

# Analisis Tingkat Stres Pengguna Twitter pada Menfess “Tanyarl” Menggunakan IndoBERT

Nadia Nur Alifiana\*, Bambang Irawan, Otong Saeful Bachri

Program Studi Teknik Informatika, Fakultas Teknik, Universitas Muhadi Setiabudi, Brebes, Indonesia

Email: <sup>1,\*</sup>nadialifiana@gmail.com, <sup>2</sup>bambangumus@gmail.com, <sup>3</sup>otongsaifulbahri@gmail.com

Email Penulis Korespondensi: nadialifiana@gmail.com\*

Submitted: 12/01/2026; Accepted: 07/02/2026; Published: 31/03/2026

**Abstrak**—Twitter merupakan salah satu media sosial berbasis teks yang banyak dimanfaatkan pengguna sebagai ruang untuk mengekspresikan kondisi emosional dan tekanan psikologis secara terbuka, khususnya melalui akun *menfess* yang bersifat anonim. Karakter anonim tersebut mendorong pengguna untuk menyampaikan keluhan, perasaan tertekan, dan stres yang dialami dalam kehidupan sehari-hari. Namun, tingginya volume dan variasi bahasa informal dan kontekstual pada *menfess* menyebabkan proses identifikasi tingkat stres secara manual menjadi tidak efisien dan rentan terhadap subjektivitas. Penelitian sebelumnya umumnya masih terbatas pada klasifikasi biner atau menggunakan metode konvensional yang kurang mampu menangkap konteks semantik bahasa Indonesia secara optimal. Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk mengembangkan sistem deteksi tingkat stres pengguna Twitter pada *menfess* “Tanyarl” dengan memanfaatkan pendekatan kecerdasan buatan melalui *fine-tuning* model IndoBERT. Data unggahan (tweet) dikumpulkan melalui proses *crawling* menggunakan *Harvest-Tweet*, kemudian diproses melalui tahapan *preprocessing* teks dan pelabelan semi-otomatis ke dalam tiga kategori tingkat stres, yaitu stres, tidak stres, dan netral. Dataset yang telah diproses selanjutnya digunakan untuk melatih model IndoBERT dengan pendekatan klasifikasi multi-kelas. Hasil evaluasi menunjukkan bahwa model IndoBERT mencapai nilai akurasi sebesar 84% serta *macro F1-score* sebesar 85%, yang menandakan performa model yang konsisten pada seluruh kelas stres. Hasil ini membuktikan bahwa IndoBERT efektif dalam memahami konteks linguistik dan ekspresi emosional pada teks media sosial berbahasa Indonesia. Kontribusi penelitian ini meliputi penerapan model *transformer* untuk mendeteksi tingkat stres berbasis *menfess* anonim, pengembangan klasifikasi multi-kelas stres, serta penyediaan bukti empiris efektivitas IndoBERT dalam analisis kesehatan mental berbasis data media sosial.

**Kata Kunci:** Deteksi Stres; Media Sosial; Twitter; *Menfess*; IndoBERT.

**Abstract**—Twitter is a text-based social media platform widely used by users to express emotional states and psychological distress, particularly through anonymous *menfess* accounts. The anonymous nature of *menfess* encourages users to share personal complaints and stress-related experiences. However, the large volume of data and the prevalence of informal and context-dependent language make manual stress identification inefficient and prone to subjectivity. Previous studies have often been limited to binary classification or conventional methods that are less effective in capturing the semantic complexity of Indonesian social media text. This study aims to develop an automated stress level detection system for Twitter users on the *menfess* account “Tanyarl” using an artificial intelligence approach based on fine-tuning the IndoBERT model. Tweet data were collected through a crawling process using *Harvest-Tweet*, followed by text preprocessing and semi-automatic labeling into three stress level categories: low, moderate, and high stress. The processed dataset was then used to train an IndoBERT-based multi-class classification model. Experimental results show that the proposed model achieved an accuracy of 84% and a *macro F1-score* of 85%, indicating stable and balanced performance across stress categories. These findings demonstrate the effectiveness of IndoBERT in capturing linguistic context and emotional expressions in Indonesian social media text. This research contributes to the application of transformer-based models for multi-class stress detection on anonymous social media data and provides empirical evidence of IndoBERT’s suitability for mental health analysis in Indonesian-language social media contexts.

**Keywords:** Stress Detection; Social Media; Twitter; *Menfess*; IndoBERT

## 1. PENDAHULUAN

Media sosial saat ini menjadi salah satu ruang utama bagi masyarakat untuk mengekspresikan emosi dan kondisi psikologis secara terbuka. Indonesia memiliki sebanyak 191 juta pengguna aktif media sosial pada Januari 2022, mengalami peningkatan sebesar 12,35% dibandingkan tahun sebelumnya, dengan Twitter sebagai salah satu platform populer yang digunakan oleh sekitar 18,45 juta pengguna pada tahun yang sama [1]. Platform Twitter, atau yang sekarang dikenal sebagai X merupakan media sosial berbasis teks yang kerap digunakan untuk menyalurkan keluhan, tekanan akademik, maupun beban pekerjaan. Selain itu, Twitter juga sering dimanfaatkan sebagai tempat meluapkan permasalahan pribadi, terutama ketika pengguna merasa stres atau mengalami kelelahan akibat aktivitas sehari-hari [2].

Fenomena tersebut berkaitan erat dengan stres sebagai respons mental maupun fisik yang muncul dari proses penilaian kognitif individu terhadap tuntutan lingkungan [3]. Dalam konteks komunikasi digital, stres kerap terefleksikan melalui ekspresi keluhan, pengungkapan emosi, serta narasi mengenai tekanan akademik, sosial, dan personal yang disampaikan pengguna di media sosial [4]. Individu yang mengalami stres cenderung mengalami tekanan emosional yang tinggi serta kesulitan dalam mengelola emosi, yang dapat dipicu oleh faktor internal maupun eksternal, seperti tekanan akademik, beban pekerjaan, kondisi keluarga, dan lingkungan sosial. Apabila

stres berlangsung dalam jangka waktu panjang dan tidak ditangani secara tepat, kondisi ini berpotensi berdampak serius terhadap kesehatan fisik dan mental, serta meningkatkan risiko terjadinya gangguan psikologis lainnya [5].

Lingkungan media sosial yang memberikan ruang ekspresi terbuka turut mendorong berkembangnya berbagai bentuk komunikasi anonim, salah satunya melalui akun *menfess* (*mention confess*). Akun *menfess* memungkinkan pengguna mengirimkan unggahan secara anonim dan dipublikasikan kepada publik, sehingga mendorong keterbukaan dalam menyampaikan perasaan, keluhan pribadi, serta tekanan emosional. Anonimitas dalam komunikasi daring diketahui dapat mengurangi hambatan sosial dan meningkatkan kecenderungan individu untuk mengekspresikan kondisi emosional yang bersifat personal dan sensitif, termasuk stres dan tekanan psikologis [6]. Menurut penelitian oleh Basuki dan Indrabayu [7], *menfess* berfungsi sebagai bentuk katarsis digital yang memungkinkan pengguna mengekspresikan keresahan emosional sekaligus memperoleh respons dan validasi dari komunitas daring. Oleh karena itu, unggahan pada akun *menfess* tidak hanya merefleksikan interaksi sosial pengguna media, tetapi juga merepresentasikan kondisi tekanan emosional masyarakat moder dalam ruang virtual.

