa mengubah makna dasarnya. Teknik ini sangat berguna ketika jumlah data berlabel terbatas, terutama dalam skenario *low-resource* atau tugas baru yang sulit mendapatkan data berlabel dalam jumlah besar. Augmentasi data dalam *NLP* dapat dilakukan pada beberapa tingkatan, antara lain seperti *Token-Level Augmentation*, *Sentence-Level*

Augmentation, Adversarial Data Augmentation dan Hidden-Space Augmentation (J. Chen dkk., 2023).

Token-Level Augmentation berfokus pada modifikasi kata atau frasa dalam kalimat seperti *Synonym Replacement, Random Insertion, Random Deletion* dan *Random Swapping*. Namun, dalam beberapa kasus hal tersebut dapat berisiko mengubah makna asli kalimat terutama jika kata yang diubah memiliki banyak arti atau tergantung pada konteks tertentu. Sementara itu, *Sentence-Level Augmentation* memodifikasi seluruh kalimat untuk menciptakan variasi yang lebih besar, tetapi tetap mempertahankan konteks keseluruhan dari teks seperti *Paraphrasing* dan *Conditional Generation*. Ada juga teknik *Adversarial Data Augmentation*, di mana data baru dibuat dengan mengganggu data asli secara terarah, dengan tujuan "menguji" ketahanan model. Terakhir *Hidden-Space Augmentation*, Augmentasi ini bekerja pada representasi tersembunyi atau internal dari data, seperti *embedding* kata, dengan melakukan gangguan atau interpolasi pada representasi tersebut seperti *Interpolation-Based Methods* (J. Chen dkk., 2023).

2.1.6. T5

T5 (Text-to-Text Transfer Transformer) adalah model transformer yang mengonversi berbagai tugas *NLP* ke dalam format *text-to-text*, input dan output selalu berupa teks, baik untuk klasifikasi, terjemahan, ringkasan, maupun penjawaban pertanyaan. Model ini menggunakan arsitektur *encoder-decoder*, di mana *encoder* memproses input untuk memahami konteks, dan *decoder* menghasilkan teks berdasarkan pemahaman tersebut. Dua komponen utama dari *T5* adalah *Sequence-to-Sequence Language Modeling* (SSLM) dan *Denoising*

Autoencoder (DAE). Pada pendekatan *SSLM*, model dilatih dengan tugas memprediksi token yang disembunyikan (*masked*) dalam urutan teks. Dalam konteks *T5*, input teks diberikan kepada model dengan beberapa bagian yang *masked*. Model kemudian diminta untuk memprediksi token yang hilang tersebut berdasarkan konteks yang ada sebelum dan sesudahnya. Pendekatan ini melibatkan *level token* dan *n-gram*, sehingga model dapat bekerja pada prediksi satu kata atau frasa (Bird dkk., 2023; Rahali & Akhloufi, 2023).

Rumus *SSLM* sebagai pada persamaan 1,

$$Loss^x_{SSLM} = - \frac{1}{|l_s|} \sum_{s=i} \log \left(P \left(\frac{X_s}{X^{Masked}, X_{i:s-1}} \right) \right) \quad (1)$$

Dimana:

x^{Masked} = teks dengan beberapa token disembunyikan

l_s = panjang dari span (rentang kata) yang di-*mask*

T5 juga memanfaatkan metode *DAE*, di mana model dilatih untuk merekonstruksi input teks yang telah dirusak. Ini dilakukan dengan memberikan versi teks yang rusak kepada model, di mana beberapa bagian dari teks asli dihilangkan atau diubah secara acak. Tujuannya adalah agar model dapat mengidentifikasi token atau kata yang hilang dan mengembalikannya ke bentuk aslinya. Pendekatan ini digunakan untuk memperkuat kemampuan model dalam menangani teks yang tidak sempurna, seperti teks yang terpotong atau berisi kesalahan (Rahali & Akhloufi, 2023).

