

Deteksi Depresi di Twitter Menggunakan Metode CNN-BiGRU dengan Fitur Ekspansi FastText

1st MUHAMMAD ARIF DWI PUTRA

line 2: Informatika

Telkom University

Bandung, Indonesia

arifdwiputra@student.telkomuniversity.ac.id

2nd Erwin Budi Setiawan

line 2: Informatika

Telkom University

Bandung, Indonesia

erwinbudisetiawan@telkomuniversity.ac.id

Abstrak — Depresi merupakan gangguan mental yang sering tidak terdeteksi dengan baik, meskipun banyak mempengaruhi individu di seluruh dunia. Media sosial, khususnya Twitter, menjadi platform yang digunakan untuk mengekspresikan emosi, termasuk gejala depresi yang tidak diungkapkan secara eksplisit. Penelitian ini bertujuan untuk mengembangkan model deteksi depresi di Twitter dengan metode hybrid CNN-BiGRU yang dilengkapi dengan ekspansi fitur FastText. CNN digunakan untuk mengekstraksi pola lokal dalam teks, sementara BiGRU memproses urutan kata dari dua arah untuk menangkap konteks yang lebih dalam. Ekspansi FastText bertujuan untuk menangani variasi kosakata dan meningkatkan akurasi dalam mendeteksi makna implisit dalam teks. Penelitian ini penting karena banyak pengguna media sosial yang tidak mendapatkan perawatan depresi yang memadai. Deteksi otomatis melalui teks Twitter dapat menjadi solusi untuk intervensi dini. Pengujian menggunakan dataset tweet berbahasa Indonesia menunjukkan bahwa model hybrid BiGRU-CNN dengan FastText mencapai akurasi tertinggi sebesar 80,65% pada korpus IndoNews dengan optimizer RMSprop. Model ini diharapkan dapat berkontribusi dalam deteksi depresi dan mendukung intervensi kesehatan mental.

Kata kunci— Depresi, Twitter, CNN-BiGRU, FastText, Deteksi Emosi, Model Hybrid.

I. PENDAHULUAN

Depresi adalah gangguan mental serius yang dapat menurunkan kualitas hidup seseorang secara signifikan. Menurut *World Health Organization* (WHO), lebih dari 280 juta orang di dunia mengalami depresi. Namun, banyak penderita yang tidak terdiagnosis atau tidak mendapatkan perawatan yang memadai, seringkali karena stigma sosial atau keterbatasan akses ke layanan kesehatan mental. Oleh karena itu, deteksi dini menjadi penting untuk intervensi yang lebih cepat dan tepat. Selain itu, Ada bukti bahwa penggunaan media sosial yang berlebihan dan kompulsif terkait dengan berbagai masalah kesehatan mental [1]. Penelitian menunjukkan Penggunaan media sosial secara berlebihan dan kompulsif sering dikaitkan dengan masalah kesehatan mental[2].

Di era digital, media sosial seperti Twitter menjadi ruang bagi individu untuk mengekspresikan emosi, termasuk tanda-tanda depresi yang mungkin tidak diekspresikan secara eksplisit. Kecemasan dan depresi adalah gangguan mental yang paling umum, dan stigmatisasi sosial sering menyebabkan penderita menghindari pengobatan [3]. Dengan berkembangnya teknologi machine learning, data teks di

Twitter berpotensi digunakan untuk mendeteksi gejala depresi secara otomatis. Namun, tantangan utama dalam deteksi depresi berbasis teks adalah kemampuan model untuk memahami konteks bahasa yang kompleks dan beragam. Banyak model yang ada saat ini hanya berfokus pada analisis fitur lokal atau sekuensial teks, tanpa mempertimbangkan keduanya secara bersamaan. Oleh karena itu, diperlukan pendekatan yang lebih canggih.

Kombinasi *Convolutional Neural Networks* (CNN) dan *Bidirectional Gated Recurrent Units* (BiGRU) menawarkan potensi yang besar dalam hal ini. CNN mampu mengenali pola lokal dalam teks dengan dua lapisan utama lapisan konvolusi dan lapisan pooling. Lapisan konvolusi digunakan untuk mengekstrak informasi input yang disebut fitur. Dan untuk peta fitur keluaran dari lapisan konvolusi di-downsampling dengan menggunakan lapisan pooling[4] [5]. Sementara BiGRU memiliki kemampuan untuk memproses dalam dua arah yang menangkap hubungan berurutan antar kalimat dalam *tweet* [6]. Ditambah dengan representasi kata FastText, yang unggul dalam menangkap makna kata dalam 8 berbagai konteks, memungkinkan pembuatan model yang kuat pada dataset yang besar dengan menyediakan representasi untuk kata-kata yang mungkin tidak ditemukan selama proses pelatihan. Model ini diharapkan dapat meningkatkan akurasi deteksi depresi di Twitter. Pendekatan ini bertujuan untuk mengembangkan model deteksi depresi berbasis teks yang lebih akurat dan dapat diandalkan. Dengan memanfaatkan CNN, BiGRU, dan FastText, diharapkan dapat tercipta sistem deteksi yang lebih efektif, yang tidak hanya mengidentifikasi tandatanda depresi, tetapi juga menyediakan alat yang berguna untuk intervensi dini dalam kesehatan mental.

II. STUDI TERKAIT

Dalam studi terkait ini terdapat beberapa penelitian yang relevan Meskipun dengan masalah deteksi depresi penelitian pertama oleh Pabian [8], Penggunaan gabungan antara CNN dan BiGRU memberikan keuntungan dalam menangkap pola spasial dan konteks temporal, yang dapat meningkatkan akurasi deteksi untuk teks berbahasa Indonesia. Metode ini juga efektif dalam memahami kompleksitas teks, sangat relevan untuk deteksi emosi atau kondisi seperti depresi. *FastText* dapat menangani kata-kata yang tidak ada dalam kamus tradisional dengan menganalisis struktur subword, yang sangat bermanfaat untuk bahasa Indonesia dengan morfologi yang kompleks. Ini dapat membantu mendeteksi nuansa emosi dalam *tweet* yang berkaitan dengan depresi.

