

BAB II

TINJAUAN PUSTAKA

2.1. Landasan Teori

2.1.1. Emosi

Emosi merupakan reaksi psikologis dan fisiologis yang muncul sebagai respons terhadap suatu peristiwa, situasi, atau interaksi sosial. Emosi dipandang sebagai fenomena dasar yang bersifat universal, memengaruhi perilaku manusia, dan dapat dikenali melalui ekspresi wajah, suara, maupun bahasa (Ekman, 1992). Dalam komunikasi berbasis teks, seperti di media sosial, emosi biasanya diekspresikan melalui kata, frasa, tanda baca, hingga simbol tertentu.

Berbagai teori telah dikembangkan untuk menjelaskan klasifikasi dasar emosi manusia. Salah satunya adalah teori yang membagi emosi ke dalam enam kategori utama, yaitu *love*, *joy*, *anger*, *sadness*, *fear*, dan *surprise* (Shaver dkk., 1987). Sementara itu, teori hierarki emosi menyusun emosi ke dalam emosi primer, sekunder, dan tersier, dengan enam emosi dasar pada tingkat primer yang serupa, yaitu *joy*, *love*, *surprise*, *anger*, *sadness*, dan *fear* (Parrott, 2001). Kedua teori ini menjadi landasan umum dalam banyak penelitian yang berkaitan dengan analisis emosi, baik di bidang psikologi maupun pemrosesan bahasa alami.

Dalam konteks penelitian ini, konsep emosi dipahami sebagai kategori psikologis yang dapat direpresentasikan melalui teks. Dengan adanya kerangka teori emosi tersebut, analisis emosi pada media sosial dapat difokuskan pada dimensi emosi yang paling mendasar sehingga lebih mudah diukur dan diolah menggunakan metode komputasi.

2.1.2. Klasifikasi Emosi

Dalam ranah pemrosesan bahasa alami (*Natural Language Processing*), klasifikasi emosi didefinisikan sebagai tugas untuk mengelompokkan teks ke dalam kategori emosi tertentu berdasarkan ekspresi linguistik yang terkandung di dalamnya. Berbeda dengan analisis sentimen yang hanya berfokus pada polaritas positif, negatif, atau netral, klasifikasi emosi memungkinkan identifikasi variasi emosional yang lebih kaya seperti marah, sedih, atau bahagia (Mohammad, 2015; Sailunaz & Alhajj, 2019). Tugas ini menjadi penting karena memberikan pemahaman yang lebih mendalam mengenai aspek psikologis dalam komunikasi manusia, khususnya pada teks yang bersumber dari media sosial.

Kategori emosi yang digunakan dalam penelitian ini didasarkan pada teori psikologi yang telah banyak digunakan dalam studi emosi, khususnya teori enam emosi dasar (Parrott, 2001; Shaver dkk., 1987). Dalam implementasinya pada analisis teks media sosial, salah satu adaptasi yang umum dilakukan adalah menambahkan kategori *neutral* untuk mewakili teks yang tidak mengekspresikan emosi tertentu.

Kategori inilah yang kemudian diadopsi dalam *Emotion Dataset from Indonesian Public Opinion* (Riccosan dkk., 2022), yaitu *anger*, *fear*, *joy*, *love*, *sadness*, dan *neutral*. Dengan menggunakan kategori ini, penelitian dapat memfokuskan analisis pada dimensi emosi yang mendasar, sekaligus tetap praktis untuk diterapkan dalam sistem NLP.

2.1.3. Aplikasi Twitter (X)

Aplikasi X, atau lebih dikenal sebagai Twitter, adalah salah satu platform media sosial terkemuka yang didirikan pada tahun 2006. Twitter memungkinkan pengguna untuk membuat dan membagikan pesan singkat yang disebut "tweet" dalam format teks, gambar, video, atau tautan. Dengan batasan jumlah karakter yang relatif kecil (280 karakter per tweet), Twitter mempromosikan komunikasi yang singkat dan langsung (Antonakaki dkk., 2021).

Twitter telah menjadi salah satu sumber utama informasi dan diskusi tentang berbagai topik, termasuk politik, berita terkini, hiburan, dan lainnya. Pengguna Twitter dapat mengikuti akun-akun yang relevan dengan minat mereka dan berinteraksi dengan pesan yang diposting oleh akun-akun tersebut melalui mekanisme retweet, balasan, dan suka. Keunggulan Twitter sebagai platform media sosial dalam konteks penelitian ini adalah kemampuannya untuk memberikan akses real-time terhadap berbagai peristiwa dan topik yang sedang tren (Yeung dkk., 2021). Dengan analisis sentimen yang dilakukan terhadap data Twitter, peneliti dapat memperoleh pemahaman yang mendalam tentang opini, sikap, dan pandangan masyarakat terkait dengan berbagai isu yang berkembang.

Penelitian ini, data dari Aplikasi X (Twitter) akan menjadi sumber utama informasi untuk melakukan analisis emosi terhadap data berbahasa Indonesia. Dengan memahami dinamika dan tren yang terjadi di Twitter, penelitian ini akan memberikan wawasan berharga tentang pola emosi dan respons masyarakat terhadap berbagai topik yang sedang dibahas secara luas (Rahmat & Rafi, 2022).

2.1.4. *Natural Language Processing (NLP)*

Natural Language Processing (NLP) adalah cabang dari kecerdasan buatan yang berfokus pada interaksi antara komputer dan bahasa manusia, baik dalam bentuk teks maupun suara. NLP bertujuan untuk memungkinkan komputer menganalisis, memahami, dan menghasilkan bahasa alami manusia secara cerdas dan bermanfaat. Aplikasi umum dari NLP meliputi pengenalan suara, penerjemahan mesin, analisis sentimen, dan pengenalan entitas bernama (Agarwal, 2019).

Awalnya, pendekatan NLP banyak bergantung pada aturan-aturan linguistik yang kompleks, tetapi dengan kemajuan teknologi pembelajaran mesin dan deep learning, pendekatan berbasis statistik kini lebih umum digunakan. Model transformer, seperti BERT dan GPT, merevolusi bidang NLP dengan kemampuannya memahami konteks kata dalam urutan teks secara dua arah (bidirectional), yang memungkinkan pemahaman konteks yang lebih mendalam (Hirschberg & Manning, 2015).

Pada NLP modern, berbagai teknik digunakan untuk memproses teks, termasuk tokenisasi, penghilangan stop words, stemming, dan lemmatisasi. Langkah-langkah ini memungkinkan model untuk mengidentifikasi fitur-fitur penting dalam teks dan mengubah teks mentah menjadi bentuk yang lebih terstruktur. Setelah tahap praproses, algoritma machine learning dan deep learning digunakan untuk menyelesaikan berbagai tugas, seperti klasifikasi teks dan analisis emosi (Geetha dkk., 2023).

Pada teks bahasa alami, seperti yang digunakan dalam media sosial, layanan pelanggan, dan data medis, sehingga menghasilkan wawasan yang bermanfaat di berbagai bidang. Model NLP berbasis transformer, seperti BERT dan GPT, memungkinkan komputer untuk memproses data teks dalam skala besar dan memahami konteks yang lebih kompleks secara lebih efektif dibandingkan metode konvensional (Qiu dkk., 2020).

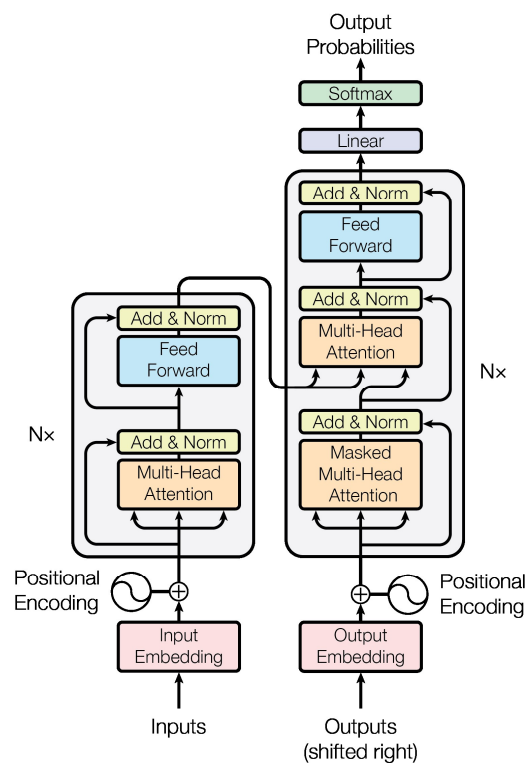
2.1.5. *Transformer*

Transformer adalah arsitektur deep learning yang diperkenalkan oleh Vaswani et al. (2017) dalam penelitian berjudul "Attention Is All You Need". Model ini dirancang untuk menangani tugas-tugas Natural Language Processing (NLP) dengan efisiensi lebih tinggi dibandingkan model berbasis Recurrent Neural Networks (RNN) atau Long Short-Term Memory (LSTM). Keunggulan utama Transformer adalah kemampuannya dalam memproses data secara paralel dan memahami hubungan antar kata dalam suatu teks secara lebih efektif melalui mekanisme self-attention (Vaswani et al., 2017).