Rumus *DAE* sebagai pada persamaan 2,

$$Loss^x_{DAE} = - \frac{1}{|x|} \sum_{i=1} \log \left(P \left(\frac{x_i}{x_{corr}, x_{<i}} \right) \right) \quad (2)$$

Dimana:

x_{corr} = versi rusak dari input asli

2.1.7. *BABE* (Bias Annotations By Experts)

Dataset *BABE* merupakan kumpulan data teranotasi yang dirancang khusus untuk mendukung penelitian bias media. *BABE* terdiri dari 3.700 kalimat yang diambil dari artikel berita kontroversial yang diterbitkan oleh berbagai platform media di Amerika Serikat antara Januari 2017 hingga Juni 2020. Kalimat-kalimat ini diseimbangkan berdasarkan topik (misalnya, politik, imigrasi) dan outlet media untuk memastikan representasi yang beragam. Setiap kalimat dilabeli untuk bias media pada tingkat kata dan kalimat, memungkinkan analisis granular terhadap penggunaan bahasa yang subjektif. Kualitas *BABE* didukung oleh anotasi yang dilakukan ahli di bidang bias media. Annotator dipilih berdasarkan pengalaman dan keahlian, lalu menjalani pelatihan intensif untuk memastikan konsistensi dan netralitas dalam pelabelan. Pendekatan ini menghasilkan tingkat kesepakatan antar-annotator yang tinggi, yang menjadi pembeda utama dari dataset lain seperti *MBIC* (Spinde dkk., 2021).

2.2. Penelitian Terkait (*State of the Art*)

Tabel 2.1 State of The Art.

No.	Konten	Deskripsi
1.	Judul, Penulis	Neural Media Bias Detection Using Distant Supervision With BABE – Bias Annotations By Experts, (Spinde dkk., 2021)
	Permasalahan	Deteksi bias media yang kompleks, kualitas dataset yang rendah, keterbatasan model deteksi dan kurangnya dataset berkualitas tinggi.
	Metode/Solusi	Pembuatan dataset <i>BABE</i> , anotasi oleh ahli, model deteksi berbasis neural.
	Kontribusi	Dataset berkualitas tinggi, model deteksi yang lebih baik dengan memperkenalkan metode distant supervision. <i>RoBERTa - distant supervision</i> mencapai <i>Weighted F1-Score</i> 0.798 dan <i>BERT – distant supervision</i> mencapai <i>Weighted F1-Score</i> 0.804.
	Batasan	Keterbatasan dataset karena ukurannya masih relatif kecil (3.700 kalimat) untuk pelatihan model deep learning yang lebih kompleks, fokus pada tingkat kalimat dan biaya anotasi.
2.	Judul, Penulis	A Domain-adaptive Pre-training Approach for Language Bias Detection in News, (Krieger dkk., 2022)
	Permasalahan	Deteksi bias dalam berita sulit karena kurangnya dataset standar yang memadai untuk melatih model besar.
	Metode/Solusi	Menggunakan pendekatan <i>pre-training</i> berbasis <i>domain-adaptif</i> pada model <i>BERT</i> , <i>RoBERTa</i> , <i>BART</i> , dan <i>T5</i> .
	Kontribusi	<i>DA-RoBERTa</i> mencapai <i>Weighted F1-Score</i> 0.814 dalam deteksi bias berbasis kalimat, menjadikannya model terdepan untuk tugas ini.
	Batasan	Tidak menggunakan augmentasi data, membatasi kemampuan model menangani keragaman bias media.
3.	Judul, Penulis	MAGPIE: Multi-Task Analysis of Media-Bias Generalization with Pre-Trained Identification of Expressions, (Horych dkk., 2024)
	Permasalahan	Deteksi bias media menimbulkan masalah yang kompleks dan multifaset yang secara tradisional ditangani menggunakan model tugas tunggal dan kumpulan data