Penelitian selanjutnya oleh Merinda Lestandy[6]. Penggunaan gabungan antara CNN dan BiGRU memberikan keuntungan dalam menangkap pola spasial dan konteks temporal, yang dapat meningkatkan akurasi deteksi untuk teks berbahasa Indonesia. Meskipun penggunaan *Reddit* sebagai sumber data sangat relevan untuk penelitian ini, hasilnya mungkin tidak sepenuhnya dapat diterapkan pada data Twitter.

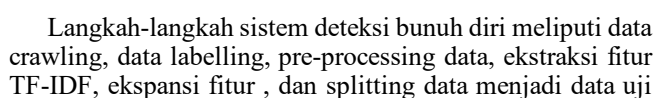
Selanjutnya penelitian yang dilakukan oleh Setiawan [10]. Ekspansi fitur melalui embedding kata memungkinkan model untuk memahami hubungan semantik antara kata-kata yang mungkin tidak ditemukan dalam *tweet*. Walaupun metode yang digunakan sangat baik dalam menangani data teks, penggunaan teknik *deep learning*.

dan data latih. Serta klasifikasi data dengan empat model: CNN, BigRU, CNN-BiGRU, BiGRU-CNN. Terakhir, kinerja sistem yang dibangun akan dievaluasi.

Mengumpulkan data melalui proses crawling[11] merupakan salah satu cara untuk mengambil informasi dari berbagai sumber digital seperti media sosial. Dalam penelitian ini, data diambil dari platform Twitter yang menggunakan bahasa Indonesia dengan bantuan API resmi dari Twitter yang mempermudah proses pengambilan data. Proses *crawling* difokuskan pada *tweet* yang berpotensi mengandung indikasi depresi, seperti ungkapan perasaan cemas, kesendirian, keputusan, atau perasaan tidak berharga. Semua perilaku ini menjadi fokus dalam proses identifikasi depresi. Proses pengumpulan data menghasilkan total 58.115 data. Sebaran data dapat dilihat pada Tabel 1.

Kata Kunci	Jumlah
cemas	9,245
kesendirian	6,842
psikis	7,189
emotional	6,318
psikologi depresi	7,915
insomnia	8,427
frustasi	6,031
melankolis	6,148
Total Data	58,115

Pelabelan data merupakan tahap berikutnya setelah proses pengumpulan data selesai. Langkah ini bertujuan untuk membantu model dalam membedakan antara data yang mengandung indikasi depresi atau tidak. Proses pelabelan dilakukan dengan menggunakan format *biner*, di mana data yang memiliki indikasi depresi diberi label "1," sedangkan data tanpa indikasi tersebut diberi label "0." Pelabelan data dilakukan oleh 5 annotator yang sebelumnya dilakukan persamaan persepsi terkait dengan pengertian depresi dan ciri



cirinya. Pelabelan ditetapkan dengan prinsip *majority vote*. Contoh pelabelan data bisa dilihat Tabel 2.

Tabel 2 Pelabelan Data

<i>Tweet</i>	Label
bisa lebih bangak lagi karena sempet mogok dengerin lagu for a whole month karena terlalu depresi	1
Dia di awal chapter tu depresi dan suicidal karena perasaannya ke seorang cewek. Terus dia jadi lebih baik karena menemukan seorang yang bener bener tulus sama dia. Eh ternyata orang itu juga aslinya suicidal dan ga punya tujuan hidup	0
Tidur sore bisa ganggu keseimbangan hormon yg berpengaruh pd kemampuan tubuh mengtsi stress dan resiko depresi bisa meningkat	1

Setelah dilakukan proses pelabelan, jumlah distribusi kelas indikasi depresi dan tanpa indikasi depresi setelah di labelling dapat dilihat pada Tabel 3.

Tabel 3 Hasil Labelling

Label	Total
1	29,421
0	28,694
Total	58,115

C. PRE-PROCESSING DATA

Data yang diperoleh dari platform media sosial seperti Twitter sering kali berbentuk data mentah yang belum terstruktur dan mengandung banyak gangguan, seperti simbol, angka, atau karakter yang tidak relevan. Oleh karena itu, langkah pertama dalam proses pre-processing adalah pembersihan (*cleaning*), dengan memakai *re* (*Regular Expression*) yang bertujuan untuk menghilangkan elemen-elemen yang mengganggu, termasuk tagar (#), nama pengguna (@), URL, dan emotikon yang sering muncul dalam *tweet*[11]. Setelah data dibersihkan, langkah selanjutnya adalah *case folding*, yaitu mengubah seluruh teks menjadi huruf kecil untuk memastikan konsistensi dalam analisis dan menghindari perbedaan makna akibat penggunaan kapitalisasi yang berbeda. Kemudian, dilakukan *tokenization*, yang membagi teks menjadi unit-unit terkecil, seperti kata-kata atau token yang lebih mudah dianalisis[11].

Setelah proses tokenisasi, terdapat tahap removed stopwords, dengan memakai library Sastrawi. yang bertujuan untuk kata-kata yang tidak memberikan kontribusi berarti terhadap analisis, seperti "dan", "yang", atau "dengan", akan dihapus. Proses ini bertujuan untuk memastikan bahwa teks yang akan dianalisis lebih terfokus pada kata-kata yang relevan dan bermakna dan juga mengurangi noise. Melalui serangkaian tahapan *pre-processing* ini, teks dari media sosial yang sebelumnya tidak terstruktur dan penuh gangguan dapat diubah menjadi format yang lebih bersih dan terstruktur. Setelah proses selesai, dataset tanpa indikasi depresi berjumlah 29,421, sementara yang mengindikasikan depresi berjumlah 28,694, sebagaimana dapat dilihat pada Tabel 4.