Salah satu akun *menfess* yang aktif dan memiliki tingkat interaksi tinggi adalah “Tanyarl”. Akun ini banyak dimanfaatkan sebagai ruang untuk berbagi pengalaman, keluhan, serta permasalahan kehidupan sehari-hari oleh pengguna dari berbagai latar belakang. Karakteristik unggahan yang anonim, spontan, serta sarat dengan ekspresi emosional menjadikan unggahan pada akun *menfess* “Tanyarl” sebagai sumber data yang relevan untuk mengkaji representasi stres masyarakat Indonesia dalam bentuk teks media sosial. Namun demikian, tingginya volume data, variasi bahasa informal, penggunaan slang, serta konteks emosional yang kompleks menimbulkan permasalahan utama dalam studi kasus ini, yaitu sulitnya melakukan identifikasi dan pengelompokan tingkat stres secara manual maupun dengan metode konvensional secara akurat dan konsisten.

Permasalahan tersebut menjadi semakin signifikan mengingat tingginya prevalensi gangguan kesehatan mental di Indonesia. Survei I-NAMHS tahun 2022 mencatat 34,9% remaja dalam 12 bulan terakhir mengalami masalah kesehatan mental dan 2,45 juta atau 5,5% remaja memiliki setidaknya salah satu gangguan mental [8]. Tahun berikutnya, Survei Kesehatan Indonesia pada 2023 menunjukkan remaja usia 15–24 tahun menjadi kelompok dengan gejala depresi tertinggi. Sebanyak 1% mengalami depresi, 3,7% mengalami gangguan kecemasan, 0,9% menunjukkan gejala PTSD, dan 0,5% mengalami ADHD [9]. Kondisi stres yang tidak tertangani berpotensi berkembang menjadi gangguan kesehatan mental yang lebih serius, salah satunya adalah depresi, yang prevalensinya di Indonesia masih tergolong tinggi. Menurut laporan Survei Kesehatan Indonesia (SKI) oleh Kementerian Kesehatan, prevalensi depresi nasional mencapai 1,4% pada tahun 2023, dengan angka tertinggi ditemukan pada kelompok usia 15–24 tahun (Generasi Z) sebesar 2%, diikuti lansia usia  $\geq 75$  tahun (1,9%), kelompok usia 65–74 tahun (1,6%), usia 25–34 tahun (1,3%), 55–64 tahun (1,2%), 45–54 tahun (1,1%), serta prevalensi terendah pada kelompok usia 35–44 tahun sebesar 1% [10]. Meskipun generasi muda menunjukkan prevalensi depresi tertinggi, hanya sekitar 10,4% yang tercatat mencari pengobatan. Kementerian Kesehatan menyatakan bahwa kondisi depresi yang tidak tertangani, khususnya pada Generasi Z, berpotensi memicu berbagai masalah sosial, seperti memburuknya kondisi kesehatan mental, meningkatnya risiko bunuh diri, penyalahgunaan zat adiktif, serta permasalahan sosial lainnya. Rendahnya tingkat deteksi dan penanganan dini, yang dipengaruhi oleh keterbatasan sumber daya serta stigma sosial terhadap isu kesehatan mental, menyebabkan banyak kasus stres berkembang menjadi gangguan yang lebih serius [11].

Sejumlah penelitian sebelumnya telah menunjukkan bahwa teks pada Twitter dapat dimanfaatkan untuk mendeteksi sinyal kesehatan mental melalui pendekatan pemrosesan bahasa alami [12]. Dalam konteks tersebut, teknik pemrosesan *Natural Language Processing* (NLP) berbasis *deep learning*, khususnya model Transformer seperti BERT, banyak digunakan dalam penelitian analisis teks. BERT (*Bidirectional Encoder Representations from Transformers*) merupakan model bahasa berbasis *transformer* yang dilatih secara dua arah (*bidirectional*) sehingga mampu memahami konteks kata berdasarkan keseluruhan kalimat, bukan hanya dari urutan sebelumnya, yang menjadikannya efektif untuk berbagai tugas pemahaman bahasa alami [13]. Selain itu, proses pelatihan BERT melalui teknik *bidirectional pre-training* pada teks tidak berlabel memungkinkan model ini menangkap hubungan semantik dan konteks lebih mendalam dibandingkan model sekuensial seperti RNN dan LSTM yang hanya memproses teks dalam satu arah [14]. Kemampuan tersebut memungkinkan BERT menangkap nuansa emosional dan makna implisit dalam teks media sosial, sehingga relevan untuk mendeteksi kondisi psikologis seperti stres melalui unggahan daring [15].

Dalam konteks bahasa Indonesia, *IndoBERT* hadir sebagai model bahasa yang telah dilatih menggunakan korpus teks yang luas dan beragam dari sumber-sumber berbahasa Indonesia, termasuk berita, media sosial, blog, situs web dan dokumen bahasa Indonesia lainnya sehingga lebih sesuai untuk memproses teks yang mengandung ragam bahasa informal dan slang [16]. Model ini dirancang untuk menangani berbagai tugas pemahaman bahasa alami seperti klasifikasi teks, analisis sentimen, dan ekstraksi informasi, dan telah menunjukkan kinerja yang unggul dibandingkan model-model sebelumnya [17]. Namun, sebagian besar penelitian terdahulu masih berfokus pada dataset berbahasa asing, tugas klasifikasi biner, atau analisis emosi dan sentimen secara umum, sementara penelitian yang secara khusus menerapkan *IndoBERT* untuk deteksi tingkat stres pada data *menfess* Twitter dengan pendekatan klasifikasi multi-kelas masih sangat terbatas [18].

Penelitian oleh J. Ria Wiyani [5] melakukan klasifikasi tingkat stres menggunakan metode *Support Vector Machine* dengan seleksi fitur *Information Gain* pada dataset berukuran kecil, dan memperoleh akurasi sebesar 59,11%. Hasil ini menunjukkan bahwa pendekatan *machine learning* konvensional masih memiliki keterbatasan

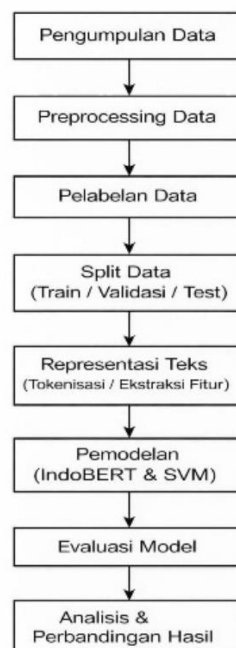
dalam menangkap kompleksitas bahasa alami. Sementara itu, penelitian oleh Suhartono, Saputra, Pratama, Nathaniel [19] menunjukkan bahwa model berbasis *Transformer* seperti BERT dan RoBERTa mampu mencapai performa yang lebih tinggi dalam mendeteksi stres pada data Twitter, dengan nilai akurasi dan *F1-score* di atas 0,8. Temuan tersebut memperkuat potensi penggunaan model BERT untuk analisis stres berbasis teks media sosial.