No.	Konten	Deskripsi
		kecil dalam domain, akibatnya kurang dapat digeneralisasi.
	Metode/Solusi	Pendekatan pra-pelatihan Multi-Task skala besar pertama yang secara eksplisit disesuaikan untuk deteksi bias media (Magpie). Untuk memungkinkan pra-pelatihan dalam skala besar, menyajikan <i>Large Bias Mixture</i> (LBM), kompilasi dari 59 tugas terkait bias.
	Kontribusi	MAGPIE mengungguli pendekatan sebelumnya dalam deteksi bias media pada kumpulan data Bias Annotation By Experts (BABE), dengan peningkatan relatif sebesar 3,3% <i>Weighted F1-Score</i> . MAGPIE juga berkinerja lebih baik daripada model sebelumnya pada 5 dari 8 tugas dalam Media Bias Identification Benchmark (MBIB).
	Batasan	Tidak melakukan augmentasi data pada dataset untuk meningkatkan metrik evaluasi.
4.	Judul, Penulis	We Can Detect Your Bias: Predicting the Political Ideology of News Articles, (Baly dkk., 2020)
	Permasalahan	Banyak artikel berita menunjukkan kecenderungan ideologi politik yang tidak eksplisit, yang dapat mempengaruhi perilaku pemilih dan polarisasi masyarakat.
	Metode/Solusi	Menggunakan model berbasis <i>transformer</i> dengan <i>adversarial media adaptation</i> atau <i>triplet loss</i> untuk mendeteksi ideologi politik artikel.
	Kontribusi	Model berbasis <i>transformer</i> dengan <i>Triplet Loss Pre-training</i> (TLP) menunjukkan peningkatan signifikan dalam mendeteksi bias politik pada artikel berita, terutama ketika data uji berasal dari sumber media yang tidak pernah dilihat sebelumnya.
	Batasan	Penggunaan anotasi politik yang eksplisit (left, center, right) terlalu sederhana untuk menangani spektrum bias yang lebih luas. Tidak melakukan augmentasi data pada dataset untuk meningkatkan metrik evaluasi.
5.	Judul, Penulis	Experiments in News Bias Detection with Pre-Trained Neural Transformers, (Menzner & Leidner, 2024)
	Permasalahan	Media bias, propaganda, dan <i>fake news</i> menjadi ancaman terhadap demokrasi. Masalahnya adalah deteksi bias pada tingkat kalimat atau dokumen yang tidak akurat.

No.	Konten	Deskripsi
	Metode/Solusi	Membandingkan beberapa model bahasa <i>pre-trained</i> seperti <i>GPT-3.5</i> , <i>GPT-4</i> , dan <i>Meta Llama2</i> untuk deteksi bias pada level kalimat. Termasuk eksperimen <i>fine-tuning</i> .
	Kontribusi	Menunjukkan kinerja yang lebih baik dari model <i>fine-tuned GPT-3.5</i> untuk deteksi bias berbasis kalimat, dengan <i>F1-Score</i> yang tinggi.
	Batasan	Tidak menggunakan augmentasi data, dan hasil deteksi bias tidak sepenuhnya mencakup konteks artikel yang lebih luas. Fokus pada kalimat saja membatasi pemahaman konteks artikel secara keseluruhan.
6.	Judul, Penulis	Exploiting Transformer-based Multitask Learning for the Detection of Media Bias in News Articles, (Spinde dkk., 2022)
	Permasalahan	Bias dalam pemberitaan dapat memengaruhi persepsi publik terhadap suatu peristiwa.
	Metode/Solusi	Menggunakan model distilBERT dengan <i>Multi-Task Learning</i> (MTL) pada enam dataset terkait bias untuk meningkatkan kinerja deteksi bias media.
	Kontribusi	<i>MTL</i> dengan arsitektur <i>transformer</i> meningkatkan kinerja deteksi bias, mencapai <i>Macro F1-Score</i> 0.776, peningkatan 3% dari <i>baseline</i> .
	Batasan	Penggunaan beberapa dataset mungkin memperkenalkan bias tak terdeteksi dalam penggabungan data, terutama untuk <i>genre</i> berita yang berbeda. Tidak menggunakan augmentasi data, sehingga kemampuan model dalam menangani keragaman bias media terbatas.
7.	Judul, Penulis	Mitigation of Diachronic Bias in Fake News Detection Dataset, (Murayama dkk., 2021)
	Permasalahan	Model deteksi berita palsu sering gagal mendeteksi berita baru akibat <i>diachronic</i> bias dari perbedaan periode data pelatihan.
	Metode/Solusi	Menggunakan metode <i>masking</i> berbasis Wikidata untuk mengurangi pengaruh bias nama orang dan meningkatkan kemampuan generalisasi model terhadap data masa depan.
	Kontribusi	Metode <i>masking</i> ini berhasil meningkatkan akurasi deteksi berita palsu pada data di luar domain dan