Transformer terdiri dari dua komponen utama, yaitu encoder dan decoder, yang masing-masing terdiri dari beberapa lapisan (stacked layers). Arsitektur ini awalnya digunakan untuk tugas penerjemahan mesin (machine translation), tetapi saat ini telah diadaptasi untuk berbagai tugas NLP lainnya (Devlin et al., 2019).

Arsitektur Transformer terdiri dari dua komponen utama, yaitu encoder dan decoder, yang masing-masing terdiri dari beberapa lapisan (stacked layers). Encoder bertanggung jawab untuk memproses input teks dan mengubahnya

menjadi representasi yang lebih abstrak, sementara decoder mengambil representasi tersebut dan menggunakannya untuk menghasilkan output dalam bentuk teks yang diprediksi. Kedua komponen ini berinteraksi dengan bantuan mekanisme self-attention dan feed-forward networks untuk memastikan pemrosesan yang optimal (Vaswani dkk., 2017). Arsitektur Transformer bisa dilihat pada gambar 2.2.



Gambar 2. 1 Arsitektur Transformer (Vaswani dkk., 2017)

Pemrosesan dalam Transformer dimulai dengan positional encoding, yang menambahkan informasi posisi ke dalam embedding kata-kata input. Karena Transformer tidak memiliki urutan eksplisit seperti RNN, positional encoding

menggunakan fungsi sinus dan cosinus untuk memberikan sinyal urutan kata dalam teks. Positional encoding dihitung menggunakan persamaan berikut:

$$PE(pos, 2i) = \sin\left(\frac{pos}{10000^{2i/d}}\right) PE(pos, 2i + 1)$$

$$= \cos\left(\frac{pos}{10000^{2i/d}}\right)$$

Persamaan 2. 1 Positional Encoding

Dimana:

pos = posisi kata dalam kalimat

i = indeks dimensi embedding

d = dimensi keseluruhan embedding

Dengan cara ini, positional encoding memungkinkan Transformer untuk mempertahankan informasi urutan kata dalam kalimat tanpa memerlukan mekanisme berbasis waktu seperti RNN.

Setelah positional encoding diterapkan, representasi ini masuk ke dalam beberapa lapisan encoder. Setiap lapisan encoder terdiri dari dua komponen utama, yaitu self-attention mechanism dan feed-forward neural network (FFNN). Mekanisme self-attention memungkinkan model untuk mempertimbangkan hubungan antara setiap kata dalam kalimat, terlepas dari jaraknya. Self-attention dihitung menggunakan persamaan berikut:

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

Persamaan 2. 2 Self-attention

Di mana Q (Query), K (Key), dan V (Value) merupakan representasi vektor dari kata-kata dalam kalimat, sedangkan dk adalah faktor normalisasi untuk menghindari nilai terlalu besar dalam softmax.

Setelah self-attention diterapkan, hasilnya diproses melalui feed-forward neural network (FFNN), yang terdiri dari dua lapisan linear dengan aktivasi non-linear di antaranya. Lapisan ini memungkinkan transformasi representasi data menjadi lebih kompleks sebelum diteruskan ke lapisan berikutnya.

Decoder bekerja dengan cara yang mirip dengan encoder, tetapi memiliki satu perbedaan utama: Masked Self-Attention. Masking ini mencegah model melihat token masa depan selama pelatihan, memastikan bahwa prediksi hanya bergantung pada informasi sebelumnya. Setelah melalui beberapa lapisan masked self-attention dan feed-forward networks, output dari decoder dikombinasikan dengan informasi dari encoder melalui encoder-decoder attention, yang memungkinkan model memahami hubungan antara input dan output yang dihasilkan.

Langkah terakhir dalam arsitektur Transformer adalah lapisan softmax, yang menghasilkan probabilitas untuk setiap kata dalam kosakata. Kata dengan probabilitas tertinggi dipilih sebagai prediksi akhir, membentuk output teks yang dihasilkan oleh model.

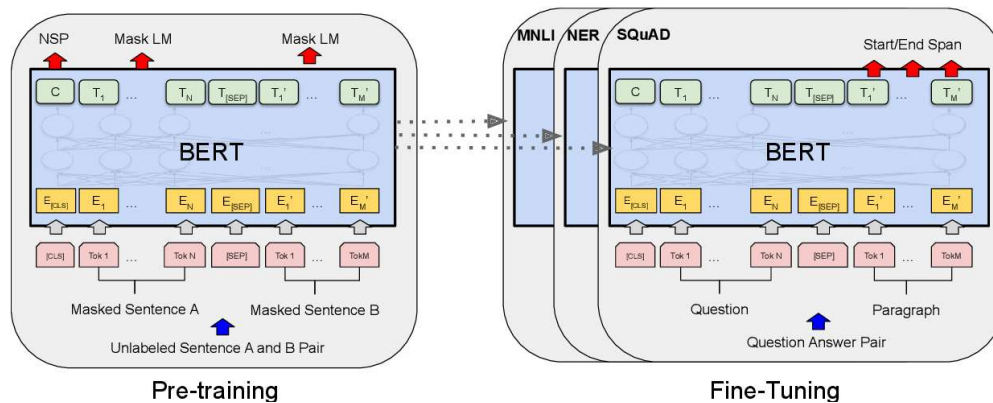
Dengan keunggulan-keunggulan ini, Transformer telah menjadi standar dalam NLP modern, termasuk dalam tugas klasifikasi emosi seperti yang dilakukan dalam penelitian ini (Vaswani dkk., 2017).

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

Bidirectional Encoder Representations from Transformers (BERT) adalah model berbasis Transformer yang dirancang untuk memahami teks dalam konteks bidirectional. Berbeda dengan model sebelumnya yang hanya membaca teks dari kiri ke kanan atau dari kanan ke kiri, BERT membaca teks secara simultan ke kedua arah, memungkinkan model untuk memahami hubungan antar kata lebih akurat dalam suatu kalimat (Devlin dkk., 2019). Pendekatan bidirectional ini memberikan keuntungan signifikan karena BERT dapat menangkap informasi kontekstual yang lebih lengkap dari teks, yang tidak dapat dicapai oleh model unidirectional.

BERT menggunakan arsitektur Transformer, yang terdiri dari beberapa lapisan encoder yang memungkinkan model untuk menangkap hubungan antar kata dalam teks dengan lebih mendalam (Vaswani dkk., 2017). Arsitektur ini bekerja dengan menggunakan mekanisme self-attention, di mana setiap kata dalam sebuah kalimat dapat melihat semua kata lain secara bersamaan, bukan hanya kata sebelum atau sesudahnya. Hal ini membuat model lebih efektif dalam menangkap makna dan konteks dari suatu kalimat dibandingkan dengan pendekatan berbasis RNN atau LSTM yang hanya memproses teks secara sekuensial. Selain itu, mekanisme self-attention ini juga memungkinkan model untuk memproses kata-kata dalam kalimat secara paralel, tanpa bergantung pada urutan sekuensial. Pendekatan paralel ini tidak hanya meningkatkan kemampuan model dalam memahami konteks secara

global, tetapi juga meningkatkan efisiensi dalam pelatihan dan inferensi, sehingga model dapat belajar lebih cepat dan menangani lebih banyak data dalam waktu yang lebih singkat.



Gambar 2. 2 Arsitektur BERT (Devlin dkk., 2019)

Pada Gambar 2.3, pada arsitektur BERT menggunakan dua teknik utama, yaitu Masked Language Modeling (MLM) dan Next Sentence Prediction (NSP) (Devlin dkk., 2019). Pada Masked Language Modeling, sekitar 15% dari token dalam sebuah teks disembunyikan (masked), dan model dilatih untuk memprediksi token yang hilang berdasarkan konteks kiri dan kanan. Teknik ini memungkinkan model untuk belajar memahami hubungan antara kata-kata dalam kalimat secara lebih efektif. Selain itu, dalam Next Sentence Prediction, model dilatih untuk memprediksi apakah dua kalimat memiliki hubungan sekuensial atau tidak. Model diberi dua kalimat sebagai input dan harus menentukan apakah kalimat kedua memang merupakan kelanjutan dari kalimat pertama atau tidak.

BERT memiliki dua varian utama yang digunakan untuk berbagai tugas NLP, yaitu BERT-Base dan BERT-Large. BERT-Base memiliki 12 lapisan

encoder, 768 hidden size, dan 12 attention heads, dengan total 110 juta parameter. Sementara itu, BERT-Large memiliki 24 lapisan encoder, 1024 hidden size, dan 16 attention heads, dengan total 340 juta parameter (Devlin et al., 2019). Model ini dapat diadaptasi untuk berbagai tugas NLP melalui proses fine-tuning, yang memungkinkan model yang sudah dilatih sebelumnya untuk disesuaikan dengan data spesifik tanpa perlu pelatihan ulang dari awal (Sun dkk., 2019).