Berdasarkan celah penelitian tersebut, novelty penelitian ini terletak pada penerapan model *IndoBERT* yang di-*fine-tune* untuk melakukan klasifikasi multi-kelas tingkat stres pada data *menfess* Twitter berbahasa Indonesia, khususnya pada akun “Tanyar!” yang memiliki karakteristik bahasa anonim dan emosional. Penelitian ini bertujuan untuk mengembangkan sistem deteksi tingkat stres pengguna Twitter dengan mengklasifikasikan *tweet* ke dalam tiga kategori tingkat stres, yaitu stres, tidak stres, dan netral. Kontribusi penelitian ini mencakup pengembangan pendekatan deteksi stres berbasis *transformer* pada data *menfess* anonim, penyediaan evaluasi empiris performa *IndoBERT* dalam klasifikasi multi-kelas stres, serta dukungan awal terhadap upaya pemantauan dan deteksi dini kondisi stres melalui analisis teks media sosial berbahasa Indonesia.

## 2. METODOLOGI PENELITIAN

### 2.1 Tahapan Penelitian

Tahapan penelitian dalam studi ini meliputi pengumpulan data, *preprocessing* data, pelabelan data, pembagian data (*split* data), representasi teks, pemodelan, evaluasi model, serta analisis dan perbandingan hasil. Pada tahap representasi teks, digunakan tokenisasi *IndoBERT* untuk model *transformer* dan ekstraksi fitur TF-IDF untuk baseline SVM. Selanjutnya, proses pemodelan dilakukan menggunakan *IndoBERT* dan SVM sebagai model pembanding. Alur lengkap tahapan penelitian disajikan pada **Gambar 1**.



**Gambar 1.** Flowchart Alur Penelitian

Pengumpulan data merupakan tahapan paling awal dalam penelitian ini. Data dikumpulkan melalui proses *crawling tweet* dari akun *menfess* “Tanyar!” menggunakan *Harvest-Tweet*. Teks unggahan yang terkumpul kemudian diproses dalam tahap *preprocessing*, meliputi pembersihan teks, normalisasi simbol, penanganan emoji, serta normalisasi kata tidak baku untuk memastikan data siap digunakan dalam pemodelan.

Tahap berikutnya adalah pelabelan data secara semi-otomatis dengan memanfaatkan pencarian berbasis kata kunci untuk setiap kategori label. Teks unggahan terlebih dahulu dikelompokkan secara awal menggunakan daftar kata kunci terkait tingkat stres, tidak stres, dan netral. Hasil pengelompokan tersebut kemudian diverifikasi secara manual oleh anotator untuk memastikan ketepatan label. Setelah itu, data dibagi menjadi tiga subset melalui tahap *split* data, yaitu data *train*, validasi, dan *test*.

Data yang telah dipisahkan selanjutnya digunakan pada tahap pemodelan. Pada penelitian ini, model utama yang digunakan adalah *IndoBERT* melalui proses tokenisasi dan *fine-tuning*. Selain itu, model *Support Vector Machine* (SVM) juga diimplementasikan sebagai model baseline untuk keperluan perbandingan performa. Tahapan akhir dalam penelitian ini adalah evaluasi model menggunakan metrik klasifikasi untuk menilai kemampuan model dalam mendeteksi tingkat stres pada unggahan *menfess*.

### 2.1.1 Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan dengan memanfaatkan *Harvest-Tweet*, yaitu *tool* berbasis Node.js yang digunakan untuk melakukan proses *crawling* teks unggahan tanpa memerlukan API resmi Twitter. Penggunaan *tool* ini dipilih untuk memudahkan proses pengambilan data secara langsung dari akun *menfess* “Tanyar!” pada rentang waktu yang telah ditentukan. Data yang berhasil dikumpulkan tidak hanya berupa teks utama unggahan, tetapi juga dilengkapi dengan data pendukung, seperti waktu unggahan, dan jumlah interaksi.

Untuk menjaga privasi pengguna serta mematuhi prinsip etika penelitian, seluruh informasi yang berpotensi mengidentifikasi individu dihapus dari data yang digunakan. Dengan demikian, data yang dianalisis hanya berfokus pada konten teks dan informasi umum yang relevan dengan tujuan penelitian.

### 2.1.2 Preprocessing Data

Tahap *preprocessing* dilakukan sebagai langkah awal untuk membersihkan dan menstandarkan data teks sebelum digunakan dalam proses pelatihan model *IndoBERT*. Pada tahap ini, peneliti menghapus elemen-elemen yang tidak relevan, seperti URL, *mention*, dan *hashtag*, dengan memanfaatkan pola ekspresi reguler. Selain itu, karakter non-alfanumerik, termasuk emoji dan simbol khusus juga turut dihilangkan karena berpotensi menambah *noise* pada data teks. Untuk menjaga konsistensi format, spasi berlebih yang muncul dalam teks juga distandarkan menjadi satu spasi. Melalui tahapan *preprocessing* tersebut, seluruh data teks berada dalam kondisi yang lebih bersih dan seragam, sehingga siap untuk diproses lebih lanjut pada tahap pemodelan.

### 2.1.3 Pelabelan Semi-otomatis

Dalam penelitian ini, proses pelabelan data dilakukan dengan menggunakan pendekatan semi-otomatis. Pendekatan tersebut dipilih karena mampu mempercepat proses anotasi data tanpa menghilangkan peran penilaian manusia. Pada tahap awal, peneliti menyusun daftar kata kunci yang disesuaikan dengan masing-masing kategori label, yaitu stres, tidak stres, dan netral. Kata-kata seperti “capek”, “kesel”, “stres”, serta beberapa istilah lain yang sering digunakan untuk mengekspresikan emosi negatif dalam konteks keluhan sehari-hari, digunakan sebagai indikator awal untuk mendeteksi unggahan yang berpotensi mengandung stres. Sedangkan unggahan dengan ungkapan kegembiraan, rasa bersyukur, dan luapan rasa bahagia digunakan untuk kategori tidak stres. Sebaliknya, unggahan yang mengungkapkan kegembiraan, rasa syukur, maupun luapan emosi positif lainnya dikelompokkan ke dalam kategori tidak stres.

Meskipun demikian, tidak seluruh unggahan yang terdeteksi melalui kata kunci secara langsung ditetapkan sebagai data stres maupun tidak stres. Beberapa unggahan menunjukkan keluhan ringan atau mengandung makna yang ambigu, misalnya ketika satu kalimat memuat lebih dari satu ekspresi emosi. Dalam kasus seperti ini, penentuan label dilakukan manual oleh anotator. Unggahan dengan karakteristik tersebut kemudian diklasifikasikan ke dalam kategori netral, berdasarkan pertimbangan konteks dan makna keseluruhan kalimat. Hasil pelabelan awal ini tidak langsung digunakan sepenuhnya. Anotator melakukan pemeriksaan ulang secara manual untuk memastikan bahwa konteks dan makna unggahan telah sesuai dengan kategori yang diberikan. Apabila ditemukan ketidaksesuaian, label pada unggahan tersebut diperbaiki. Dengan cara ini, proses anotasi tetap terkontrol dan kualitas data yang digunakan dalam penelitian dapat terjaga.

### 2.1.4 Split Data

Dataset yang telah dilabeli kemudian dibagi menjadi tiga bagian, yaitu 70% sebagai data *train*, 15% sebagai data validasi, dan 15% sebagai data *test*. Data *training* digunakan untuk melatih model, data validasi untuk memantau performa serta mencegah terjadinya *overfitting*, sedangkan data *testing* berfungsi mengukur performa model pada data baru yang belum pernah diuji sebelumnya [20]. Pembagian data dilakukan dengan menjaga proporsi setiap kelas agar distribusi label tetap seimbang.