No.	Konten	Deskripsi
		mengurangi bias kronologis dalam model deteksi berita palsu.
	Batasan	<i>Masking</i> pada <i>proper nouns</i> mungkin tidak efektif dalam menangani bias yang lebih mendalam atau tidak langsung dalam narasi berita. Menggunakan metode <i>masking</i> dengan cara menyembunyikan atau mengaburkan informasi tertentu (seperti nama orang). Namun, tidak melakukan variasi kalimat tambahan lainnya seperti penambahan dataset serupa.
8.	Judul, Penulis	Detecting Political Bias in News Articles Using Headline Attention, (Reddy dkk., 2019)
	Permasalahan	Judul berita sering kali digunakan untuk menyampaikan bias politik secara lebih jelas daripada konten artikel itu sendiri.
	Metode/Solusi	Menggunakan <i>Headline Attention Network</i> (HAN) untuk mendeteksi bias politik dalam artikel berdasarkan perhatian utama yang diberikan pada judul berita.
	Kontribusi	Meningkatkan deteksi bias politik dengan memperhitungkan judul berita sebagai fokus utama, sehingga meningkatkan akurasi prediksi bias secara signifikan.
	Batasan	Terbatas pada bahasa Telugu. Tidak menggunakan teknik augmentasi data.
9.	Judul, Penulis	Revealing Media Bias in News Articles: NLP Techniques for Automated Frame Analysis, (Hamborg, 2023)
	Permasalahan	Penelitian sebelumnya kurang mampu mengidentifikasi bias secara otomatis dan menyampaikan <i>frame</i> yang signifikan dalam berita.
	Metode/Solusi	Mengusulkan metode <i>Person-Oriented Framing Analysis</i> (PFA) untuk mengidentifikasi <i>frame</i> menggunakan teknik <i>NLP</i> , bagaimana media menggambarkan orang dalam berita.
	Kontribusi	<i>PFA</i> berhasil mengidentifikasi <i>frame</i> yang signifikan dalam berita.
	Batasan	Tidak menggunakan augmentasi data, hanya berfokus pada framing analisis.