Tabel 4 Data setelah Preprocessing

Label	Total
1	29,421
0	28,694
Total	58,115

D. Feature Extraction TF-IDF

Ekstraksi fitur merupakan langkah untuk menghitung bobot setiap kata dalam teks dan mengubah kata-kata tersebut menjadi representasi vektor digital. Proses ini adalah tahap pertama dalam klasifikasi teks. Dalam penelitian ini, metode yang digunakan untuk ekstraksi fitur adalah TF-IDF. *Term Frequency-Inverse Document Frequency* (TF-IDF) adalah teknik yang digunakan untuk memberikan vektor pada kata-kata (*term*) dalam sebuah dokumen[11]. *Term Frequency* (TF) mengukur frekuensi kemunculan suatu kata dalam dokumen, sedangkan *Inverse Document Frequency* (IDF) menghitung logaritma dari kebalikan proporsi dokumen yang mengandung kata tersebut dalam seluruh korpus. Dengan menggabungkan keduanya, TF-IDF memberikan cara untuk menilai sejauh mana pentingnya suatu kata dalam dokumen tertentu dibandingkan dengan dokumen-dokumen lainnya dalam korpus. Berikut adalah langkah untuk menghitung vektor dalam metode TF-IDF:

$$tf_t = 1 + \log (tf_t) \quad (1)$$

$$idf_t = \log \left(\frac{D}{df_t} \right) \quad (2)$$

$$W_{t,d} = tf_t \times idf_t \quad (3)$$

Bobot dokumen ke-i terhadap kata ke-j ($W_{t,d}$) menunjukkan pentingnya kata ke-j dalam dokumen ke-i. Nilai *Term Frequency* (TF_t) mencerminkan frekuensi kemunculan kata ke-t dalam dokumen, semakin sering muncul semakin tinggi nilainya, dihitung dengan rumus $1 + \log (tf_t)$. Sebaliknya, *Inverse Document* (IDF_t) menunjukkan kelangkaan kata ke-t di seluruh dokumen dalam korpus, dihitung dengan rumus $\log \left(\frac{D}{df_t} \right)$, di mana D adalah total jumlah dokumen dan (DF_t) adalah jumlah dokumen yang mengandung kata ke-t. Semakin jarang kata muncul di banyak dokumen, semakin tinggi nilai IDF-nya.

E. Feature Expansion

Penelitian ini menerapkan metode word embedding dengan menggunakan FastText sebagai representasi kata. FastText adalah model yang dikembangkan oleh *Facebook*, yang memperluas konsep Word2vec dengan menambahkan representasi berbasis sub-kata. Berbeda dengan Word2vec yang hanya memperhitungkan kata utuh, FastText juga memperhatikan struktur internal kata, seperti suku kata atau karakter-karakter dalam kata tersebut. Pendekatan ini memungkinkan model untuk mengatasi kata-kata yang tidak ada dalam korpus pelatihan dengan lebih efektif, dengan memanfaatkan konteks kata-kata di sekitarnya [8] [12].

Dalam penelitian ini, FastText digunakan untuk mengukur kemiripan antar kata dengan menghasilkan vektor kata yang lebih kaya informasi, serta mampu menangani kata-kata yang tidak ada dalam korpus, dengan bergantung pada konteks

sekitar kata tersebut. Tujuan utama penggunaan FastText adalah untuk mengurangi ketidaksesuaian kosakata dan menangani kata-kata baru yang mungkin muncul, melalui representasi sub-kata dalam "word embeddings"[13]

Penelitian ini menggunakan dua korpus, yaitu *tweet* dan *indonews* lanjutan yang dikumpulkan dari berbagai sumber untuk memastikan representasi kata yang lebih komprehensif dan mencakup berbagai topik serta variasi bahasa. Korpus *tweet* diambil dari platform media sosial Twitter, sementara korpus *IndoNews* diperoleh dari berita dan artikel online dalam bahasa Indonesia. Kedua korpus ini memungkinkan model untuk menangkap konteks bahasa yang lebih kaya dan memperluas cakupan analisis, terutama dalam mendeteksi pola atau indikasi depresi dalam teks berbahasa Indonesia, bisa dilihat pada Tabel 5.

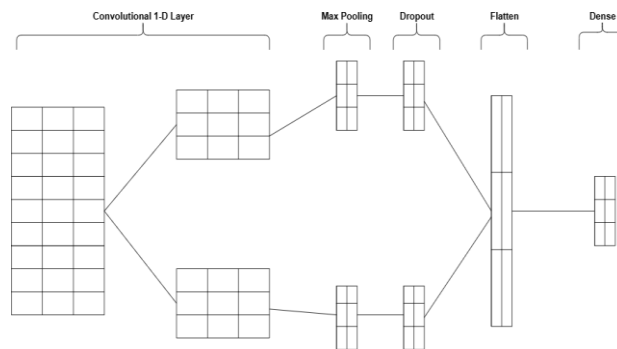
Tabel 5 Build Corpus

Corpus	Total
<i>Tweet</i>	58,115
Indonews	111,526
Indonews+Tweets	169,641

F. Model Klasifikasi

Setelah data melewati tahapan pre-processing dan dipresentasikan dalam bentuk vektor, data akan diklasifikasikan menggunakan beberapa algoritma klasifikasi, seperti *Convolutional Neural Network* (CNN), *Bidirectional Gated Recurrent Unit* (BiGRU), model *hybrid deep learning* CNN-BiGRU, dan model *hybrid deep learning* BiGRU-CNN. *Convolutional Neural Network* (CNN) adalah salah satu jaringan yang banyak digunakan dalam bidang pemrosesan bahasa alami (NLP), khususnya untuk tugas-tugas yang melibatkan data teks. CNN memanfaatkan lapisan konvolusi untuk mengekstrak vektor dari data input, yang melibatkan penggunaan filter atau kernel yang bergerak di atas data input untuk mendeteksi vektor fitur lokal. Arsitektur CNN ini terdiri dari beberapa layer, antara lain input, convolutional layer, pooling layer, dropout layer, flattening, fully connected layer, dan output[14].