Dalam implementasinya, BERT menangani input teks dengan membaginya menjadi token melalui BERT Tokenizer, yang menggunakan teknik *WordPiece Tokenization*. Teknik ini memungkinkan kata-kata yang jarang muncul untuk dipecah menjadi sub-token yang lebih kecil, sehingga model dapat menangani kata-kata tak dikenal dengan lebih efektif (Wu dkk., 2016). Model menerima input dalam bentuk urutan token yang dikodekan dalam vektor, yang kemudian diproses oleh lapisan self-attention dalam Transformer. Setiap token juga diberikan representasi tambahan berupa segment embeddings dan positional embeddings, yang membantu model memahami hubungan antar kata dan urutan dalam kalimat (Devlin dkk., 2019).

2.1.7. IndoBERTweet

IndoBERTweet adalah model transformer berbasis BERT (*Bidirectional Encoder Representations from Transformers*) yang dirancang khusus untuk memproses teks dalam Bahasa Indonesia, terutama yang bersumber dari media sosial seperti Twitter. Model ini dikembangkan untuk mengatasi keterbatasan model multibahasa seperti mBERT dalam menangkap karakteristik unik bahasa

Indonesia, termasuk gaya bahasa informal, slang, dan struktur linguistik khas media sosial (Koto dkk., 2021).

IndoBERTweet dibangun dengan pendekatan yang memungkinkan adaptasi efisien terhadap perbedaan kosakata antara domain data. Hal ini dilakukan melalui teknik inisialisasi kosakata domain-spesifik yang mempercepat proses pelatihan dan meningkatkan kinerja model dalam berbagai tugas berbasis teks media sosial. Model ini menunjukkan performa yang signifikan dalam klasifikasi teks, analisis sentimen, dan pengenalan entitas bernama (NER) untuk data Twitter berbahasa Indonesia (Koto dkk., 2021).

Keunggulan *IndoBERTweet* terletak pada kemampuannya untuk menangani teks tidak formal yang sering ditemukan di Twitter. Penelitian menunjukkan bahwa model ini lebih efektif dibandingkan dengan transformer generik dalam memahami konteks bahasa Indonesia di Twitter. Sebagai contoh, *IndoBERTweet* secara signifikan mengurangi masalah miskomunikasi yang disebabkan oleh slang atau singkatan dalam analisis sentimen dan klasifikasi emosi (Koto dkk., 2021).

Meskipun menawarkan banyak keunggulan, *IndoBERTweet* juga memiliki keterbatasan. Model ini memerlukan sumber daya komputasi yang besar untuk pelatihan dan inferensi, dan performanya bergantung pada kualitas dataset pelatihan. Namun demikian, *IndoBERTweet* tetap menjadi salah satu model terbaik untuk tugas pemrosesan bahasa alami di Indonesia, khususnya pada data berbasis media sosial.

2.1.8. *Imbalanced Data*

Imbalanced data adalah kondisi ketika distribusi kelas dalam sebuah dataset tidak merata. Satu atau lebih kelas memiliki jumlah sampel yang jauh lebih sedikit dibandingkan kelas lainnya. Situasi ini umum terjadi pada berbagai tugas klasifikasi, termasuk klasifikasi emosi, di mana beberapa emosi seperti *joy* atau *sadness* lebih dominan, sedangkan emosi lain seperti *surprise* atau *fear* cenderung jarang muncul (He & Garcia, 2009).

Ketidakseimbangan data menimbulkan sejumlah permasalahan serius dalam pembelajaran mesin. Pertama, model cenderung menghasilkan akurasi tinggi hanya dengan memprediksi kelas mayoritas, meskipun performanya buruk pada kelas minoritas. Kedua, model menjadi bias terhadap kelas mayoritas sehingga gagal mengenali pola yang jarang muncul pada kelas minoritas. Ketiga, kondisi ini mengurangi kemampuan model untuk melakukan generalisasi, sehingga performa pada data baru dapat menurun drastis (Haixiang dkk., 2017).

Berbagai pendekatan telah dikembangkan untuk mengatasi masalah ini dan secara garis besar terbagi menjadi dua, yaitu pendekatan pada tingkat data, seperti oversampling, undersampling, dan *Synthetic Minority Over-sampling Technique* (SMOTE) yang menambah data sintetis untuk kelas minoritas (Chawla dkk., 2002), serta pendekatan pada tingkat algoritma, seperti *cost-sensitive learning* dan modifikasi fungsi loss (misalnya *focal loss*) yang memberikan bobot lebih besar pada kelas minoritas tanpa mengubah distribusi data (Krawczyk, 2016).

Dalam konteks klasifikasi emosi di media sosial, masalah imbalanced data menjadi tantangan penting karena distribusi emosi tidak merata. Oleh karena itu,

diperlukan strategi khusus untuk memastikan model mampu mengenali seluruh kategori emosi, termasuk kelas yang jarang muncul.

2.1.9. *Focal Loss*

Focal Loss adalah sebuah modifikasi dari fungsi *loss Cross-Entropy* standar yang dirancang khusus untuk mengatasi masalah ketidakseimbangan kelas (*class imbalance*) yang ekstrem dalam tugas klasifikasi. Diperkenalkan oleh (Lin dkk., 2018), pendekatan ini sangat relevan untuk dataset di mana jumlah sampel pada kelas mayoritas jauh lebih banyak daripada kelas minoritas, seperti yang sering ditemukan pada data ulasan pelanggan atau deteksi emosi.

Dalam klasifikasi tradisional yang menggunakan *Cross-Entropy Loss*, semua sampel memberikan kontribusi yang sama terhadap total *loss*, terlepas dari apakah sampel tersebut mudah atau sulit untuk diklasifikasikan. Akibatnya, sejumlah besar sampel yang mudah dari kelas mayoritas dapat mendominasi proses training dan membuat model mengabaikan kelas minoritas yang lebih sulit.

Focal Loss mengatasi masalah ini dengan menambahkan dua mekanisme utama pada *loss function*. Mekanisme pertama adalah pembobotan kelas (α), yang mirip dengan *cost-sensitive learning* dalam menyeimbangkan pentingnya setiap kelas secara statis. Namun, inovasi utamanya terletak pada mekanisme kedua, yaitu faktor pemfokus (γ). Faktor ini secara dinamis mengurangi kontribusi *loss* dari sampel-sampel yang sudah terklasifikasi dengan baik (sampel mudah), sehingga memaksa model untuk lebih memfokuskan proses pembelajarannya pada sampel-

sampel yang sulit dan sering salah diklasifikasikan, yang umumnya berasal dari kelas minoritas.

Fungsi *Cross-Entropy Loss* standar dapat ditulis sebagai berikut:

$$CE(P_t) = -\log(P_t)$$

Persamaan 2. 3 *Cross-Entropy* (Lin dkk., 2018)

Focal Loss memodifikasi persamaan ini menjadi:

$$FL(P_t) = -\alpha(1 - P_t)^\gamma \log(P_t)$$

Persamaan 2. 4 *Focal Loss* (Lin dkk., 2018)

Dimana:

P_t = Probabilitas hasil prediksi model untuk kelas yang benar.

γ (gamma) = Faktor pemfokus (focusing parameter). Ketika $\gamma > 0$, kontribusi *loss* dari sampel yang mudah (P_t mendekati 1) akan dikurangi secara signifikan, sehingga model fokus pada sampel yang sulit (P_t mendekati 0).

α (alpha) = Faktor penyeimbang (balancing factor) yang mengatasi ketidakseimbangan kelas secara langsung dengan memberikan bobot yang berbeda untuk setiap kelas.

Dengan mekanisme ini, *Focal Loss* memberikan solusi yang lebih dinamis dan efektif dibandingkan sekadar pembobotan kelas statis, menjadikannya pendekatan yang kuat untuk meningkatkan kinerja model pada dataset yang sangat tidak seimbang.

2.1.10. *Confusion Matrix*

Confusion Matrix adalah tabel yang digunakan untuk mengevaluasi kinerja model klasifikasi dengan membandingkan hasil prediksi model terhadap label yang sebenarnya. Matriks ini terdiri dari 4 kategori yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN), yang mencerminkan jumlah prediksi yang benar dan salah pada setiap kelas. Matriks ini berguna untuk memahami kesalahan model secara rinci dan menghitung berbagai metrik evaluasi. Beberapa metrik yang digunakan pada confusion matrix yaitu *accuracy*, *precision*, *recall*, dan *f1-score*.