### 2.1.5 Tokenisasi *IndoBERT*

Proses tokenisasi dilakukan menggunakan tokenizer bawaan *IndoBERT*. Tokenizer bertugas mengonversi teks unggahan menjadi urutan token yang sesuai dengan format input model. Pada proses ini, *IndoBERT* secara otomatis menambahkan beberapa token khusus untuk menandai struktur input.

Token [CLS] ditempatkan pada awal setiap teks unggahan sebagai representasi keseluruhan teks. *Embedding* dari token ini digunakan oleh *IndoBERT* untuk melakukan klasifikasi tingkat stres menjadi tiga kategori, yaitu stres, tidak stres, dan netral. Token [SEP] kemudian ditambahkan di akhir teks unggahan sebagai penanda batas akhir kalimat, mengingat penelitian ini termasuk tugas *single-sentence classification* sehingga hanya membutuhkan satu token [SEP]. Token [PAD] digunakan untuk menyeragamkan panjang setiap tweet agar sesuai dengan panjang maksimum token yang ditetapkan. Token ini ditambahkan pada bagian akhir teks yang lebih pendek dan diabaikan oleh model melalui *attention mask* [21].

Pada penelitian ini, nilai maksimal panjang token yang digunakan adalah 160 token, sehingga teks unggahan yang lebih pendek akan ditambahkan token [PAD], sedangkan tweet yang lebih panjang akan dipotong agar sesuai dengan batas tersebut.

### 2.1.6 Fine-tuning IndoBERT

Tahapan selanjutnya adalah melakukan *fine-tuning* pada model *IndoBERT*. *Fine-tuning* merupakan proses penyesuaian model pelatihan agar lebih optimal dalam menyelesaikan tugas atau dataset tertentu [22]. Pada tahap ini, pemilihan dan pengaturan *hyperparameter* menjadi penting karena konfigurasi yang tepat dapat meningkatkan performa model secara signifikan. Untuk itu, dilakukan eksplorasi beberapa kombinasi *hyperparameter* guna memperoleh pengaturan yang paling sesuai. Berdasarkan penelitian terkait model *BERT*, terdapat tiga *hyperparameter* utama yang umum dioptimalkan pada proses *fine-tuning*, yaitu *batch size*, *learning rate*, dan jumlah epoch [23]. Model dasar yang digunakan adalah ‘indobenchmark/Indobert-base-pl’, yang telah dilatih sebelumnya pada korpus bahasa Indonesia.

### 2.1.7 Model Baseline Support Vector Machine (SVM)

Sebagai model pembanding, penelitian ini mengimplementasikan *Support Vector Machine* (SVM) dengan representasi fitur TF-IDF. Model SVM dilatih menggunakan dataset yang sama serta skema pembagian data yang identik dengan eksperimen *IndoBERT* untuk memastikan perbandingan performa yang adil. Kernel linear digunakan karena efektif untuk klasifikasi teks berdimensi tinggi dan bersifat sparse. Model SVM berfungsi sebagai baseline untuk mengevaluasi peningkatan performa yang diberikan oleh pendekatan berbasis transformer.

### 2.1.7 Evaluasi Model

Evaluasi model dilakukan menggunakan data *test* yang tidak pernah digunakan pada tahapan pelatihan maupun validasi. Metrik evaluasi yang digunakan meliputi akurasi, *recall*, *precision*, dan *F1-score*. Perhitungan masing-masing metrik dilakukan menggunakan persamaan (1)-(4) [24].

$$\text{Akurasi} = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (2)$$

$$\text{Precision} = \frac{TP}{TP+FP} \quad (3)$$

$$\text{F1-score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

Keterangan:

TP (*True Positive*) : Jumlah data positif yang diklasifikasikan dengan benar

FP (*False Positive*) : Jumlah data negatif yang salah diklasifikasikan sebagai positif

TN (*True Negative*) : Jumlah data negatif yang diklasifikasikan sebagai negatif

FN (*False Negative*) : Jumlah data positif yang salah diklasifikasikan sebagai negatif

## 3. HASIL DAN PEMBAHASAN

### 3.1 Hasil Pengolahan Data

Pengumpulan data yang dilakukan dengan proses *crawling* menggunakan *Harvest-Tweet* menghasilkan dataset yang berisi 5.808 unggahan pada *menfess* “Tanyarl”. Dataset ini memiliki 14 kolom atribut yang berisi *created\_at*, *favorite\_count*, *full\_text*, *id\_str*, *image\_url*, *in\_reply\_to\_screen\_name*, *lang*, *location*, *quote\_count*, *reply\_count*, *retweet\_count*, *tweet\_url*, *user\_id\_str*, dan *username*. Pada penelitian ini hanya atribut *full\_text* yang digunakan karena berisi teks unggahan pengguna Twitter yang relevan untuk proses analisis tingkat stres. Visualisasi contoh data hasil *crawling* beserta pelabelannya ditampilkan pada **Gambar 2**.

Sample data:

|   | tweet                                             | label        |
|---|---------------------------------------------------|--------------|
| 0 | @wawalostyou happy national best friend day sa... | tidak stress |
| 1 | @Helmylaz 3. Gak perlu khawatir kalau memang g... | tidak stress |
| 2 | ni kl gue bikin a day in my life hari ini past... | tidak stress |
| 3 | @nanaduungies life is still going on lagunya ...  | tidak stress |
| 4 | Sunday Funday! Sehat dan berkah selalu ya ĩ_□ ... | tidak stress |
| 5 | @moviemnfs and the second after i finished thi... | tidak stress |
| 6 | omak kalo kata aku mah kamu cakep banget. paca... | tidak stress |
| 7 | My pretty baby scoupsies I love youuu sayangku... | tidak stress |
| 8 | Day 15 Life after breakup hari ini tuh aku ten... | tidak stress |
| 9 | gue sadar kalau lu bisa dapet kerjaan yang gak... | tidak stress |

**Gambar 2.** Contoh data unggahan hasil *crawling* dan pelabelan

Struktur dataset hasil crawling disajikan pada **Tabel 1**, namun hanya lima atribut utama yang ditampilkan, yaitu *conversation\_id\_str*, *created\_at*, *favorite\_count*, *full\_text*, dan *id\_str*, karena atribut tersebut dianggap representatif untuk menggambarkan karakteristik data dan kebutuhan analisis.

**Tabel 1.** Dataset Tweet

| <i>conversation_id_str</i> | <i>created_at</i>                    | <i>favorite_count</i> | <i>full_text</i>                                                                                                                | <i>id_str</i> |
|----------------------------|--------------------------------------|-----------------------|---------------------------------------------------------------------------------------------------------------------------------|---------------|
| 1.87E+18                   | Sat Dec 14<br>14:57:25 +0000<br>2024 | 0                     | mau life update hari ini. best day banget pokonya di bulan ini. intinya seneng banget good things finally come to me hehe       | 1.87E+18      |
| 1.96423E+18                | Sat Sep 06<br>07:30:07 +0000<br>2025 | 15                    | mendem biar ga nangis ga sedih tp pasti ada moment tangisan itu pecah juga dan itu lega bgt                                     | 1.9642E+18    |
| 1.96421E+18                | Sat Sep 06<br>05:55:07 +0000<br>2025 | 273                   | Rasa tertinggal dipertemanan tuh ada aja ya sedih deh kenapa cm aku yg gagal ya? gmn sih biar ga iri sama pencapaian orang lain | 1.9642E+18    |
| 1.95052E+18                | Wed Jul 30 11:20:09<br>+0000 2025    | 42                    | kamar kalian tuh vibesnya adem gitu ga sih??? btw ruangan di rumah sender tuh biasa aja tapi pas di kamar beneran adem.....     | 1.95052E+18   |

Dataset yang sudah dikumpulkan, selanjutnya melalui tahap *preprocessing* untuk membersihkan teks dari URL, *mention*, *hashtag*, spasi berlebih, emoji dan simbol khusus. Contoh teks yang sudah melalui tahapan preprocessing ditunjukkan pada **Tabel 2**.