No.	Konten	Deskripsi
10.	Judul, Penulis	Comparing Topic-Aware Neural Networks for Bias Detection of News, (Jiang dkk., 2020)
	Permasalahan	Deteksi bias dalam berita sering tidak akurat karena bias bersifat halus dan kompleks dalam berbagai topik, khususnya artikel dengan ukuran berbeda.
	Metode/Solusi	Menggunakan model <i>Neural Network</i> berbasis topik seperti <i>LDA-HAN</i> untuk menangani distribusi topik dalam dokumen. Membandingkan kinerja <i>CNN</i> , <i>RNN</i> , <i>LDA-Transformer</i> .
	Kontribusi	Peningkatan akurasi dalam deteksi bias, terutama dengan pendekatan <i>Hierarchical framework</i> yang memperhitungkan distribusi topik dalam dokumen.
	Batasan	Model hanya diuji pada berita dalam bahasa Inggris dan tidak mengeksplorasi berbagai perspektif internasional atau <i>genre</i> berita lainnya. Tidak menggunakan augmentasi data, sehingga kurang mampu menangani variasi bias media.
11.	Judul, Penulis	Identifying Media Bias beyond Words: Using Automatic Identification of Persuasive Techniques for Media Bias Detection, (Rodrigo-Ginés dkk., 2023)
	Permasalahan	Model deteksi bias saat ini sulit menggeneralisasi di berbagai wilayah dan gaya jurnanisme karena kompleksitas bahasa.
	Metode/Solusi	Sistem berbasis <i>transformer</i> dengan identifikasi teknik persuasi retorika menggunakan dataset krisis Ukraina. Pendekatan sistem <i>cascade transformer</i> meningkatkan akurasi deteksi bias media sebesar 6%.
	Kontribusi	Model <i>cascade transformer</i> ini unggul dalam mendeteksi teknik retorika dan persuasi, meningkatkan akurasi deteksi bias media di berbagai negara.
	Batasan	Teknik persuasif yang dianalisis belum cukup menyeluruh dan tidak mencakup berbagai strategi retorika yang lebih halus. Tidak menggunakan augmentasi data.
12.	Judul, Penulis	Detecting Media Bias in News Articles using Gaussian Bias Distributions, (W. F. Chen dkk., 2020)

No.	Konten	Deskripsi
	Permasalahan	Deteksi bias pada level artikel sulit karena informasi bias tidak konsisten antara level kalimat dan artikel, khususnya pada peristiwa baru.
	Metode/Solusi	Menggunakan <i>Gaussian Mixture Model</i> (GMM) untuk mendeteksi distribusi bias berdasarkan posisi, frekuensi, dan urutan kalimat bias di artikel.
	Kontribusi	Mengembangkan pendekatan baru dengan menggunakan informasi bias pada level kalimat untuk meningkatkan deteksi bias artikel secara keseluruhan.
	Batasan	<i>Gaussian Bias Distributions</i> kurang cocok untuk menangani konteks yang lebih kompleks atau topik yang lebih halus dalam berita. Tidak menggunakan augmentasi data, yang membatasi variasi data.
13.	Judul, Penulis	Analyzing Political Bias and Unfairness in News Articles at Different Levels of Granularity, (W.-F. Chen dkk., 2020)
	Permasalahan	Kurangnya pemahaman mendalam tentang bagaimana bias politik dan ketidakadilan dimanifestasikan di berbagai tingkatan teks (kata, kalimat, paragraf).
	Metode/Solusi	Model <i>RNN</i> dengan analisis fitur terbaik untuk mendeteksi bias politik dan ketidakadilan di berbagai tingkatan granularitas teks (kata, kalimat, paragraf, wacana).
	Kontribusi	Memberikan pemahaman mendalam tentang bagaimana bias dan ketidakadilan termanifestasi di berbagai tingkatan granularitas teks.
	Batasan	Model hanya diuji pada tingkat granularitas kata dan wacana dan tidak menggunakan augmentasi data.
14.	Judul, Penulis	Viable Threat on News Reading: Generating Biased News Using Natural Language Models, (Gupta dkk., 2020)
	Permasalahan	Penyalahgunaan model bahasa untuk menghasilkan berita bias yang dapat memengaruhi opini pembaca.
	Metode/Solusi	Menggunakan dan fine-tuning <i>GPT-2</i> dan <i>GROVER</i> untuk menghasilkan berita bias yang terkontrol. <i>RoBERTa</i> digunakan untuk mengevaluasi bias pada teks yang dihasilkan.