Arsitektur CNN dimulai dengan lapisan input yang menerima data dalam format vektor. Lapisan Conv1D dengan 32 filter, kernel 5, dan fungsi aktivasi ReLU digunakan untuk mengekstrak fitur teks yang penting. Lapisan MaxPooling1D berukuran 5 kemudian melakukan down-sampling. Untuk mencegah overfitting, lapisan Dropout berukuran 0.3 digunakan. Setelah output konvolusi dan pooling diratakan dengan flattening, dua lapisan densitas digunakan untuk memprosesnya. Lapisan pertama memiliki 32 neuron dan ReLU, dan lapisan kedua memiliki 2 neuron dan sigmoid untuk klasifikasi biner. Model ini menggunakan loss function binary crossentropy dan Adam *optimizer*, dengan EarlyStopping digunakan untuk mencegah overfitting. Arsitektur CNN dapat dilihat pada Gambar 2.



Gambar 3 Arsitektur CNN

Bidirectional Gated Recurrent Unit (BiGRU) adalah jenis *Recurrent Neural Network* (RNN) yang dirancang untuk memproses data berurutan dengan memperhatikan konteks dari kedua arah, yaitu dari masa lalu dan masa depan. BiGRU menggantikan neuron tradisional dalam RNN dengan sel GRU yang lebih efisien[15]. Arsitektur BiGRU terdiri dari dua komponen utama, yaitu komponen yang memproses data dari arah waktu ke depan (forward) dan komponen yang memproses data dari arah waktu ke belakang (backward).

Pada BiGRU, mekanisme gerbang mengatur aliran informasi. Terdapat dua gerbang utama dalam BiGRU, yaitu *Update Gate* dan *Reset Gate*. *Update Gate* mengontrol seberapa banyak informasi yang perlu dipertahankan dari langkah sebelumnya dan seberapa banyak yang perlu diperbarui. Sementara itu, *Reset Gate* menentukan seberapa banyak informasi dari langkah sebelumnya yang perlu dilupakan saat memproses input baru. Dengan memproses data dari kedua arah, BiGRU dapat menangkap informasi yang lebih luas dan relevan, yang sangat berguna dalam memahami urutan data yang lebih kompleks [15].

Model ini menggunakan Bidirectional GRU (BiGRU) dengan dua lapisan untuk menangkap hubungan jangka panjang dalam data sekuensial. Pada tahap pertama, lapisan Bidirectional GRU dengan 64 neuron digunakan untuk mengolah urutan data dan mempertimbangkan informasi dari kedua arah (masa lalu dan masa depan). Pada tahap kedua, lapisan Bidirectional GRU dengan 32 neuron digunakan untuk melanjutkan pemrosesan data dengan mempertimbangkan konteks lebih lanjut. Masing-masing lapisan GRU diikuti oleh lapisan Dropout dengan tingkat 0.3 dan 0.2 untuk mengurangi risiko overfitting. Selanjutnya, lapisan Dense dengan 32 neuron dan aktivasi ReLU digunakan untuk menangkap hubungan non-linear dalam data dan memperkuat representasi yang dihasilkan oleh GRU.

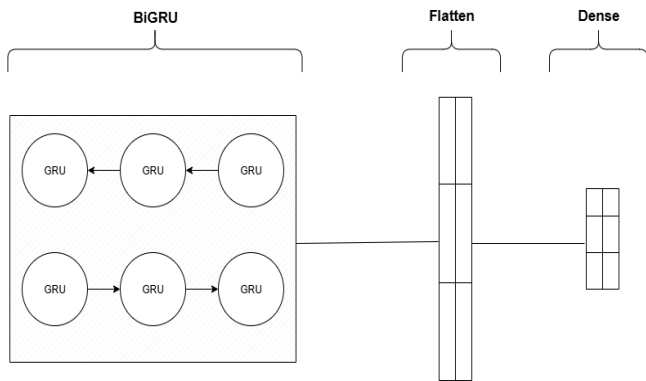
Pada tahap akhir, lapisan Dense dengan jumlah neuron sesuai dengan jumlah kelas (3 kelas) dan aktivasi softmax digunakan untuk klasifikasi multi-kelas. Model ini dikompilasi menggunakan Adam *optimizer* dengan categorical_crossentropy sebagai fungsi loss untuk klasifikasi multi-kelas. Jika validation loss tidak menunjukkan perbaikan selama pelatihan, EarlyStopping akan diterapkan untuk menghentikan pelatihan lebih awal, sehingga mencegah overfitting. Arsitektur BiGRU ini digunakan untuk menangkap pola dan konteks dalam data sekuensial, dengan tujuan meningkatkan akurasi dalam klasifikasi teks, seperti deteksi depresi dalam *tweet*.

Arsitektur BiGRU ini dapat dilihat pada Gambar 4, yang menggambarkan bagaimana kedua lapisan Bidirectional

GRU bekerja bersama-sama dengan lapisan Dense untuk menghasilkan output yang optimal.