**Tabel 2.** Preprocessing Data

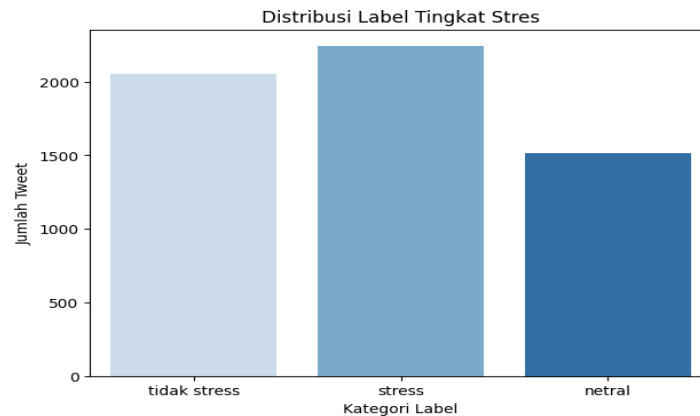
| <i>Full_text</i>                                                                                                                                                                                     | <i>Clean_text</i>                                                                                                                                                                              |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| fluent banget inggris nya :_____) aku seneng dengernya si adek selalu keren my adek my love <a href="https://t.co/G7UshtYHEN">https://t.co/G7UshtYHEN</a>                                            | fluent banget inggris nya aku seneng dengernya si adek selalu keren my adek my love                                                                                                            |
| PLS GW KESEL BGT LG NAIK TRANS DI DEKET LAMPU MERAH ADA MOBIL YG PARKIRINI MENGHAMBAT BGT PLISSS AAAAAAAAAAAAAQ KESEL SM MOBIL YG SUKA PARKIR SEMBARANGANNNN                                         | pls gw kesel bgt lg naik trans di dekat lampu merah ada mobil yg parkirini menghambat bgt plissss aaaaaaaaaaaaaq kesel sm mobil yg suka parkir sembarangannnn                                  |
| nyesel banget kenal dia tapi klo ga ketemu dia gw ga pernah tau rasanya patah hati yang bener.~.~² sakit banget tuh gimana wkwwkwk. tapi aku beneran udah bahagia liat kita jalan di masing masing:D | nyesel banget kenal dia tapi klo ga ketemu dia gw ga pernah tau rasanya patah hati yang bener sakit banget tuh gimana wkwwkwk. tapi aku beneran udah bahagia liat kita jalan di masing masingg |

Pelabelan data dilakukan melalui dua tahap, yaitu pelabelan otomatis berbasis kata kunci dan verifikasi manual oleh anotator. Daftar kategori dan kata kunci yang digunakan ditunjukkan pada **Tabel 3**.

**Tabel 3.** Kategori dan *Keywords*

| <i>Kategori</i> | <i>Keyword</i>                                                                                        |
|-----------------|-------------------------------------------------------------------------------------------------------|
| Stress          | Sedih, capek, kecewa, kesel, males, stress, marah, depresi, nangis, badmood, emosi                    |
| Tidak Stres     | seneng, bahagia, bersyukur, <i>happy</i> , semangat, selamat, puas, terharu, <i>excited</i> , gembira |
| Netral          | yaudah, biasa aja, <i>no comment</i> , bodo amat, <i>nothing</i> , <i>fine</i> aja, okelah, gitu aja  |

Distribusi hasil pelabelan bahwa label stres memiliki jumlah tweet paling banyak dengan jumlah 2.242 unggahan, sedangkan pada label tidak stress memiliki jumlah 2.050 tweet, dan label netral sebanyak 1.516 unggahan, sebagaimana ditampilkan pada **Gambar 3**.



**Gambar 3.** Diagram Pelabelan Semi-otomatis

Dominasi unggahan berlabel stres pada **Gambar 3** menunjukkan bahwa akun menfess “Tanyarl” menjadi medium ekspresi emosional bagi pengguna, khususnya dalam menyalurkan tekanan psikologis secara anonim. Dataset yang telah diberi label kemudian dibagi menjadi data *train* (70%), validasi (15%), dan *test* (15%) secara acak dengan tetap menjaga proporsi setiap kelas. Hasil pembagian data ditampilkan pada **Tabel 4**.

**Tabel 4.** Hasil Split Data

| Data Train | Data Validasi | Data Test |
|------------|---------------|-----------|
| 4.065      | 871           | 872       |
| 70%        | 15%           | 15%       |

### 3.2 Hasil Tokenisasi dan *Fine-tuning* IndoBERT

Tahapan tokenisasi dilakukan menggunakan tokenizer *IndoBERT* untuk mengonversi teks unggahan menjadi representasi numerik. Setiap teks unggahan diawali [CLS], diakhiri dengan [SEP], serta menggunakan token [PAD] untuk penyesuaian panjang input hingga maksimum 160 token. Hasil tokenisasi berupa *token\_ids* dan *attention\_mask* ditampilkan pada **Tabel 5**.

**Tabel 5.** Hasil Tokenisasi IndoBERT[illegible]

Proses *fine-tuning* dilakukan menggunakan model ‘indobenchmark/indobert-base-pl’ dengan optimizer AdamW dengan *learning rate* sebesar  $2 \times 10^{-5}$ . Untuk mengatasi ketidakseimbangan data, digunakan *weighted cross-entropy loss*. Konfigurasi *hyperparameter* yang digunakan ditampilkan pada **Tabel 6**.

Tabel 6. Konfigurasi Parameter

| Parameter     | Nilai                        |
|---------------|------------------------------|
| Learning rate | $2 \times 10^{-5}$           |
| Optimizer     | AdamW                        |
| Batch size    | 16                           |
| Epoch         | 4                            |
| Scheduler     | Linear warmup                |
| Loss function | Cross-Entropy + Class Weight |

Model dilatih selama 4 epoch dengan total 860 langkah pembaruan parameter. Nilai training loss akhir yang diperoleh sebesar 0.3125, dengan waktu pelatihan total adalah 410,33 detik. Model hasil *fine-tuning* selanjutnya digunakan untuk proses evaluasi pada data *test*.

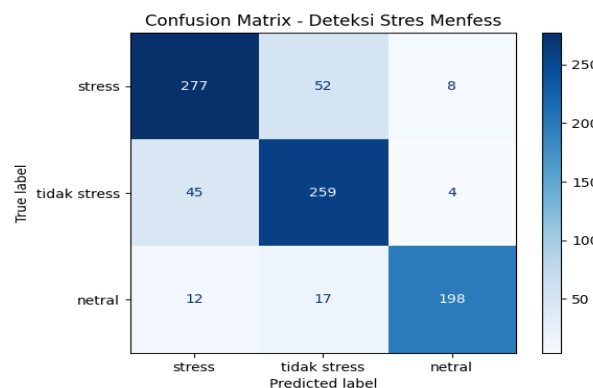
### 3.3 Hasil Evaluasi Model IndoBERT

Evaluasi model dilakukan menggunakan data *test* sebanyak 872 unggahan. Hasil *classification report* ditampilkan pada Tabel 7.