No.	Konten	Deskripsi
	Kontribusi	Menunjukkan bahwa model bahasa yang ada dapat dengan mudah diadaptasi untuk menghasilkan berita bias yang terlihat sangat nyata dan sulit dibedakan dengan berita yang ditulis manusia.
	Batasan	Penggunaan model bahasa untuk menghasilkan berita bias hanya diuji pada berita dalam bahasa Inggris, dan mungkin tidak berlaku untuk konteks lintas budaya atau bahasa. Tidak ada augmentasi data atau generalisasi untuk domain yang lebih luas.
15.	Judul, Penulis	Uncovering the essence of diverse media biases from the semantic embedding space, (Huang dkk., 2024)
	Permasalahan	Analisis bias media sering terbatas pada satu jenis bias dan memerlukan upaya manual yang intensif.
	Metode/Solusi	Menggunakan <i>embedding media</i> untuk mengukur bias berdasarkan pemilihan peristiwa dan bahasa, serta menerapkan metode analisis semantik menggunakan <i>Word2Vec</i> dan <i>Truncated SVD</i> .
	Kontribusi	<i>Framework</i> ini mampu mengungkap bias pada skala besar dari berbagai jenis bias, termasuk bias politik, gender, dan pendapatan.
	Batasan	<i>Embedding</i> semantik hanya mampu menangkap hubungan bias yang ada di permukaan dan tidak selalu dapat menangani bias yang lebih halus atau terselubung. Menggunakan <i>up-sampling</i> untuk memastikan setiap korpus media memiliki jumlah artikel yang sama, namun tidak menggunakan augmentasi data tambahan untuk menambah berbagai variasi dataset.
16.	Judul, Penulis	Disentangling Structure and Style: Political Bias Detection in News by Inducing Document Hierarchy, (Hong dkk., 2023)
	Permasalahan	Model deteksi bias politik sebelumnya sering kali bergantung pada gaya penulisan outlet berita, yang menyebabkan <i>overfitting</i> dan terbatasnya generalisasi model.
	Metode/Solusi	Model <i>Hierarki Multi-Head Attention</i> yang memperhitungkan struktur retorika dokumen dan semantik tingkat kalimat.

No.	Konten	Deskripsi
	Kontribusi	Model ini mengatasi masalah ketergantungan domain dan menyoroti struktur wacana jurnalistik dengan lebih baik untuk klasifikasi bias politik yang lebih akurat.
	Batasan	Pemisahan gaya dan struktur dalam model ini bisa terlalu terfokus pada aspek tekstual dan mengabaikan faktor-faktor luar seperti konteks sosial dan politik. Tidak menggunakan augmentasi data, yang membatasi kemampuan model dalam menangani keragaman bias media.
17.	Judul, Penulis	From Pretraining Data to Language Models to Downstream Tasks: Tracking the Trails of Political Biases Leading to Unfair NLP Models, (Feng dkk., 2023)
	Permasalahan	Bias politik dalam data <i>pre-training LMs</i> mempengaruhi tugas <i>downstream</i> seperti deteksi ujaran kebencian dan misinformasi.
	Metode/Solusi	Mengembangkan <i>framework</i> untuk mengukur bias politik di <i>LMs</i> , dan mengevaluasi pengaruh bias politik dalam data <i>pre-training</i> pada performa tugas <i>downstream</i> .
	Kontribusi	Mengidentifikasi bahwa bias politik dalam <i>pre-training LMs</i> mempengaruhi performa model, terutama dalam tugas sosial seperti deteksi ujaran kebencian dan misinformasi.
	Batasan	Metode ini lebih fokus pada bias dalam data <i>pre-training</i> dan tidak sepenuhnya membahas bias yang muncul selama inferensi model dalam berbagai domain. Tidak menggunakan augmentasi data. Pendekatan lebih fokus pada dampak bias politik dan tidak mengeksplorasi teknik augmentasi data untuk meningkatkan generalisasi model.
18.	Judul, Penulis	Media Bias, the Social Sciences, and NLP: Automating Frame Analyses to Identify Bias by Word Choice and Labeling, (Hamborg, 2020)
	Permasalahan	Media bias melalui pemilihan kata dan pelabelan sering sulit dideteksi secara otomatis.
	Metode/Solusi	Menggunakan metode <i>Cross-Document Coreference Resolution</i> (CDCR) dan <i>Target-dependent Sentiment Classification</i> (TSC) untuk mengidentifikasi bias melalui pilihan kata dan pelabelan di berbagai berita.