Akurasi mengukur seberapa tepat model dalam memprediksi kelas yang benar, baik untuk kelas positif maupun negatif. Metrik ini dinyatakan sebagai rasio jumlah prediksi yang benar (True Positive dan True Negative) terhadap total jumlah data. Akurasi memberikan gambaran umum kinerja model secara keseluruhan, tetapi metrik ini kurang representatif dalam kasus data tidak seimbang. Rumus akurasi dapat dinyatakan sebagai berikut:

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

Persamaan 2. 5 Akurasi

Precision mengukur seberapa tepat model dalam memprediksi kelas positif, yaitu rasio jumlah prediksi positif yang benar dibandingkan dengan total prediksi positif. Precision berguna untuk mengukur akurasi model dalam kasus di mana kesalahan prediksi positif palsu (False Positive) harus diminimalkan. Rumus precision dapat dituliskan sebagai berikut:

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

Persamaan 2. 6 Presisi

Recall, atau disebut juga sensitivitas, mengukur kemampuan model untuk mendeteksi semua data positif yang sebenarnya ada. Recall dihitung sebagai rasio jumlah data positif yang benar-benar terdeteksi (True Positive) dibandingkan dengan total jumlah data positif yang ada. Rumus recall adalah:

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

Persamaan 2. 7 Recall

F1-Score adalah rata-rata harmonis antara precision dan recall, memberikan keseimbangan antara kemampuan model dalam memprediksi dengan tepat (precision) dan mendeteksi semua data positif (recall). F1-Score berguna dalam evaluasi model pada data yang tidak seimbang, karena mempertimbangkan kedua metrik tersebut secara proporsional. Rumus F1-Score adalah:

$$\textbf{F1} = \frac{2 \times \textbf{Precision} \times \textbf{Recall}}{\textbf{Precision} + \textbf{Recall}} \quad 2.8$$

Persamaan 2. 8 F1-score

2.2. Penelitian Terkait

Tabel 2. 1 Penelitian terkait

No.	Konten	Deskripsi
-----	--------	-----------

1	Judul paper, Penulis	Cost-Sensitive BERT for Generalisable Sentence Classification with Imbalanced Data (Madabushi dkk., 2020)
	Permasalahan	Penelitian ini mengidentifikasi masalah ketidakseimbangan kelas dalam klasifikasi teks, di mana kategori minoritas cenderung diabaikan oleh model. Selain itu, penelitian ini menghadapi tantangan generalisasi model ketika data pelatihan dan pengujian berbeda secara signifikan, seperti dalam kasus artikel berita yang topiknya sering berubah.
	Kontribusi	Studi ini memperkenalkan teknik sensitivitas biaya pada model BERT untuk meningkatkan performa klasifikasi pada dataset yang tidak seimbang dan tidak serupa. Penelitian ini juga memberikan metode untuk mengukur kesamaan dataset secara statistik menggunakan uji Wilcoxon. Selain itu, penelitian ini mengevaluasi efektivitas teknik augmentasi data pada model berbasis BERT.
	Metode/ Solusi	Penelitian menggunakan model BERT dengan penyesuaian sensitivitas biaya pada fungsi kehilangan (loss function) untuk meningkatkan akurasi klasifikasi pada kelas minoritas. Teknik augmentasi data seperti oversampling, synonym insertion, dan random word deletion dievaluasi untuk mengatasi ketidakseimbangan kelas. Uji statistik Wilcoxon digunakan untuk mengukur kesamaan antara dataset pelatihan dan pengujian. Metode ini diterapkan pada dataset Propaganda Techniques Corpus untuk tugas klasifikasi propaganda.
	Batasan	Studi ini terbatas pada klasifikasi teks propaganda, sehingga penerapannya pada domain lain (misalnya, klasifikasi emosi) memerlukan adaptasi tambahan.
2	Judul paper, Penulis	Indonesian Tweet Emotion Detection Using IndoBERT (Masaling dkk., 2024)
	Permasalahan	Penelitian ini mengangkat permasalahan deteksi emosi pada teks berbahasa Indonesia, khususnya pada data dari media sosial Twitter yang memiliki karakteristik teks pendek, informal, dan bernuansa kompleks. Hal ini menjadi tantangan dalam bidang Pemrosesan Bahasa Alami (NLP) untuk dapat mengidentifikasi emosi secara akurat.
	Kontribusi	Kontribusi utama penelitian ini adalah menyajikan analisis komprehensif mengenai kinerja model IndoBERT untuk tugas deteksi emosi pada tweet berbahasa Indonesia. Penelitian ini memvalidasi efektivitas IndoBERT dengan metrik evaluasi yang rinci, sehingga menambah pemahaman tentang penerapan praktis model bahasa modern untuk

		aplikasi di dunia nyata, khususnya pada bidang analisis sentimen dan kecerdasan emosional.
	Metode/ Solusi	Metode yang digunakan adalah model transformer-based IndoBERT yang dilatih secara spesifik selama 50 epoch. Untuk mengukur keberhasilan solusi, kinerja model dievaluasi menggunakan metrik standar dan berhasil mencapai skor akurasi 0.74, dengan rata-rata presisi, recall, dan F1-score masing-masing sebesar 0.79, 0.78, dan 0.78.
	Batasan	Batasan utama dalam penelitian ini adalah fokus spesifik pada evaluasi model IndoBERT untuk tugas deteksi emosi yang hanya menggunakan dataset dari cuitan (tweet) berbahasa Indonesia. Konsekuensinya, penelitian ini tidak akan membandingkan performa IndoBERT dengan arsitektur model lain dan hasil dari model yang dikembangkan memiliki generalisasi yang terbatas, artinya kemampuannya mungkin tidak optimal jika diterapkan pada sumber teks lain di luar Twitter (seperti Facebook, artikel berita, atau ulasan produk) atau pada teks yang mengandung bahasa selain bahasa Indonesia.
3	Judul paper, Penulis	Emotion Classification in Indonesian Language: A CNN Approach with Hyperband Tuning (Baihaqi dkk., 2023)
	Permasalahan	Penelitian ini mengatasi masalah rendahnya akurasi metode klasifikasi emosi yang ada untuk Bahasa Indonesia dan kesulitan dalam melakukan <i>tuning hyperparameter</i> pada model <i>deep learning</i> seperti CNN. Kebutuhan akan teknik yang andal sangat tinggi, namun metode sebelumnya seringkali tidak mencapai performa yang memuaskan.
	Kontribusi	Kontribusi utama penelitian ini adalah penerapan model <i>Convolutional Neural Network</i> (CNN) yang dioptimalkan dengan <i>Hyperband Tuner</i> (HT) untuk meningkatkan akurasi klasifikasi emosi. Studi ini juga mengevaluasi berbagai teknik ekstraksi fitur dan menunjukkan bahwa kombinasi <i>Keras Tokenizer</i> (KT) dengan CNN yang dioptimalkan berhasil mengungguli metode-metode sebelumnya, dengan akurasi mencapai 71,5655%.
	Metode/ Solusi	Solusi yang diusulkan adalah menggunakan dataset 7.080 kalimat berbahasa Indonesia yang telah melalui pra-pemrosesan data seperti <i>stemming</i> dan penghapusan <i>stop words</i> . Model CNN kemudian diterapkan, dengan proses optimasi <i>hyperparameter</i> dilakukan secara otomatis oleh <i>Hyperband Tuner</i> (HT). Kombinasi terbaik kemudian dievaluasi menggunakan metrik akurasi, presisi, <i>recall</i> , dan F1-score.

	Batasan	Meskipun hasilnya unggul, penelitian ini memiliki batasan. Tantangan utamanya meliputi kelangkaan sumber daya linguistik untuk Bahasa Indonesia, kesulitan memprediksi kelas minoritas akibat data yang tidak seimbang, kompleksitas klasifikasi enam emosi, dan adanya <i>noise</i> dari kalimat yang panjang dalam dataset.
4	Judul paper, Penulis	Emotion dataset from Indonesian public opinion (Riccosan dkk., 2022)
	Permasalahan	kelangkaan dataset emosi berbahasa Indonesia yang tersedia untuk publik. Hal ini menghambat kemajuan riset analisis sentimen, meskipun Twitter di Indonesia merupakan sumber data opini yang sangat potensial.
	Kontribusi	Penelitian ini berkontribusi dengan merilis dataset publik baru berisi 7.080 tweet opini publik Indonesia. Dataset ini telah dianotasi dengan enam label emosi (marah, takut, senang, cinta, sedih, dan netral) dan dapat diakses secara bebas melalui GitHub untuk mendukung penelitian selanjutnya.
	Metode/ Solusi	Dataset ini dibuat dengan mengumpulkan tweet berbahasa Indonesia menggunakan kata kunci spesifik per emosi. Data tersebut kemudian dibersihkan dari duplikat dan hashtag namun dengan sengaja tidak menghilangkan stop-word. Proses anotasi dilakukan oleh tiga orang, dan kualitasnya divalidasi dengan statistik Kappa yang menunjukkan tingkat kesepakatan moderat.
	Batasan	Dataset ini memiliki beberapa batasan. Pertama, sumbernya secara eksklusif hanya dari Twitter, sehingga mungkin tidak mewakili bahasa di domain lain. Kedua, tingkat kesepakatan antar-anotator hanya berada pada level "moderat", yang mengindikasikan adanya subjektivitas label. Ketiga, dataset tidak melalui normalisasi
5	Judul paper, Penulis	Addressing Class Imbalance in Customer Review: Analysis Using <i>Focal Loss</i> and SVM with BERT (Li & Shimada, 2024)
	Permasalahan	ketidakseimbangan kelas pada dataset ulasan pelanggan, di mana ulasan penting seperti "Permintaan" dan "Keluhan" sangat minoritas. Hal ini menyebabkan model klasifikasi menjadi bias dan tidak efektif, terutama saat dataset berukuran kecil.
	Kontribusi	Kontribusi utamanya adalah mengusulkan metode baru yang mengintegrasikan <i>Focal Loss</i> saat fine-tuning BERT dan mengganti lapisan klasifikasi akhirnya dengan SVM. Pendekatan ini terbukti secara signifikan meningkatkan