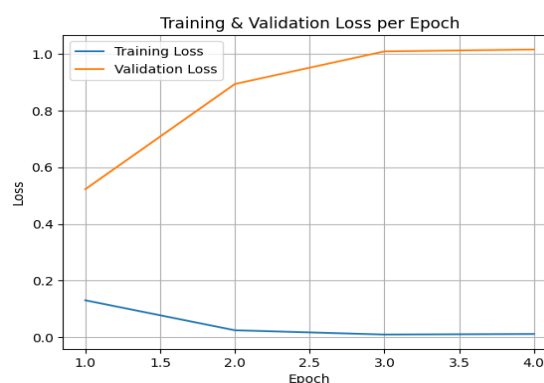
Tabel 7. Classification Report IndoBERT

| Kategori Kelas | Precision | Recall | F1-score | Jumlah Data |
|----------------|-----------|--------|----------|-------------|
| Stres          | 0,83      | 0,82   | 0,83     | 337         |
| Tidak Stress   | 0,79      | 0,84   | 0,81     | 308         |
| Netral         | 0,94      | 0,87   | 0,91     | 227         |
| Accuracy       |           |        | 0,84     | 872         |
| Macro avg      | 0,85      | 0,85   | 0,85     | 872         |
| Weighted avg   | 0,84      | 0,84   | 0,84     | 872         |

Model IndoBERT mencapai nilai akurasi sebesar 0,84, dengan *macro F1-score* mencapai 0,85 dan *weighted average F1-score* sebesar 0,84. *Confusion matrix* ditampilkan pada Gambar 4, sedangkan grafik *training loss* dan *validation loss* per epoch disajikan pada Gambar 5.



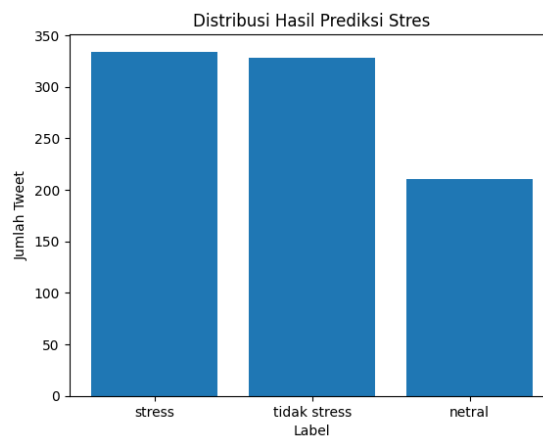
Gambar 4. Confusion Matrix



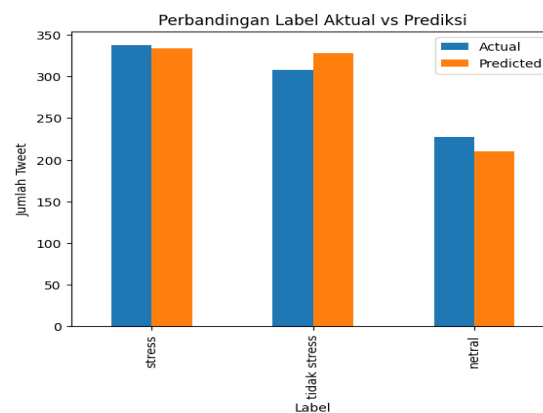
Gambar 5. Grafik Training Loss dan Validating Loss per Epoch

Terlihat bahwa *training loss* menurun secara signifikan seiring bertambahnya epoch, sedangkan *validation loss* menunjukkan kecenderungan meningkat. Pola tersebut mengindikasikan adanya kecenderungan overfitting, namun performa model pada data test tetap stabil sehingga kemampuan generalisasi masih terjaga. Untuk mengurangi potensi overfitting, digunakan dropout pada layer klasifikasi serta pembatasan jumlah epoch,

sementara penerapan early stopping dapat dipertimbangkan pada penelitian selanjutnya. **Gambar 6** menunjukkan distribusi hasil prediksi tingkat stres yang dihasilkan oleh model *IndoBERT* pada data *test*. Perbandingan antara label aktual dan label prediksi ditunjukkan pada **Gambar 7**.



**Gambar 6.** Distribusi Prediksi Tingkat Stres



**Gambar 7.** Perbandingan Pelabelan

Pada **Gambar 6** menunjukkan distribusi hasil prediksi tingkat stres yang dihasilkan oleh model *IndoBERT* pada data uji. Terlihat bahwa prediksi model relatif seimbang pada kategori stres dan tidak stres, sementara kategori netral memiliki jumlah prediksi yang lebih rendah. Selanjutnya, **Gambar 7** menyajikan perbandingan antara label aktual dan label prediksi untuk masing-masing kategori tingkat stres. Hasil perbandingan tersebut menunjukkan bahwa distribusi prediksi model mendekati distribusi label aktual, yang mengindikasikan bahwa model mampu meniru pola distribusi kelas secara cukup baik.

### 3.4 Hasil Evaluasi Model Baseline SVM

Sebagai model pembanding, penelitian ini mengimplementasikan Support Vector Machine (SVM) dengan representasi fitur TF-IDF dan kernel linear. Evaluasi dilakukan menggunakan data test yang sama dengan eksperimen *IndoBERT*, yaitu sebanyak 872 unggahan, untuk memastikan perbandingan performa yang adil. Berikut hasil evaluasi yang disajikan pada **Tabel 8**.

**Tabel 8.** Classification Report SVM

| Kategori Kelas | Precision | Recall | F1-score | Jumlah Data |
|----------------|-----------|--------|----------|-------------|
| Stres          | 0,76      | 0,78   | 0,77     | 308         |
| Tidak Stress   | 0,76      | 0,76   | 0,76     | 337         |
| Netral         | 0,94      | 0,89   | 0,91     | 227         |
| Accuracy       |           |        | 0,80     | 872         |
| Macro avg      | 0,82      | 0,81   | 0,82     | 872         |
| Weighted avg   | 0,81      | 0,80   | 0,80     | 872         |

Berdasarkan hasil evaluasi, model SVM memperoleh nilai akurasi sebesar 0,80. **Tabel 8** menyajikan *classification report* dari model SVM yang menunjukkan bahwa kelas netral memiliki performa paling baik dengan nilai *precision* sebesar 0,94, *recall* 0,89, dan *F1-score* 0,91. Sementara itu, kelas stres dan tidak stres masing-masing memperoleh nilai *F1-score* sebesar 0,77 dan 0,76.

### 3.5 Perbandingan Kinerja IndoBERT dan Baseline SVM

Perbandingan performa antara model *IndoBERT* dan baseline SVM ditunjukkan pada **Tabel 9**. Secara keseluruhan, *IndoBERT* menunjukkan performa yang lebih unggul pada seluruh metrik evaluasi. *IndoBERT* berhasil meningkatkan akurasi sebesar 4% dibandingkan SVM, serta memperoleh nilai *macro* dan *weighted F1-score* yang lebih tinggi.