Model *deep learning hybrid* diterapkan dengan menggabungkan dua arsitektur, yaitu CNN (*Convolutional Neural Network*) dan BiGRU (*Bidirectional Gated Recurrent Unit*). Keunggulan utama dari CNN-BiGRU terletak pada kemampuannya untuk mengombinasikan CNN, yang berfungsi untuk mengekstraksi fitur lokal dari teks, seperti pola kata, huruf, dan frasa, dengan BiGRU yang dapat memahami konteks dan urutan informasi dalam data. CNN berfokus pada pengidentifikasian pola lokal dalam teks, sementara BiGRU memanfaatkan konteks dari kedua arah untuk menangkap hubungan jangka panjang antar elemen dalam urutan data teks[15].

Dengan memadukan kedua model ini, dihasilkan representasi teks yang lebih komprehensif dan informatif, memungkinkan model untuk lebih efektif dalam menangkap hubungan antar kata dan memahami konteks kalimat secara lebih mendalam. Kombinasi kekuatan CNN dalam ekstraksi fitur lokal dan BiGRU dalam menangkap ketergantungan urutan data menjadikan model ini sangat efektif untuk tugas-tugas pemrosesan bahasa alami (NLP), seperti klasifikasi teks, analisis sentimen, atau deteksi depresi dalam teks. Model CNN-BiGRU ini dirancang untuk mengoptimalkan pemrosesan teks yang kompleks dengan menghasilkan representasi yang lebih kuat dan akurat.



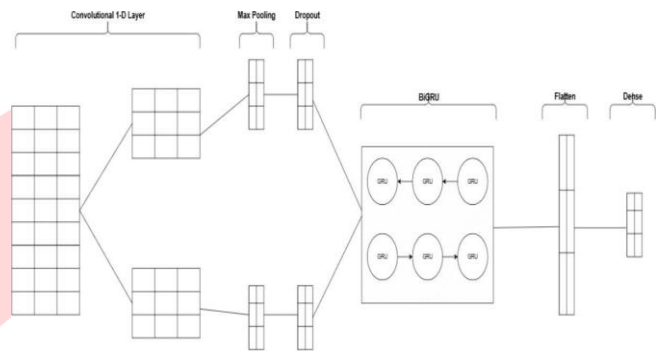
Gambar 4 Arsitektur BiGRU

Model CNN-BiGRU hybrid ini dimulai dengan Embedding layer yang mengonversi input data menjadi vektor kata[8]. Selanjutnya, model menggunakan dua lapisan Conv1D untuk mengekstraksi fitur lokal dari teks. Pada lapisan pertama, Conv1D memiliki 32 filter dengan ukuran kernel 3 dan menggunakan fungsi aktivasi ReLU, diikuti oleh lapisan MaxPooling1D dengan ukuran pool 2 dan stride 2 untuk mengurangi dimensi fitur. Lapisan kedua Conv1D juga memiliki 32 filter, tetapi dengan ukuran kernel 4, diikuti oleh lapisan MaxPooling1D lainnya untuk lebih mengurangi dimensi fitur. Sebagai langkah pencegahan overfitting, lapisan SpatialDropout1D dengan rate 0.2 ditambahkan setelah lapisan konvolusi. Lapisan pertama BiGRU memiliki 64 unit dan mengolah data secara bidirectional, memproses informasi dari kedua arah (masa lalu dan masa depan). Lapisan ini diikuti oleh lapisan Dropout dengan rate 0.3 untuk mencegah overfitting. Lapisan kedua BiGRU memiliki 32 unit dan juga menggunakan Dropout dengan rate 0.2.

Model kemudian dilanjutkan dengan lapisan Flatten untuk mengubah data 2D menjadi 1D yang bisa diterima

oleh lapisan Dense. Lapisan Dense pertama memiliki 32 unit dan menggunakan aktivasi ReLU untuk menangkap hubungan non-linear dalam data. Di bagian akhir, lapisan output Dense digunakan dengan jumlah neuron sesuai dengan jumlah kelas (3 kelas dalam hal ini) dan fungsi aktivasi softmax untuk klasifikasi multi-kelas.

Model ini dikompilasi menggunakan Adam *optimizer* dengan learning rate 1e-4 dan categorical crossentropy sebagai fungsi loss untuk klasifikasi multi-kelas. Untuk mencegah overfitting selama pelatihan, EarlyStopping bisa diterapkan jika validation loss tidak menunjukkan perbaikan.



Gambar 5 Hybrid Deep Learning Arsitektur CNN-BiGRU

G. Evaluasi Performansi

Evaluasi kinerja dalam penelitian ini menggunakan confusion matrix, yang merupakan alat penting untuk menilai efektivitas model klasifikasi. Confusion matrix dapat memberikan gambaran seberapa baik model membedakan antar kelas dengan membandingkan antara hasil klasifikasi yang sebenarnya dan prediksi yang dihasilkan oleh model. Confusion matrix terdiri dari empat kombinasi hasil klasifikasi, yaitu *True Positive* (TP) yang menunjukkan data positif yang diprediksi dengan benar, *False Positive* (FP) yang menunjukkan data positif yang salah diprediksi, *True Negative* (TN) yang menunjukkan data negatif yang diprediksi dengan benar, dan *False Negative* (FN) yang menunjukkan data negatif yang salah diprediksi[14]. Nilai-nilai yang dihasilkan dari confusion matrix digunakan untuk mengevaluasi kinerja model klasifikasi melalui perhitungan Accuracy, Precision, Recall, dan F1-score, yang dapat diberi bobot. Keempat metrik tersebut dapat dihitung menggunakan rumus berikut:

$$Accuracy = \frac{(TP + TN)}{(TP + TN) + (TP + FP + FN + TN)} \quad (4)$$

$$Precision = \frac{TP}{(TP + FP)} \quad (5)$$

$$Recall = \frac{TP}{(TP + FN)} \quad (6)$$

$$F1 - Score = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (7)$$

IV. HASIL DAN PEMBAHASAN

Pada penelitian ini akan dilakukan beberapa skenario pengujian terhadap empat model klasifikasi yaitu CNN, BiGRU, *hybrid deep learning* CNN-BiGRU dan *hybrid deep learning* BiGRU-CNN. Skenario pertama dilakukan untuk menentukan model CNN dan BiGRU terbaik. Skenario kedua menerapkan Max-Feature pada baseline dari skenario pertama dengan baseline digunakan untuk ekstraksi fitur menggunakan TF-IDF. Eksperimen dilakukan dengan membandingkan hasil akurasi pada jumlah maksimum fitur TF-IDF sebesar 2500, 5000, 7500, 10000. Skenario ketiga menguji model baseline CNN dan BiGRU yang diperoleh dari skenario kedua, kemudian dikombinasikan dengan ekspansi fitur menggunakan Skenario keempat menguji efektivitas dari tiga *optimizer* populer yang sering digunakan dalam training model deep learning, yaitu Adam, Nadam, dan RMSprop.

A. Skenario 1

Skenario pertama bertujuan untuk menguji eksperimen dalam menentukan model baseline terbaik. Pengujian dilakukan dengan mencari nilai *split ratio* terbaik serta menentukan rasio pembagian data latih dan data uji. Model yang digunakan dalam pengujian ini adalah CNN dan BiGRU. Rasio pembagian data yang diuji meliputi 90:10, 80:20, dan 70:30. Pada rasio 90:10, 90% data digunakan untuk pelatihan dan 10% sisanya digunakan untuk pengujian, bisa dilihat pada Tabel 6.

Tabel 6 Hasil Uji Skenario Pertama

Splitting Ratio	Accuracy			
	CNN	BiGRU	CNN-BiGRU	BiGRU-CNN
70:30	74,46%	74,51%	73,69%	73,61%
80:20	74,56%	75,13%	73,85%	74,02%
90:10	74,74%	75,18%	73,83%	74,45%

B. Skenario 2

Skenario kedua adalah menerapkan hasil terbaik dari baseline dengan vektor fitur 2500, 5000, 7500, 10000. menunjukkan hasil akurasi terbaik menggunakan vektor dengan model CNN adalah 7500 dengan akurasi 79,98 % dan BiGRU 80,03 % untuk model hybrid CNN-BiGRU dan hybrid CNN-BiGRU menunjukkan hasil akurasi terbaik menggunakan vektor fitur 7500 masing masing model yaitu hybrid CNN-BiGRU 79,18 % dan hybrid CNN-BiGRU 80,30 % .Dengan hasil tersebut akan digunakan fitur terbaik untuk skenario pengujian selanjutnya pada table 7 bisa di lihat.

Tabel 7 Hasil Uji Skenario Kedua

Max Feature	Accuracy			
	CNN	BiGRU	CNN-BiGRU	BiGRU-CNN
2500	79,64 %	79,64 %	79,15 %	80,29 %
5000	79,76 %	79,79 %	78,79 %	80,10 %
7500	79,98 %	80,03 %	79,18 %	80,30 %
10000	79,85 %	79,68 %	78,54 %	80,17 %

C. Skenario 3

Skenario ketiga merupakan eksperimen yang dilaksanakan dengan mengaplikasikan ekspansi fitur pada model baseline

terbaik dari skenario sebelumnya. Ekspansi fitur ini dilakukan menggunakan FastText untuk menentukan kesamaan makna kata berdasarkan peringkat teratas pada korpus kesamaan kata yang telah dibangun. Pengujian dilakukan dengan memanfaatkan dua jenis korpus, yaitu tweet dan *IndoNews*. Evaluasi dilakukan pada peringkat teratas yang mencakup Top 1, Top 5, dan Top 10 pada table 8 bisa di lihat.

Tabel 8 Hasil Uji Skenario Ketiga

Model	Rank	Accuracy		
		Tweet	Indonews	Tweet + Indonews
CNN	Top 1	79,71 %	80,00 %	78,92 %
	Top 5	79,17 %	79,53 %	79,84 %
	Top 10	79,57 %	79,82 %	79,04 %
BiGRU	Top 1	79,71 %	80,10 %	79,69 %
	Top 5	79,02 %	79,98 %	79,62 %
	Top 10	79,23 %	79,64 %	79,19 %
CNN-BiGRU	Top 1	79,16 %	79,20 %	79,35 %
	Top 5	78,65 %	78,65 %	79,11 %
	Top 10	78,98 %	78,98 %	78,94 %
BiGRU-CNN	Top 1	79,78 %	80,33 %	80,37 %
	Top 5	79,98 %	80,30 %	79,75 %
	Top 10	79,78 %	79,30 %	79,02 %

D. Skenario 4

Pada scenario keempat tujuan utama adalah untuk membandingkan kinerja dari tiga *optimizer* dalam deep learning yaitu Adam, Nadam, dan RMSprop Ketiga *optimizer* ini memiliki karakteristik yang berbeda dalam cara mereka memperbarui bobot selama proses pelatihan, yang dapat memengaruhi kinerja model. Model yang digunakan untuk eksperimen ini adalah CNN, BiGRU, hybrid CNN-BiGRU dan hybrid BiGRU-CNN. Model ini diuji menggunakan peringkat Top 1 dengan korpus IndoNews, dan diterapkan algoritma optimasi seperti Adam. Nadam, RMSprop. Hasil bisa dilihat pada Tabel 9.