		performa klasifikasi untuk kelas "Permintaan" yang sangat minoritas.
	Metode/ Solusi	Solusinya adalah melakukan fine-tuning model BERT pada dataset ulasan yang tidak seimbang menggunakan <i>Focal Loss</i> untuk meningkatkan fokus pada kelas minoritas. Kemudian, embedding dari token [CLS] diekstraksi dan digunakan sebagai input untuk SVM, yang lebih efektif pada dataset kecil.
	Batasan	Batasan utamanya adalah meskipun metode ini signifikan meningkatkan performa kelas "Permintaan", performa untuk kelas minoritas lainnya, yaitu "Keluhan", justru sedikit menurun. Hal ini menunjukkan bahwa <i>Focal Loss</i> mungkin terlalu fokus pada kelas yang paling minoritas sehingga mengabaikan kelas minoritas lainnya.
6	Judul paper, Penulis	Text Mining and Emotion Classification on Monkeypox Twitter Dataset: A Deep Learning-Natural Language Processing (NLP) Approach (Olusegun dkk., 2023)
	Permasalahan	Penelitian ini mengidentifikasi kurangnya pemahaman terhadap emosi publik terkait wabah monkeypox melalui platform Twitter. Sebagian besar studi sebelumnya berfokus pada epidemiologi penyakit, namun belum banyak yang meneliti dampak emosional secara langsung. Selain itu, adanya ketidakseimbangan kelas dalam data emosi menjadi tantangan signifikan dalam klasifikasi emosi.
	Kontribusi	Studi ini memperkenalkan analisis klasifikasi emosi pada data Twitter terkait wabah monkeypox menggunakan pendekatan pembelajaran mendalam. Penelitian ini menawarkan perspektif baru dalam memahami dampak psikologis dari penyakit ini pada masyarakat umum. Selain itu, penelitian ini menggunakan data multibahasa, memberikan gambaran yang lebih luas dan mengurangi bias bahasa tertentu.
	Metode/ Solusi	Penelitian menggunakan teknik deep learning dengan model CNN, LSTM, BiLSTM, dan CNN-LSTM untuk klasifikasi emosi berdasarkan delapan kategori emosi Plutchik. Teknik SMOTE dan random undersampling diterapkan untuk menangani ketidakseimbangan data. NRCLex digunakan untuk labeling emosi, dan data diambil dari Twitter dengan bantuan API Tweepy.
	Batasan	Studi ini terbatas pada emosi publik terkait wabah monkeypox, sehingga mungkin tidak dapat digeneralisasikan ke penyakit lain tanpa adaptasi tambahan. Penggunaan dataset terbatas pada tweet terkait monkeypox,

		dan keterbatasan bahasa non-Inggris, meskipun diatasi sebagian dengan Google Translate, dapat memengaruhi akurasi klasifikasi emosi lintas bahasa.
7	Judul paper, Penulis	Improving Multi-label Classification Performance on Imbalanced Datasets Through SMOTE Technique and Data Augmentation Using IndoBERT Model (Cahya dkk., 2024)
	Permasalahan	Penelitian ini mengidentifikasi tantangan dalam klasifikasi emosi di media sosial berbahasa Indonesia, terutama terkait ketidakseimbangan data antar kategori emosi. Ketidakseimbangan ini menyebabkan model cenderung bias terhadap emosi mayoritas dan mengabaikan kelas minoritas.
	Kontribusi	Studi ini memberikan kontribusi dengan membandingkan efektivitas SMOTE dan data augmentasi untuk mengatasi ketidakseimbangan data pada model berbasis IndoBERT. Hasil penelitian menunjukkan bahwa SMOTE lebih unggul daripada augmentasi dalam meningkatkan akurasi dan F1-score pada data tidak seimbang.
	Metode/ Solusi	Penelitian ini menggunakan IndoBERT sebagai model utama untuk klasifikasi emosi. SMOTE diterapkan untuk mengatasi ketidakseimbangan data kelas minoritas dengan menghasilkan sampel sintetik. Model dievaluasi dengan metrik akurasi, presisi, recall, dan F1-score untuk membandingkan metode SMOTE dan augmentasi.
	Batasan	Studi ini terbatas pada data berbahasa Indonesia dari satu platform media sosial. Hasil mungkin tidak dapat digeneralisasi ke platform atau bahasa lain. Selain itu, penggunaan SMOTE dan IndoBERT tidak mencakup pendekatan berbasis deep learning lainnya yang mungkin juga efektif.
8	Judul paper, Penulis	EMOTION RECOGNITION INDONESIAN LANGUAGE FROM TWITTER, (William dkk., 2024)
	Permasalahan	Penelitian ini mengidentifikasi kurangnya penelitian klasifikasi emosi dalam bahasa Indonesia dibandingkan bahasa lain seperti Inggris, terutama dalam domain media sosial seperti Twitter. Salah satu masalah utama adalah keterbatasan dataset publik untuk klasifikasi emosi dalam bahasa Indonesia, yang menyulitkan pengembangan model NLP yang akurat untuk emosi.
	Kontribusi	Studi ini berkontribusi dengan membuat dataset emosi bahasa Indonesia dari Twitter yang berisi 7.629 tweet, menggunakan enam emosi dasar Ekman (Bahagia, Takut, Marah, Jijik, Sedih, Kaget) serta satu label netral. Selain itu, penelitian ini membandingkan model IndoBERT yang di-

		fine-tune dengan model Bi-LSTM untuk klasifikasi emosi, menunjukkan keunggulan IndoBERT.
	Metode/ Solusi	Data dikumpulkan menggunakan API Tweepy dengan filter kata kunci terkait emosi, diikuti oleh proses pelabelan manual oleh dua anotator dan pengecekan inter-rater reliability (kappa score 87.7%). Model yang dibandingkan adalah IndoBERT yang di-fine-tune dan Bi-LSTM. IndoBERT dioptimalkan dengan hyperparameter seperti ukuran batch dan learning rate, sementara Bi-LSTM menggunakan FastText untuk embedding.
	Batasan	Studi ini terbatas pada data berbahasa Indonesia dari Twitter, sehingga hasilnya mungkin tidak dapat digeneralisasikan ke platform lain. Selain itu, penelitian ini hanya mencakup satu model deep learning tambahan (Bi-LSTM) dan tidak membandingkan dengan model transformer lain yang mungkin memiliki performa kompetitif pada tugas klasifikasi emosi.
9	Judul paper, Penulis	ANALISIS EMOSI PADA MEDIA SOSIAL TWITTER MENGGUNAKAN METODE MULTINOMIAL NAÏVE BAYES DAN SYNTHETIC MINORITY OVERSAMPLING TECHNIQUE, (Agung dkk., 2023)
	Permasalahan	Penelitian ini mengidentifikasi masalah utama dalam klasifikasi emosi pada Twitter, yaitu ketidakseimbangan data antar kelas emosi, yang menyebabkan bias model terhadap kelas mayoritas dan menurunkan performa klasifikasi pada kelas minoritas. Selain itu, kualitas data Twitter, yang sering mengandung slang, noise, dan kesalahan ejaan, menjadi tantangan dalam analisis teks.
	Kontribusi	Penelitian ini menggabungkan Multinomial Naïve Bayes (MNB) dengan SMOTE untuk menangani ketidakseimbangan data. SMOTE digunakan untuk memperluas data kelas minoritas secara sintetik, sedangkan MNB digunakan sebagai algoritma klasifikasi utama. Studi ini juga menguji pengaruh parameter k pada SMOTE terhadap performa klasifikasi dan menyimpulkan bahwa k tidak signifikan memengaruhi hasil.
	Metode/ Solusi	Penelitian menggunakan pipeline: preprocessing (case folding, data cleaning, konversi slangword, tokenization, dll.), ekstraksi fitur n-gram, pembobotan dengan term frequency (TF), SMOTE untuk penyeimbangan data, dan klasifikasi menggunakan MNB. Evaluasi dilakukan menggunakan 10-fold cross-validation dengan metrik akurasi dan F1-score. Hasil menunjukkan bahwa SMOTE meningkatkan F1-score sebesar 1%.