No.	Konten	Deskripsi
	Kontribusi	Menyediakan kerangka otomatis untuk mendeteksi bias dengan cara mengidentifikasi perbedaan dalam pemilihan kata yang digunakan untuk menggambarkan peristiwa yang sama di berita.
	Batasan	Penelitian hanya berfokus pada pemilihan kata dan pelabelan, tanpa memperhitungkan bias struktural atau framing lebih kompleks. Tidak menggunakan augmentasi data dan teknik yang digunakan terbatas pada deteksi bias berbasis pilihan kata dan pelabelan tanpa mempertimbangkan konteks visual atau multimodalitas.
19.	Judul, Penulis	Media Bias Detecting based on Word Embedding, (Geng, 2022)
	Permasalahan	Kurangnya teknik <i>word embedding</i> yang efektif dan diterima secara luas untuk mendeteksi bias media.
	Metode/Solusi	Membandingkan tiga teknik <i>TF-IDF</i> , <i>Word2Vec</i> , dan <i>Doc2Vec</i> dalam mendeteksi bias di artikel berita, menggunakan <i>SVM</i> dan <i>Random Forest</i> sebagai model klasifikasi.
	Kontribusi	Menemukan bahwa <i>TF-IDF</i> memberikan hasil terbaik dalam mendeteksi bias media, dengan performa lebih baik dibandingkan <i>Word2Vec</i> dan <i>Doc2Vec</i> dalam eksperimen ini.
	Batasan	Penggunaan frekuensi kata sebagai prediktor utama bias mungkin tidak memadai dalam menangani bias yang lebih terselubung atau halus dalam narasi berita. Tidak menggunakan augmentasi data, hanya fokus pada <i>embedding</i> kata dan teknik yang diuji, tanpa mengeksplorasi model terbaru seperti <i>BERT</i> .
20.	Judul, Penulis	Determination of News Biasedness Using Content Sentiment Analysis Algorithm, (Shri Bharathi & Geetha, 2019)
	Permasalahan	Deteksi bias berita merupakan masalah kompleks yang dipengaruhi oleh faktor sosial, politik, dan ekonomi.
	Metode/Solusi	Menggunakan algoritma analisis sentimen untuk mendeteksi bias berdasarkan skor polaritas (+1 hingga -1) dari artikel berita. Sistem ini mengumpulkan berita dari berbagai situs berita di U.S., U.K., dan India.

No.	Konten	Deskripsi
	Kontribusi	Menyediakan sistem yang otomatis mendeteksi bias dalam berita menggunakan analisis sentimen, yang dapat membantu menyediakan informasi yang lebih objektif.
	Batasan	Penggunaan analisis sentimen mungkin terlalu sederhana untuk menangkap dimensi bias yang lebih kompleks dalam berita. Tidak menggunakan augmentasi data, hanya fokus pada analisis sentimen tanpa melibatkan teknik <i>NLP</i> lain atau generalisasi pada dataset yang lebih luas.

Tabel 2.1 menyajikan informasi mengenai judul, penulis, permasalahan, metode/solusi, kontribusi, dan batasan masalah. Penelitian-penelitian tersebut memiliki fokus yang serupa, yaitu untuk deteksi/analisis bias media. Beberapa model yang sudah digunakan sebelumnya antara lain seperti *BERT*, *RoBERTa*, *RNN* hingga *GPT-3.5*. Kemudian, terdapat penelitian yang menggunakan pendekatan lain untuk meningkatkan kinerja model, diantaranya seperti penambahan metode *Distant Supervision Learning*, *Domain Adaptif (DA)*, *Triplet Loss Pre-training (TLP)*, *Multi-Task Learning (MTL)* dan lain-lain. Ekstraksi fitur yang digunakan diantaranya seperti *TF-IDF*, *Word2Vec*, *BoW* hingga ekstraksi fitur modern seperti *Contextualized Word Embeddings*. Adapun dalam transformasi data, penelitian sebelumnya menggunakan beberapa pendekatan seperti metode *Drop*, *Masking* hingga *Up-Sampling*. Selain itu, tabel tersebut menyajikan informasi mengenai batasan atau celah penelitian dari masing-masing penelitian yang telah dipaparkan. Hal tersebut bertujuan untuk memperoleh pemahaman yang lebih mendalam tentang kesenjangan penelitian dalam bidang deteksi/analisis bias media saat ini.