Tabel 9 Hasil Uji Skenario Keempat

Model	Optimizer	Accuracy Indonews Top 1
CNN	Adam	80,00 %
	Nadam	80,07 %
	RMSprop	79,07 %
BiGRU	Adam	80,10 %
	Nadam	79,85 %
	RMSprop	79,86 %
CNN-BiGRU	Adam	79,20 %
	Nadam	79, 24 %
	RMSprop	80,57 %
BiGRU-CNN	Adam	80,33 %
	Nadam	80,24 %
	RMSprop	80,65 %

E. Analisis Hasil Pengujian

Penelitian ini bertujuan untuk mengembangkan model deteksi depresi di Twitter menggunakan metode CNN-BiGRU yang dipadukan dengan ekspansi fitur FastText. Hasil pengujian menunjukkan bahwa kombinasi CNN dan BiGRU memberikan kontribusi signifikan terhadap peningkatan akurasi deteksi depresi pada dataset Twitter[8].

Pada skenario pertama, hasil pengujian menunjukkan bahwa model baseline terbaik diperoleh dengan splitting ratio 90:10, di mana akurasi model CNN mencapai 74,74%, model BiGRU sebesar 75,18%, model CNN-BiGRU mencapai 73,83%, dan model BiGRU-CNN mencapai 74,45%.

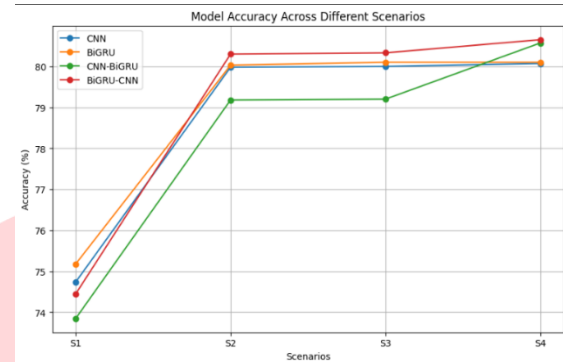
Pada skenario kedua, penggunaan vektor fitur sebanyak 7.500 meningkatkan akurasi pada model. Pada skenario ini, model CNN mencapai akurasi 79,98 % model BiGRU mencapai 80,03 %, model CNN-BiGRU mencatatkan akurasi 79,18 %, dan model BiGRU-CNN mencapai 80,30 %. Peningkatan akurasi dibandingkan baseline pada skenario pertama adalah sebesar untuk model CNN, 5,24% untuk model BiGRU, 4,85% untuk model CNN-BiGRU, dan 5,35% untuk model BiGRU-CNN 5,85%. Penghapusan fitur sangat diperlukan ketika vektor fitur yang digunakan sangat besar, seperti pada skenario kedua dengan 7.500 fitur. Meskipun ada peningkatan akurasi, jumlah fitur yang berlebihan dapat menyebabkan masalah seperti *overfitting*, kompleksitas komputasi yang lebih tinggi, dan penurunan kemampuan interpretasi model. Selain itu, tidak semua fitur memberikan kontribusi yang signifikan terhadap prediksi model mekanisme fitur yang saya pakai adalah sebagai berikut yaitu seleksi fitur berbasis jumlah (*feature selection based on quantity*) atau penyaringan fitur secara bertahap.

Pada skenario ketiga, penggabungan FastText menunjukkan peningkatan yang lebih signifikan. Penggunaan FastText sebagai metode ekspansi fitur membantu model menangkap konteks dan hubungan antar kata dalam teks. Pendekatan ini memperluas representasi fitur, membantu model menangani variasi bahasa dan kompleksitas data teks, Peningkatan akurasi lebih lanjut dicapai dengan penggunaan korpus peringkat teratas (*Indonews* top 1). Model CNN mencatat akurasi 80,00 %, model BiGRU mencapai 80,10 %, model CNN-BiGRU mencatatkan akurasi 79,20 %, dan model BiGRU-CNN mencapai 80,33 %. Dibandingkan dengan baseline pada skenario pertama, peningkatan akurasi sebesar untuk model CNN, 5,26% untuk model BiGRU, 4,92% untuk model CNN-BiGRU, dan 5,37% untuk model BiGRU-CNN 5,88%.

Dalam Skenario Keempat, perbandingan antara tiga *optimizer* populer, Adam, Nadam, dan RMSprop, menunjukkan bahwa RMSprop memberikan hasil terbaik pada model BiGRU-CNN, dengan akurasi 80,65%. Hal ini menegaskan bahwa pemilihan *optimizer* memiliki dampak signifikan pada performa model, meskipun perbedaan antara *optimizer* lainnya tidak terlalu besar. Hal ini menunjukkan bahwa meskipun optimisasi dapat meningkatkan akurasi, faktor-faktor lain seperti parameter

model dan dataset tetap memainkan peran yang sangat penting.

Secara keseluruhan, analisis ini menunjukkan bahwa penggunaan teknik CNN-BiGRU dengan ekspansi fitur FastText memberikan hasil yang baik dalam mendeteksi depresi di Twitter, meskipun masih dipengaruhi oleh berbagai faktor seperti jumlah fitur, jenis data, dan pemilihan *optimizer*.



Gambar 6 Peningkatan akurasi semua skenario

Tabel 10 Peningkatan akurasi semua skenario

	Akurasi (%)			
	S1	S2	S3	S4
CNN	74.74	79.98	80.00	80.07
BiGRU	75.18	80.03	80.10	80.10
CNN-BiGRU	73.85	79.18	79.20	80.57
BiGRU-CNN	74.45	80.30	80.33	80.65

V. KESIMPULAN

Penelitian ini mengembangkan model untuk mendeteksi depresi di Twitter dengan menggabungkan CNN- BiGRU dan ekspansi fitur FastText. Hasil pengujian menunjukkan bahwa kombinasi ini sangat efektif dalam meningkatkan akurasi deteksi depresi. CNN berhasil menangkap pola lokal dalam teks, sementara BiGRU membantu memahami urutan kata dengan memproses data dari kedua arah. Penggunaan FastText memperkaya representasi kata, sementara penambahan fitur hingga 7.500 vektor juga meningkatkan akurasi meskipun berisiko menyebabkan *overfitting*. Optimizer RMSprop memberikan hasil terbaik pada model BiGRU-CNN, yang mencatatkan akurasi tertinggi sebesar 80,65%. Secara keseluruhan, performa model sangat dipengaruhi oleh jumlah fitur, jenis data, dan pemilihan *optimizer*.