	Batasan	Performa terbaik (F1-score rata-rata 66%) masih rendah untuk analisis emosi, menunjukkan keterbatasan Multinomial Naïve Bayes dibanding model deep learning modern. Noise data, seperti slang dan kesalahan ejaan, juga memengaruhi hasil preprocessing. Selain itu, dataset kecil (4401 tweet) mungkin kurang representatif untuk generalisasi.
10	Judul paper, Penulis	Comparative Analyses of BERT, RoBERTa, DistilBERT, and XLNet for Text-based Emotion, (Adoma dkk., 2020)
	Permasalahan	Penelitian ini mengidentifikasi tantangan dalam memilih model transformer terbaik untuk klasifikasi emosi berbasis teks, serta mengevaluasi performa masing-masing model dengan dataset seimbang seperti ISEAR.
	Kontribusi	Penelitian ini membandingkan empat model transformer (BERT, RoBERTa, DistilBERT, dan XLNet) untuk tugas klasifikasi emosi berbasis teks. Penelitian juga menyediakan benchmark kinerja model pada dataset ISEAR menggunakan metrik akurasi, presisi, recall, dan F1-score.
	Metode/ Solusi	Dataset ISEAR digunakan dengan preprocessing standar seperti penghapusan karakter khusus dan padding teks hingga panjang 200 token. Model yang digunakan adalah BERT, RoBERTa, DistilBERT, dan XLNet, yang dilatih dengan batch size 16, learning rate 4×10^{-5} , dan optimizer Adam selama 10 epoch. Evaluasi dilakukan menggunakan akurasi, presisi, recall, dan F1-score.
	Batasan	Penelitian ini hanya menggunakan dataset ISEAR, sehingga hasilnya mungkin tidak berlaku untuk dataset lain. Tidak mencakup model transformer dan tidak mengeksplorasi teknik optimasi yang lain.
11	Judul paper, Penulis	Performance comparison of tf-idf and word2vec models for emotion text classification, (Cahyani & Patasik, 2021)
	Permasalahan	Penelitian ini mengidentifikasi ketidakseimbangan data emosi pada dataset Twitter yang menyebabkan kesulitan dalam mengenali emosi minoritas seperti sad, scared, dan surprised. Selain itu, belum ada penelitian yang membandingkan kinerja TF-IDF dan Word2Vec secara langsung dalam klasifikasi emosi pada dataset kecil dan tidak seimbang.
	Kontribusi	Penelitian ini membandingkan performa TF-IDF dan Word2Vec sebagai teknik ekstraksi fitur untuk klasifikasi emosi teks, menunjukkan bahwa TF-IDF memberikan kinerja lebih baik pada dataset kecil dan tidak seimbang.

		Selain itu, studi ini memberikan evaluasi kinerja berbasis metrik seperti precision, recall, dan F1-measure.
	Metode/ Solusi	Penelitian ini menggunakan model Support Vector Machine (SVM) dan Multinomial Naïve Bayes (MNB) untuk klasifikasi, dengan ekstraksi fitur berbasis TF-IDF dan Word2Vec. Data diambil dari tweet Commuter Line dan Transjakarta dalam bahasa Indonesia, diklasifikasikan menjadi emotion dan no emotion di tahap pertama, lalu diklasifikasikan ke lima emosi (happy, angry, sad, scared, surprised) di tahap kedua. Evaluasi dilakukan dengan accuracy, precision, recall, dan F1-measure.
	Batasan	Penelitian ini terbatas pada penggunaan metode tradisional seperti TF-IDF dan Word2Vec, sehingga belum mengadopsi model mutakhir seperti BERT. Selain itu, dataset yang digunakan tidak seimbang dan memiliki jumlah data yang kecil untuk emosi minoritas seperti scared dan surprised.
12	Judul paper, Penulis	An Ensemble Technique to Classify Multi-Class Textual Emotion. (Parvin & Hoque, 2021)
	Permasalahan	Penelitian ini menghadapi tantangan dalam mengembangkan sistem klasifikasi emosi otomatis untuk bahasa Bengali, yang merupakan bahasa dengan sumber daya rendah. Tantangan lainnya termasuk ketidakseimbangan dataset dan kompleksitas linguistik Bengali.
	Kontribusi	Penelitian ini mengembangkan dataset emosi Bengali yang terdiri dari 8047 teks. Model ensemble berbasis Logistic Regression (LR), Random Forest (RF), dan Support Vector Machine (SVM) juga dikembangkan untuk meningkatkan akurasi klasifikasi emosi multi-kelas.
	Metode/ Solusi	Data diproses dengan preprocessing (menghapus tanda baca, karakter khusus, dll.). Fitur diekstraksi menggunakan Bag-of-Words (BoW) dan TF-IDF. Model seperti LR, RF, SVM, Decision Tree, dan Multinomial Naïve Bayes dibandingkan, dengan model ensemble memberikan hasil terbaik.
	Batasan	Penelitian ini menggunakan metode tradisional berbasis pembelajaran mesin dan belum mengadopsi model berbasis transformer seperti BERT. Distribusi kelas yang tidak merata membatasi performa.
13	Judul paper, Penulis	Klasifikasi Emosi Pada Pengguna Twitter Menggunakan Metode Klasifikasi Decision Tree, (Nurfauzan & Maharani, 2021)
	Permasalahan	Penelitian ini menghadapi tantangan dalam klasifikasi emosi berbasis teks Twitter, dengan keterbatasan dataset berbahasa

		Indonesia dan penggunaan algoritma tradisional untuk tugas multi-kelas emosi.
	Kontribusi	Penelitian ini mengimplementasikan algoritma <i>Decision Tree</i> dengan pembobotan TF-IDF untuk klasifikasi emosi, menggunakan dataset Twitter berbahasa Indonesia yang terdiri dari 4.401 data dengan 5 kelas emosi (<i>anger, happy, sadness, fear, love</i>).
	Metode/ Solusi	Dataset Twitter dengan 4.401 data diproses menggunakan TF-IDF untuk pembobotan kata. Algoritma Decision Tree digunakan untuk klasifikasi emosi, dengan evaluasi performa melalui metrik akurasi, precision, recall, dan F1-score.
	Batasan	Dataset yang relatif kecil membatasi kemampuan model untuk menangkap pola kompleks, sementara algoritma Decision Tree adalah metode tradisional yang kurang optimal dibandingkan model mutakhir seperti BERT atau IndoBERT.
14	Judul paper, Penulis	Emotion and sentiment analysis of tweets using BERT, (Chiorrini dkk., 2021)
	Permasalahan	Penelitian ini menghadapi tantangan dalam klasifikasi emosi berbasis teks Twitter, dengan keterbatasan dataset berbahasa Indonesia dan penggunaan algoritma tradisional untuk tugas multi-kelas emosi.
	Kontribusi	Penelitian ini mengimplementasikan algoritma <i>Decision Tree</i> dengan pembobotan TF-IDF untuk klasifikasi emosi, menggunakan dataset Twitter berbahasa Indonesia yang terdiri dari 4.401 data dengan 5 kelas emosi (<i>anger, happy, sadness, fear, love</i>).
	Metode/ Solusi	Dataset Twitter dengan 4.401 data diproses menggunakan TF-IDF untuk pembobotan kata. Algoritma Decision Tree digunakan untuk klasifikasi emosi, dengan evaluasi performa melalui metrik akurasi, precision, recall, dan F1-score.
	Batasan	Algoritma Decision Tree adalah metode tradisional yang kurang optimal dibandingkan model mutakhir seperti BERT atau IndoBERT.
15	Judul paper, Penulis	Studi Perbandingan pada Metode CNN-LSTM dan LSTM dalam Mendeteksi Emosi pada Data Teks Berbahasa Indonesia pada Media Sosial Twitter (Kustiwa dkk., 2024)
	Permasalahan	Permasalahan utamanya adalah menguji klaim bahwa model CNN-LSTM lebih unggul dari LSTM, karena klaim tersebut