2.3. Matriks Penelitian

(Raza dkk., 2024)	DistilBERT, BERT								√					Precision, Recall, Accuracy, F1-Score, Disparate Impact (DI), Generalized Mean of Bias AUCs (G-AUC)
(Reddy dkk., 2019)	Headline Attention Network						√	√						Accuracy
(Hong dkk., 2023)	BERT, Multi-Head Hierarchical Attention Model									√				AUC, AUC-ROC
(Spinde dkk., 2022)	DistilBERT								√					Precision, Recall, Macro F1-Score, Micro F1-Score, Binary F1-Score, Cross-Entropy (CE) Loss

(Zhu dkk., 2022)	BERT+ENDEF	✓	✓					✓					Accuracy, AUC, Macro F1-Score, spAUC, F1real, F1fake
(Murayama dkk., 2021)	BERTBASE		✓					✓					Accuracy
(W. F. Chen dkk., 2020)	Logistic Regression, Naïve Bayes, SVM							✓		✓			Accuracy
(Feng dkk., 2023)	BERT							✓					F1-Score, Balanced Accuracy (BACC)
(Rodrigo-Ginés dkk., 2023)	DistilBERT, Logistic Regression						✓	✓					Precision, Recall, F1-Score
(Hamborg, 2023)	BERT, RoBERTa							✓					Precision, Recall, Accuracy, F1-Score, Macro F1-Score

(Shri Bharathi & Geetha, 2019)	SVM						√						Precision, Recall, Accuracy
(Nadeem & Raza, 2021)	SimCSE						√	√			√		Accuracy, Spearman Correlation Coefficient
(Geng, 2022)	SVM, Random Forest						√	√					Precision, Recall, F1-Score, Accuracy
(Hamborg, 2020)	BERT							√					F1-Score, Accuracy
(W.-F. Chen dkk., 2020)	GRU							√					F1-Score
(Menzner & Leidner, 2024)	GPT-3.5								√				Precision, Recall, F1-Score
(Huang dkk., 2024)	Word2Vec + LSA + Truncated SVD				√			√				√	Evaluasi berbasis embedding
Our	RoBERTa DA-RoBERTa				√				√				Accuracy Weighted F1-Score (Standar Deviation) Precision Recall

															Loss
--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	------

Tabel 2.2 memberikan perbandingan antara penelitian terdahulu dan penelitian yang akan dilakukan. Dari hasil temuan penelitian-penelitian terdahulu, berbagai pendekatan telah digunakan dalam deteksi/analisis bias media mulai dari model, ekstraksi fitur dan transformasi data. Pendekatan-pendekatan tersebut memiliki relevansi langsung dengan masalah penelitian ini. Namun, masih terdapat kesenjangan penelitian yang perlu di isi, yaitu sebagian besar penelitian sebelumnya belum secara spesifik menargetkan kinerja model dengan augmentasi data berbasis parafrase menggunakan *T5*. Penelitian ini merupakan salah satu upaya untuk mengisi kesenjangan tersebut dengan mengembangkan pendekatan yang lebih optimal dalam deteksi bias media. Penyempurnaan dilakukan terhadap metode sebelumnya yaitu *RoBERTa* dan *DA-RoBERTa* (Krieger dkk., 2022), melalui teknik augmentasi data berbasis parafrase dengan model *T5* pada dataset *BABE*.