**Tabel 9.** Perbandingan Model

| Model    | Akurasi | Precision | Recall | F1-score |
|----------|---------|-----------|--------|----------|
| IndoBERT | 0,84    | 0,85      | 0,85   | 0,85     |
| SVM      | 0,80    | 0,82      | 0,81   | 0,82     |

Keunggulan tersebut menunjukkan bahwa pendekatan berbasis transformer pada *IndoBERT* lebih efektif dalam memodelkan konteks semantik teks dibandingkan metode konvensional berbasis fitur statistik seperti SVM dengan TF-IDF. Sementara SVM masih mampu memberikan performa yang cukup kompetitif, keterbatasannya dalam merepresentasikan konteks bahasa yang kompleks menyebabkan performanya relatif lebih rendah dibandingkan *IndoBERT*. Temuan ini menegaskan bahwa model prelatih berbasis bahasa memiliki potensi yang lebih besar untuk tugas deteksi tingkat stres pada teks media sosial berbahasa Indonesia.

### 3.6 Pembahasan

Tujuan utama penelitian ini adalah untuk mengevaluasi kemampuan model *IndoBERT* dalam mendeteksi tingkat stres pengguna Twitter berdasarkan teks *menfess*. Hasil eksperimen menunjukkan bahwa *IndoBERT* mencapai performa yang lebih baik dibandingkan model baseline SVM, dengan akurasi sebesar 84%, *macro F1-score* 0,85, dan *weighted F1-score* 0,84. Capaian ini menunjukkan bahwa pendekatan berbasis transformer lebih efektif dalam menangani klasifikasi tingkat stres pada teks media sosial berbahasa Indonesia. Temuan ini sejalan dengan studi-studi sebelumnya yang menyatakan bahwa model transformer lebih unggul dalam analisis emosi dan kondisi psikologis pada teks media sosial dibandingkan metode berbasis fitur statistik. Hasil evaluasi menunjukkan bahwa kelas netral memiliki performa paling stabil pada kedua model, baik *IndoBERT* maupun SVM. Namun demikian, perbedaan performa yang signifikan terlihat pada kelas stres dan tidak stres. Model *IndoBERT* mampu mengurangi kesalahan klasifikasi pada kedua kelas tersebut dibandingkan SVM, yang masih menunjukkan tumpang tindih prediksi akibat kemiripan ekspresi emosional ringan.

Temuan ini sejalan dengan penelitian oleh Kenny Yane et al. [25] membahas analisis sentimen komentar netizen Twitter terhadap isu kesehatan mental masyarakat Indonesia menggunakan metode klasifikasi tradisional Naïve Bayes. Penelitian tersebut mengklasifikasikan sentimen ke dalam tiga kelas, yaitu positif, negatif, dan netral, dengan akurasi sekitar 79%. Hasil ini menunjukkan bahwa metode machine learning konvensional mampu menangkap kecenderungan opini publik, namun masih memiliki keterbatasan dalam memahami konteks emosional yang kompleks pada teks media sosial. Selanjutnya, penelitian oleh M. Sirojudin Mahdi Faza dan Moch. Taufik [24] menerapkan model *IndoBERT* untuk mendeteksi potensi sumber stres dalam teks media sosial. Penelitian tersebut menggunakan pendekatan deep learning berbasis transformer dengan klasifikasi sumber stres tertentu dan menghasilkan nilai *precision*, *recall*, dan *F1-score* yang sangat tinggi. Meskipun demikian, fokus penelitian tersebut lebih menekankan pada identifikasi kategori sumber stres, bukan tingkat stres secara umum. Berbeda dengan kedua penelitian tersebut, penelitian ini menggunakan model *IndoBERT* untuk mendeteksi tingkat stres pengguna Twitter dalam konteks *menfess* “Tanyarl”. Penelitian ini mengklasifikasikan data ke dalam tiga kelas, yaitu stres, tidak stres, dan netral, sehingga lebih menekankan pada kondisi psikologis pengguna. Dengan demikian, penelitian ini mengisi celah penelitian dengan mengombinasikan kekuatan model transformer dan pendekatan klasifikasi tingkat stres yang lebih umum pada teks media sosial berbahasa Indonesia.

Analisis kesalahan klasifikasi pada *IndoBERT* menunjukkan bahwa tweet dengan makna implisit atau ambigu masih sulit diprediksi secara tepat. Contoh kesalahan klasifikasi ditemukan pada teks unggahan: “capek banget rasanya pengen berhenti tapi harus kuat”. Unggahan tersebut memiliki label aktual Stres, namun diprediksi sebagai Tidak Stres oleh model. Kesalahan ini terjadi karena keberadaan kata “kuat” yang sering diasosiasikan dengan konteks positif, sehingga memengaruhi keputusan model. Temuan ini menunjukkan bahwa meskipun *IndoBERT* memiliki performa yang baik, model masih memiliki keterbatasan dalam memahami nuansa emosional yang bersifat implisit atau kontekstual.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa model *IndoBERT* yang telah melalui proses *fine-tuning* mampu mendeteksi tingkat stres pada unggahan *menfess* “Tanyarl” secara lebih akurat dibandingkan baseline SVM. Secara keseluruhan, penelitian ini berkontribusi dengan menyajikan evaluasi komprehensif antara

model transformer dan baseline konvensional dalam mendeteksi tingkat stres pada teks *menfess* berbahasa Indonesia. Meskipun demikian, pengembangan lanjutan masih diperlukan, khususnya untuk meningkatkan kemampuan model dalam memahami ekspresi emosional yang bersifat implisit dan kontekstual.

## 4. KESIMPULAN

Penelitian ini membuktikan bahwa pendekatan *fine-tuning IndoBERT* efektif dalam mendeteksi tingkat stres pada unggahan *menfess* “Tanyarl” berbahasa Indonesia, dengan capaian akurasi sebesar 84% dan *macro F1-score* sebesar 85%. Hasil perbandingan dengan baseline *Support Vector Machine* (SVM) menunjukkan bahwa *IndoBERT* secara konsisten memberikan performa yang lebih baik pada seluruh metrik evaluasi, sehingga menegaskan keunggulan model berbasis transformer dibandingkan metode klasifikasi konvensional. Secara ilmiah, hasil ini memperkuat temuan bahwa model bahasa pralatih berbasis transformer yang dikembangkan khusus untuk bahasa Indonesia mampu menangkap konteks emosional dan psikologis dalam teks media sosial yang bersifat informal. Dari sisi praktis, penelitian ini berkontribusi sebagai dasar pengembangan sistem pemantauan stres berbasis media sosial yang dapat dimanfaatkan sebagai alat bantu awal dalam mengidentifikasi indikasi stres pada pengguna, khususnya pada platform anonim seperti *menfess*. Sistem semacam ini berpotensi mendukung pihak terkait, seperti peneliti, institusi akademik, atau praktisi kesehatan mental, dalam melakukan analisis tren kondisi psikologis masyarakat secara lebih cepat dan berbasis data. Keterbatasan penelitian ini terletak pada ketidakseimbangan distribusi data antar kelas, yang memengaruhi kemampuan model dalam mengenali kelas netral secara optimal. Selain itu, ambiguitas dan variasi gaya bahasa pada unggahan *menfess* menyebabkan tumpang tindih pola linguistik antar kelas, sehingga masih terjadi kesalahan klasifikasi pada beberapa kasus. Oleh karena itu, penelitian selanjutnya disarankan untuk menggunakan dataset yang lebih besar dengan distribusi kelas yang seimbang guna meningkatkan stabilitas dan generalisasi model. Penelitian mendatang juga dapat menerapkan skema pelabelan yang lebih terstruktur, seperti pelabelan *multi-annotator* dengan pengukuran *inter-annotator agreement*, serta mengeksplorasi arsitektur model lain atau pendekatan ensemble untuk meningkatkan kemampuan model dalam membedakan ekspresi stres yang bersifat ambigu.