		belum divalidasi secara statistik pada dataset Twitter berbahasa Indonesia yang memiliki karakteristik unik.
	Kontribusi	Tidak ada perbedaan kinerja yang signifikan antara CNN-LSTM dan LSTM pada metrik akurasi, F1-score, precision, dan recall. Satu-satunya keunggulan signifikan CNN-LSTM hanya pada nilai <i>loss</i> yang lebih rendah.
	Metode/ Solusi	membandingkan model LSTM dan CNN-LSTM pada dataset Twitter berbahasa Indonesia. Kinerja kedua model dievaluasi secara ketat menggunakan 30-fold cross-validation, dan hasilnya dianalisis dengan Uji T Berpasangan (Paired t-test) untuk menentukan signifikansi statistik dari perbedaan performa.
	Batasan	cakupannya yang hanya terbatas pada satu dataset Twitter berbahasa Indonesia. Selain itu, proses optimasi hyperparameter dilakukan secara manual dan melibatkan data uji karena keterbatasan sumber daya.
16	Judul paper, Penulis	Klasifikasi Emosi pada Teks Berbahasa Inggris Menggunakan Pendekatan Ensemble Bagging (Erfian Junianto dkk., 2024)
	Permasalahan	Penelitian ini berangkat dari tantangan dalam mengklasifikasikan emosi pada teks berbahasa Inggris, khususnya data media sosial yang tidak terstruktur. Algoritma umum seperti Naïve Bayes, Logistic Regression, dan KNN memiliki berbagai kelemahan, seperti asumsi independensi fitur yang tidak realistis, sensitivitas terhadap outlier, serta risiko overfitting. Selain itu, penelitian terdahulu menunjukkan akurasi yang bervariasi dan cenderung rendah, sehingga diperlukan pendekatan yang mampu meningkatkan stabilitas dan akurasi klasifikasi emosi.
	Kontribusi	Penelitian ini berkontribusi dengan mengusulkan penggunaan metode ensemble bagging yang menggabungkan tiga algoritma dasar (NB, LR, dan KNN) untuk meningkatkan performa klasifikasi emosi. Penelitian juga memberikan pemodelan emosi untuk tujuh kategori emosi, serta menyediakan evaluasi komprehensif yang menunjukkan peningkatan akurasi dibandingkan model individual. Selain itu, hasilnya membuka peluang untuk eksplorasi lebih lanjut terkait teknik ensemble lain guna meningkatkan keandalan klasifikasi.
	Metode/ Solusi	Metode yang digunakan mencakup pengumpulan data teks dari berbagai sumber, dilanjutkan dengan pembersihan dan preprocessing menggunakan teknik text mining. Tiga

		algoritma pembelajaran mesin—Naïve Bayes, Logistic Regression, dan KNN—diterapkan sebagai model dasar yang kemudian digabungkan menggunakan teknik ensemble bagging. Kinerja model dievaluasi menggunakan akurasi, precision, recall, dan f1-score, dengan perbandingan antara model individual dan ensemble untuk menilai efektivitas pendekatan.
	Batasan	Penelitian ini memiliki beberapa batasan, seperti penggunaan dataset yang hanya berfokus pada teks berbahasa Inggris sehingga kurang dapat digeneralisasikan ke bahasa lain. Model yang digunakan juga terbatas pada tiga algoritma dasar dan belum melibatkan metode yang lebih kompleks seperti deep learning. Selain itu, data media sosial yang digunakan mengandung banyak noise, dan pendekatan ensemble bagging memerlukan sumber daya komputasi yang lebih besar. Penelitian juga belum menguji model pada dataset yang lebih luas atau domain yang berbeda.
17	Judul paper, Penulis	Improving Multi-Label Emotion Classification on Imbalanced Social Media Data With BERT and Clipped Asymmetric Loss (Ramakrishnan & Dhinesh Babu, 2025)
	Permasalahan	Penelitian ini membahas kompleksitas multi-label emotion classification yang diperparah oleh ketidakseimbangan data, di mana model cenderung bias terhadap emosi yang sering muncul. Loss function tradisional seperti cross-entropy tidak mampu menangani kondisi ini dengan baik sehingga performa pada kelas minoritas menjadi rendah.
	Kontribusi	Penelitian ini mengusulkan Clipped Asymmetric Loss yang diadaptasi ke domain NLP dan dikombinasikan dengan BERT untuk meningkatkan deteksi emosi minoritas serta meningkatkan nilai F1-score secara keseluruhan.
	Metode/ Solusi	Metode yang digunakan adalah fine-tuning BERT dengan Asymmetric Loss yang memiliki parameter berbeda untuk kelas positif dan negatif serta mekanisme clipping untuk menjaga stabilitas pelatihan.
	Batasan	Penelitian ini belum mempertimbangkan hubungan antar label secara eksplisit dan masih membutuhkan dataset besar serta memiliki kompleksitas model yang tinggi.
18	Judul paper, Penulis	AFL-BERT : Enhancing Minority Class Detection in Multi-Label Text Classification with Adaptive Focal Loss and BERT (Labd & Housni, 2025)
	Permasalahan	Permasalahan utama yang diangkat adalah ketidakseimbangan data serta keterbatasan focal loss standar

		yang menggunakan parameter tetap sehingga tidak mampu menyesuaikan dengan distribusi kelas yang beragam.
	Kontribusi	Kontribusi penelitian ini adalah pengembangan Adaptive Focal Loss yang mampu menyesuaikan parameter secara dinamis berdasarkan frekuensi kelas sehingga lebih efektif dalam mendeteksi kelas minoritas.
	Metode/ Solusi	Pendekatan yang digunakan adalah integrasi BERT dengan Adaptive Focal Loss di mana parameter gamma disesuaikan secara adaptif selama proses pelatihan.
	Batasan	Penelitian ini memiliki keterbatasan pada ruang lingkup evaluasi yang terbatas, membutuhkan proses tuning yang kompleks, serta berpotensi mengalami instabilitas selama pelatihan.
19	Judul paper, Penulis	Ustnlp16 at SemEval-2025 Task 9: Improving Model Performance through Imbalance Handling and Focal Loss (Cai dkk., 2025)
	Permasalahan	Penelitian ini menghadapi tantangan berupa data yang sangat tidak seimbang, teks yang pendek dan tidak terstruktur, serta adanya kategori yang saling tumpang tindih sehingga menyulitkan proses klasifikasi.
	Kontribusi	Kontribusi utama penelitian ini adalah penggunaan kombinasi teknik augmentasi data dan focal loss dalam sebuah pipeline terstruktur untuk meningkatkan performa klasifikasi terutama pada kelas minoritas.
	Metode/ Solusi	Metode yang digunakan melibatkan model transformer seperti BERT dan RoBERTa dengan tahapan preprocessing, oversampling, Easy Data Augmentation, dan penerapan focal loss dalam proses pelatihan.
	Batasan	Pendekatan ini berisiko menghasilkan noise dari proses augmentasi serta overfitting akibat oversampling, dan tidak menawarkan inovasi baru pada fungsi loss.
20	Judul paper, Penulis	F_MixBERT: Sentiment Analysis Model using Focal Loss for Imbalanced E-commerce Reviews (Pang dkk., 2024)
	Permasalahan	Penelitian ini mengangkat permasalahan tambahan berupa inkonsistensi label, keberadaan data tidak berlabel dalam jumlah besar, serta ketidakseimbangan distribusi kelas dalam dataset e-commerce.
	Kontribusi	Kontribusinya adalah pengembangan framework F_MixBERT yang menggabungkan semi-supervised learning melalui MixMatch dengan focal loss untuk memanfaatkan data berlabel dan tidak berlabel secara bersamaan.

	Metode/ Solusi	Metode yang digunakan adalah kombinasi BERT dengan teknik MixMatch untuk menghasilkan pseudo-label serta penerapan focal loss untuk menangani kelas minoritas.
	Batasan	Keterbatasan utama terletak pada risiko kesalahan dari pseudo-label, kompleksitas metode yang tinggi, serta ketergantungan pada kualitas data tidak berlabel.
21	Judul paper, Penulis	Fine-Tuning BERT for Automated News Classification (Salih dkk., 2025)
	Permasalahan	Penelitian ini membahas keterbatasan metode tradisional dalam memahami konteks teks serta kelemahan model sebelumnya dalam menangkap dependensi jangka panjang dalam teks.
	Kontribusi	Kontribusi penelitian ini adalah menunjukkan bahwa fine-tuning BERT mampu memberikan performa yang lebih baik dibandingkan metode klasik dalam klasifikasi teks.
	Metode/ Solusi	Metode yang digunakan berupa fine-tuning model BERT dengan preprocessing standar dan penggunaan cross-entropy loss pada dataset berita.
	Batasan	Penelitian ini tidak membahas masalah ketidakseimbangan data maupun multi-label classification serta tidak menawarkan inovasi metode baru.

Tabel 2.1 menyajikan rangkuman dari berbagai penelitian relevan yang telah dilakukan dalam bidang klasifikasi emosi berbasis teks. Selain studi yang berfokus langsung pada klasifikasi emosi, tinjauan ini juga sengaja menyertakan penelitian dari domain lain yang relevan secara metodologis. Setiap entri dalam tabel menguraikan informasi kunci yang mencakup judul penelitian, penulis, permasalahan yang diangkat, metode atau solusi yang diusulkan, kontribusi utama, serta batasan dari masing-masing studi.