Untuk penelitian selanjutnya, penting untuk fokus pada peningkatan kualitas data, serta mengeksplorasi model-model lain seperti Transformer atau BERT yang mungkin lebih efektif dan percobaan dengan *optimizer* lainnya. Pengembangan teknik seleksi fitur yang lebih baik dan eksperimen dengan *optimizer* lain juga dapat memperbaiki hasil, sementara pengembangan model multibahasa akan memperluas kemampuan deteksi depresi ini.

REFERENSI

- [1] Lin, L. Y., Sidani, J. E., Shensa, A., Radovic, A., Miller, E., Colditz, J. B., ... & Primack, B. A. (2016). *Association Between Social Media Use and*

Depression Among US Young Adults. Depression and Anxiety, 33(4), 323-331.

- [2] Nugroho, K. S., Akbar, I., & Suksmawati, A. N. (2023). Deteksi Depresi dan Kecemasan Pengguna Twitter Menggunakan Bidirectional LSTM. *arXiv preprint arXiv:2301.04521*.
- [3] H. Zogan, I. Razzak, S. Jameel, and G. Xu. (2021) "DepressionNet: Learning Multi-modalities with User Post Summarization for Depression Detection on Social Media," in *SIGIR 2021 - Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, Association for Computing Machinery, Inc, Jul., pp. 133–142. doi: 10.1145/3404835.3462938.
- [4] R. A. Rudiyanto and E. B. Setiawan. (2024) "Sentiment Analysis Using Convolutional Neural Network (CNN) and Particle Swarm Optimization on Twitter," *JITK (Jurnal Ilmu Pengetahuan dan Teknologi Komputer)*, vol. 9, no. 2, pp. 188–195, Feb., doi: 10.33480/jitk.v9i2.5201.
- [5] Sofia, R. N., & Supriyadi, D. (2021). Komparasi metode machine learning dan deep learning untuk deteksi emosi pada text di sosial media. *JUPITER: Jurnal Penelitian Ilmu dan Teknologi Komputer*, 13(2), 130-139.
- [6] Merinda Lestandy, Amrul Faruq, Adhi Nugraha, and Abdurrahim. (2024) *Analyzing Reddit Data: Hybrid Model for Depression Sentiment using FastText Embedding*, *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 8, no. 2, pp. 288–297, Apr., doi: 10.29207/resti.v8i2.5641.
- [7] Aisyiyah, S., & Maharani, W. (2023). Analisis Berbasis Emosional pada Depresi di Media Sosial Menggunakan Pendekatan Convolutional Neural Network. *eProceedings of Engineering*, 10(2).
- [8] R. Pabian and E. B. Setiawan. (2025) "Optimizing Learning Rates and Feature Expansion with FastText in Hybrid CNN-BiGRU for Indonesian Cyberbullying Detection on X," in *International Conference on Advancement in Data Science, E-learning and Information System (ICADEIS)*, 2025, pp. 1–8. doi: 10.1109/ICADEIS65852.2025.10933150.
- [9] E. J. Yeun, Y. M. Kwon, and J. A. Kim (2012), *Psychometric testing of the Depressive Cognition Scale in Korean adults, Applied Nursing Research*, vol. 25, no. 4, pp. 264–270, Nov. doi: 10.1016/j.apnr.2011.04.003.
- [10] Zogan, H., Razzak, I., Wang, X., Jameel, S., & Xu, G. (2022). *Explainable depression detection with multi-aspect features using a hybrid deep learning model on social media. World Wide Web*, 25(1), 281- 304.
- [11] Setiawan, E. B., Widyantoro, D. H., & Surendro, K. (2016, October). Feature expansion using word embedding for tweet topic classification. In *2016 10th International Conference on Telecommunication Systems Services and Applications (TSSA)* (pp. 1-5). IEEE.
- [12] V. S and J. R. (2016) *Text Mining: open Source Tokenization Tools – An Analysis*, *Advanced Computational Intelligence: An International Journal (ACII)*, vol. 3, no. 1, pp. 37–47, Jan., doi: 10.5121/acii.2016.3104.
- [13] Merinda Lestandy, Amrul Faruq, Adhi Nugraha, and Abdurrahim. (2024) *Analyzing Reddit Data: Hybrid Model for Depression Sentiment using FastText Embedding*, *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 8, no. 2, pp. 288–297, Apr., doi: 10.29207/resti.v8i2.5641.
- [14] I. A. Asqolani and E. B. Setiawan, "A Hybrid deep learning Approach Leveraging Word2Vec Feature Expansion for Cyberbullying Detection in Indonesian Twitter," *Ingenierie des Systemes d'Information*, vol. 28, no. 4, pp. 887–895, Aug. 2023, doi: 10.18280/isi.280410.
- [15] H. Kour and M. K. Gupta.(2022) *An hybrid deep learning approach for depression prediction from user tweets using feature-rich CNN and bi-directional LSTM*, *Multimed Tools Appl*, vol. 81, no. 17, pp. 23649–23685, Jul., doi: 10.1007/s11042-022-12648-y.
- [16] D. She and M. Jia. (2024) *A BiGRU method for remaining useful life prediction of machinery, Measurement (Lond)*, vol. 167, Jan. 1, doi: 10.1016/j.measurement.2020.108277.