## REFERENCES

- [1] M. Fachriza and M. Munawar, “Analisis Sentimen Kalimat Depresi Pada Pengguna Twitter Dengan Naive Bayes, Support Vector Machine, Random Forest,” *Komputek*, vol. 7, no. 2, pp. 49–58, 2023, doi: 10.24269/jkt.v7i2.2218.
- [2] I. R. Hidayat and W. Maharani, “General Depression Detection Analysis Using IndoBERT Method,” *Int. J. Inf. Commun. Technol.*, vol. 8, no. 1, pp. 41–51, Aug. 2022, doi: 10.21108/ijoict.v8i1.634.
- [3] M. Johan and S. Aurelia Azka, “Prediction of Alleged Stress Symptoms based on Indonesian Sentiment Lexicon using Multilayer Perceptron,” *G-Tech J. Teknol. Terap.*, vol. 7, no. 3, pp. 958–966, Jul. 2023, doi: 10.33379/gtech.v7i3.2611.
- [4] C. L. Odgers and M. R. Jansen, “Adolescent mental health in the digital age,” in *Handbook of Adolescent Digital Media Use and Mental Health*, J. Nesi, E. H. Telzer, and M. J. Prinstein, Eds., Oxford University Press, 2022. doi: 10.1017/9781108976237.
- [5] J. R. Wiyani, I. Indriati, and S. Sutrisno, “Klasifikasi Stres berdasarkan Unggahan pada Media Sosial Twitter menggunakan Metode Support Vector Machine dan Seleksi Fitur Information Gain,” *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 6, no. 12, pp. 6003–6009, 2022, [Online]. Available: <http://j-ptiik.ub.ac.id>
- [6] J. Suler, “The Online Disinhibition Effect Revisited: Implications for Self-disclosure and Mental Health,” *Cyberpsychology, Behav. Soc. Netw.*, vol. 26, no. 1, pp. 7–13, 2023, doi: 10.1089/cyber.2022.0235.
- [7] S. Basuki, R. Indrabayu, and N. A. Effendy, “Automatic Categorization of Mental Health Frame in Indonesian X (Twitter) Text using Classification and Topic Detection Techniques,” *Khazanah Inform. J. Ilmu Komput. dan Inform.*, vol. 10, no. 2, pp. 104–110, 2025, doi: 10.23917/khif.v10i2.3328.
- [8] Gloriabarus, “Hasil Survei I-NAMHS: Satu dari Tiga Remaja Indonesia Memiliki Masalah Kesehatan Mental,” Universitas Gadjah Mada. [Online]. Available: <https://ugm.ac.id/id/berita/23086-hasil-survei-i-namhs-satu-dari-tiga-remaja-indonesia-memiliki-masalah-kesehatan-mental?>
- [9] S. Winurini, “Pemeriksaan Kesehatan Mental Gratis bagi Remaja,” 2025.
- [10] Kementerian Kesehatan Republik Indonesia, “Survei Kesehatan Indonesia 2023,” Jakarta, 2023. [Online]. Available: <https://www.badankebijakan.kemkes.go.id/hasil-ski-2023/>
- [11] Y. Tolla and Kusriani, “Deteksi Stres dan Depresi Unggahan Media Sosial dengan Machine Learning,” *J. FASILKOM*, vol. 15, no. 1, pp. 84–92, Apr. 2025.
- [12] T. Nguyen, D. M. Pham, T. H. Nguyen, H. Pham, R. A. Patel, and M. Dredze, “Online Social Media Language and Mental Health Detection,” *J. Med. Internet Res.*, vol. 24, no. 6, 2022, doi: 10.2196/34578.
- [13] O. Campesato, *Transformer, BERT, and GPT*. Packt Publishing, 2023. [Online]. Available: <https://buku.io/book/170842/transformer-bert-and-gpt%0A>
- [14] S. Aftan and H. Shah, “A Survey on BERT and Its Applications Sulaiman,” *IEEE*, Jul. 2023.
- [15] A. Fitro, A. P. Ristadi Pinem, O. S. Bachri, and Chartini, “Cyberbullying Detection on Instagram using IndoBERTa Model,” in *International Conference on Digital Business Innovation and Technology Management*, ICONBIT, 2025, pp. 1441–1447.
- [16] G. F. Situmorang and R. Purba, “Deteksi Potensi Depresi dari Unggahan Media Sosial X Menggunakan IndoBERT,” *Build. Informatics, Technol. Sci.*, vol. 6, no. 2, pp. 649–661, Sep. 2024, doi: 10.47065/bits.v6i2.5496.

- [17] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, “IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP,” *COLING 2020 - 28th Int. Conf. Comput. Linguist. Proc. Conf.*, pp. 757–770, Nov. 2020, doi: 10.18653/v1/2020.coling-main.66.
- [18] A. C. Saputra, A. S. Saragih, and D. Ronaldo, “Prediksi Emosi Dalam Teks Bahasa Indonesia Menggunakan Model IndoBERT,” *J. Teknol. Inf. J. Keilmuan dan Apl. Bid. Tek. Inform.*, vol. 19, no. 1, pp. 1–15, 2025, doi: 10.47111/JTI.
- [19] D. Suhartono, I. F. Saputra, A. R. Pratama, and G. Nathaniel, “Psychological Stress Detection Using Transformer-Based Models,” *ComTech Comput. Math. Eng. Appl.*, vol. 15, no. 1, pp. 65–71, Jun. 2024, doi: 10.21512/comtech.v15i1.11105.
- [20] F. Koto, J. H. Lau, and T. Baldwin, “IndoBERTweet: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization,” *EMNLP 2021 - 2021 Conf. Empir. Methods Nat. Lang. Process. Proc.*, pp. 10660–10668, Nov. 2021, doi: 10.18653/v1/2021.emnlp-main.833.
- [21] S. Gupta, “Fine-Tuning Pre-trained Language Models for Text Classification: A Comparative Study,” *J. Comput. Linguist. Nat. Lang. Process.*, vol. 9, no. 1, pp. 67–75, 2021, doi: 10.1016/j.jclnlp.2021.01.005.
- [22] G. Xu, Y. Meng, X. Qiu, Z. Yu, and X. Wu, “Sentiment Analysis of Social Media Texts Using Deep Learning Models,” *Inf. Process. Manag.*, vol. 59, no. 2, 2022, doi: 10.1016/j.ipm.2021.102928.
- [23] C. Sun, X. Qiu, Y. Xu, and X. Huang, “How to fine-tune BERT for text classification?,” in *Chinese Computational Linguistics*, 2021, pp. 194–206. doi: 10.1007/978-3-030-32381-3\_16.
- [24] M. S. Mahdi Faza and M. Taufik, “Penerapan Model Indobert Untuk Deteksi Potensi Sumber Stres Dalam Teks Media Sosial,” *J. Rekayasa Sist. Inf. dan Teknol.*, vol. 3, no. 1, pp. 271–283, Aug. 2025, doi: 10.70248/jrsit.v3i1.3046.
- [25] K. Yan, D. Arisandi, and T. Tony, “Analisis Sentimen Komentar Netizen Twitter Terhadap Kesehatan Mental Masyarakat Indonesia,” *J. Ilmu Komput. dan Sist. Inf.*, vol. 10, no. 1, 2022, doi: 10.24912/jiksi.v10i1.17865.