Secara keseluruhan, mayoritas penelitian tersebut memiliki fokus yang sama, yaitu mengklasifikasikan emosi yang diekspresikan pengguna dalam data tekstual, khususnya dari platform media sosial seperti *Twitter*. Meskipun memiliki tujuan yang serupa, setiap penelitian menerapkan pendekatan yang beragam. Berbagai model telah dieksplorasi, mulai dari algoritma *machine learning* tradisional seperti *Multinomial Naïve Bayes* dan *Decision Tree*, hingga arsitektur *deep learning* yang lebih kompleks seperti *CNN*, *LSTM*, dan *BiLSTM*. Sebagian besar penelitian modern memanfaatkan model berbasis transformer untuk menangkap konteks kata secara mendalam, seperti *BERT*, *ROBERTa*, *XLNet*, dan model khusus Bahasa Indonesia, *IndoBERT*.

Selain itu, untuk mengatasi tantangan umum seperti ketidakseimbangan kelas data, beberapa penelitian menerapkan teknik-teknik khusus. Studi seperti penelitian *cost-sensitive* pada BERT (Madabushi dkk., 2020) dan penerapan *Focal Loss* pada model *transformer* untuk ulasan konsumen (Li & Shimada, 2024) diikutsertakan karena, meskipun tidak menganalisis emosi, keduanya menjadi acuan penting dalam penerapan cost-sensitive learning dan *Focal Loss*.

Teknik lain yang juga umum digunakan adalah *data sampling* dengan *SMOTE* dan *random undersampling*. Dalam hal representasi fitur, rentang metode yang digunakan juga bervariasi, mulai dari pendekatan tradisional seperti *Bag-of-Words* (*BoW*), *TF-IDF*, dan *Word2Vec*, hingga penggunaan *embedding* kontekstual yang dihasilkan oleh model-model *transformer*.

Terakhir, tabel ini juga secara sistematis menyajikan batasan dari setiap studi yang dirangkum. Pemaparan batasan ini bertujuan untuk mengidentifikasi celah atau kesenjangan (*research gap*) dalam literatur yang ada, baik secara topikal maupun metodologis, sehingga memvalidasi posisi dan kontribusi penelitian yang akan dilakukan.

2.3. Matriks Penelitian

Berikut adalah matriks penelitian

Tabel 2. 2 Matriks Penelitian

Penulis	Model	Teknik Imbalanced	Karakteristik Dataset	Hasil
(Madabushi dkk., 2020)	BERT	<i>Cost-sensitive learning</i>	Propaganda (PTC), 2 kelas	Performa tinggi, peringkat 2 pada task propaganda
(Olusegun dkk., 2023)	CNN, LSTM, BiLSTM, CNN-LSTM	SMOTE, <i>Random Undersampling</i>	Twitter (Multibahasa), 8 emosi	Akurasi tertinggi 96% (CNN)
(Chiorrini dkk., 2021)	BERT	Undersampling	Twitter (Inggris), 4 emosi	Akurasi 0.92 (sentiment), 0.90 (emotion)
(Adoma dkk., 2020)	BERT, RoBERTa, XLNet	-	Twitter (ISEAR), 7 emosi	Akurasi terbaik 0.7431 (RoBERTa)

Penulis	Model	Teknik Imbalanced	Karakteristik Dataset	Hasil
(Parvin & Hoque, 2021)	Ensemble (LR, RF, SVM)	-	Media Sosial (Bengali), 6 emosi	F1-score 62.39%
(Cahyani & Patasik, 2021)	SVM, Multinomial Naïve Bayes	-	Twitter (Indonesia), 5 emosi	SVM + TF-IDF terbaik
(Nurfauzan & Maharani, 2021)	Decision Tree	-	Twitter (Indonesia), 5 emosi	Akurasi 55%
(Cahyadkk., 2024)	IndoBERT	SMOTE, Augmentasi	Ulasan Aplikasi (Indonesia), 6 emosi	Akurasi 82%, F1 86%
(Agung dkk., 2023)	Multinomial Naïve Bayes	SMOTE	Twitter (Indonesia), 5 emosi	Akurasi 65%
(William dkk., 2024)	IndoBERT, Bi-LSTM	-	Twitter (Indonesia), 7 emosi	Akurasi 92.3%
(Li & Shimada, 2024)	BERT, SVM. BERT-SVM	<i>Focal Loss</i>	Ulasan (Jepang), 3 kelas	Minor class meningkat
(Erfian Junianto dkk., 2024)	Ensemble	-	Framenet(inggris), 7 emosi	Akurasi hingga 98%
(Pang dkk., 2024)	BERT	Focal Loss, MixMatch	E-commerce reviews, unlabeled, imbalanced	Akurasi meningkat dibanding baseline

Penulis	Model	Teknik Imbalanced	Karakteristik Dataset	Hasil
(Ramakrishnan & Dhinesh Babu, 2025)	BERT	Clipped Asymmetric Loss	GoEmotions, SemEval, multi-label, imbalanced	Macro F1 meningkat (0.46 menjadi 0.54)
(Salih dkk., 2025)	BERT	-	Reuters news, multi-class	Akurasi 91.77%
(Cai dkk., 2025)	BERT, RoBERTa	Oversampling, EDA, Focal Loss	Food hazard text, pendek, tidak terstruktur, imbalanced	Peningkatan akurasi dan F1-score
(Labd & Housni, 2025)	BERT	Adaptive Focal Loss	CMU Movie Summary, multi-label, imbalanced	F1-score 0.5, ROC 0.79
(Masaling dkk., 2024)	IndoBERT	-	Twitter (Indonesia), 6 emosi	acc (74%), prec (79%), rec (78%), f1 (78%)
(Kustiwa dkk., 2024)	CNN, LSTM	-	Twitter (Indonesia), 6 emosi	acc (71%), prec (72%), rec (72%), f1 (72%)
(Baihaqi dkk., 2023)	CNN	Class Weighting (ICF)	Twitter (Indonesia), 6 emosi	acc (72%), prec (72%), rec (72%), f1 (72%)
(Ours, 2025)	<i>IndoBERT Tweet</i>	<i>Focal Loss</i>	Twitter (Indonesia), 6 emosi	-

Tabel 2.2 menyajikan matriks perbandingan dari berbagai penelitian relevan yang menjadi landasan bagi studi ini. Tabel tersebut merangkum informasi teknis utama dari setiap penelitian, mencakup model yang digunakan, teknik penanganan data tidak seimbang (*imbalance*) yang diterapkan, karakteristik dataset, serta hasil

kinerja yang dilaporkan. Secara keseluruhan, tinjauan ini memetakan lanskap penelitian klasifikasi emosi yang sangat beragam, baik dari segi metodologi maupun konteks data.

Dari matriks tersebut, terlihat jelas bahwa terdapat rentang model yang digunakan, mulai dari pendekatan machine learning klasik seperti *Naïve Bayes* dan *SVM*, arsitektur *deep learning* seperti *CNN* dan *LSTM*, hingga model berbasis *transformer* seperti *BERT* dan variasinya untuk bahasa spesifik seperti *IndoBERT*. Untuk mengatasi tantangan utama berupa ketidakseimbangan kelas, berbagai strategi telah dieksplorasi. Beberapa penelitian menerapkan teknik pada level data seperti *SMOTE* dan *undersampling*, sementara studi lain yang relevan secara metodologis (meskipun pada domain berbeda) menggunakan pendekatan pada level algoritma dengan memodifikasi loss function, seperti *Cost-Sensitive Learning* dan *Focal Loss*.

Berdasarkan analisis terhadap tabel tersebut, terlihat bahwa kinerja pada dataset *benchmark* dalam penelitian-penelitian sebelumnya (Baihaqi dkk., 2023; Kustiwa dkk., 2024; Masaling dkk., 2024) masih memiliki ruang untuk optimasi. Oleh karena itu, tujuan dari penelitian ini adalah untuk meningkatkan performa klasifikasi emosi pada *benchmark* dataset tersebut dengan mengusulkan kombinasi model dan teknik yang lebih spesifik. Untuk mencapai tujuan ini, model *IndoBERTweet* dipilih karena arsitekturnya secara khusus telah dilatih pada data Twitter berbahasa Indonesia, sehingga secara teoretis lebih unggul dalam memahami konteks dan nuansa data yang digunakan dibandingkan *IndoBERT* generik. Selanjutnya, *Focal Loss* dipilih sebagai teknik penanganan *imbalance*

karena kemampuannya untuk secara langsung memodifikasi proses pembelajaran model agar lebih fokus pada sampel-sampel dari kelas minoritas yang sulit, sebuah pendekatan level algoritma yang menjadi alternatif dari metode level data seperti *SMOTE*.