

Identifikasi Cuitan Depresi pada Mahasiswa selama Melaksanakan Kuliah Daring di masa Pandemi COVID-19

Tugas Akhir
diajukan untuk memenuhi salah satu syarat
memperoleh gelar sarjana
pada Program Studi S1 Informatika
Fakultas Informatika
Universitas Telkom

1301172540
Naura Khairunnisa



Program Studi Sarjana Informatika
Fakultas Informatika
Universitas Telkom
Bandung
2021

LEMBAR PENGESAHAN

**Identifikasi Cuitan Depresi pada Mahasiswa selama Melaksanakan Kuliah Daring
di masa Pandemi COVID-19**

*Identification of Depression Tweets in College Students during Online Lectures during
the COVID-19 Pandemic*

NIM : 1301172540

Naura Khairunnisa

Tugas akhir ini telah diterima dan disahkan untuk memenuhi sebagian syarat memperoleh
gelar pada Program Studi Sarjana Informatika

Fakultas Informatika

Universitas Telkom

Bandung, 16 September 2021

Menyetujui

Pembimbing I,



Dade Nurjanah, S.T., M.T., Ph.D.

NIP: 93670013-1

Pembimbing II,



Hani Nurrahmi, S.Kom., M.Kom

NIP: 15900040-1

Ketua Prodi
Sarjana Informatika,

Dr. Erwin Budi Setiawan, S.Si., M.T.

NIP : 00760045

LEMBAR PERNYATAAN

Dengan ini saya, Naura Khairunnisa, menyatakan sesungguhnya bahwa Tugas Akhir saya dengan judul Identifikasi *Tweets* Depresi pada Mahasiswa selama Melaksanakan Kuliah *Online* di masa Pandemi *COVID-19* beserta dengan seluruh isinya adalah merupakan hasil karya sendiri, dan saya tidak melakukan penjiplakan yang tidak sesuai dengan etika keilmuan yang berlaku dalam masyarakat keilmuan. Saya siap menanggung resiko/sanksi yang diberikan jika di kemudian hari ditemukan pelanggaran terhadap etika keilmuan dalam Laporan TA atau jika ada klaim dari pihak lain terhadap keaslian karya.

Bandung, 16 September 2021

Yang Menyatakan



Naura Khairunnisa

Identifikasi *Tweets* Depresi pada Mahasiswa selama Melaksanakan Kuliah Online di masa Pandemi COVID-19

Naura Khairunnisa¹, Dade Nurjanah, S.T., M.T., Ph.D. ², Hani Nurrahmi, S.Kom., M.Kom ³

^{1,2,3}Fakultas Informatika, Universitas Telkom, Bandung

¹naurakh@student.telkomuniversity.ac.id, ²dadenurjanah@telkomuniversity.ac.id,

³haninurrahmi@telkomuniversity.ac.id

Abstrak

Adanya pandemi COVID-19 ini menyebabkan sistem kegiatan belajar mengajar berubah menjadi dilakukan secara daring. Salah satu hal yang menjadi perhatian ialah bagaimana sekolah daring dapat mempengaruhi kesehatan mental. Dalam mengeluhkan kesulitannya, pelajar khususnya mahasiswa turut mengungkapkannya melalui media sosial Twitter. Keluhan yang diungkapkan dalam bentuk kicauan ini dapat berpotensi menjadi kicauan atau *tweet* depresi. Untuk mengetahui apakah mahasiswa tersebut berpotensi depresi selama kuliah *online*, dapat dilakukan identifikasi depresi melalui *tweets*nya dalam tiga bulan. Dataset yang akan digunakan merupakan hasil *crawling tweets* dari kata kunci “kuliah *online*” dan dibuat sejak bulan Oktober 2020 menggunakan *library snsrape*. Sebelum proses klasifikasi, dilakukan tahapan *preprocessing* yang meliputi *data cleaning*, *case folding*, tokenisasi, *stop words removal*, normalisasi, dan *stemming*. Selanjutnya, untuk mengidentifikasi depresi, penelitian ini menggunakan metode *word embedding Fast Text* dan pengklasifikasi LSTM.

1. Pendahuluan

Latar Belakang

Adanya pandemi COVID-19 ini memberikan dampak yang signifikan terhadap berbagai macam sektor di Indonesia, tidak terkecuali sektor pendidikan. Dampak terhadap sektor pendidikan yang diakibatkan oleh pandemi COVID-19 adalah kegiatan belajar mengajar secara langsung menjadi dibatasi. Pemerintah, khususnya Kemendikbud mengumumkan bahwa seluruh kegiatan belajar mengajar dialihkan menjadi sistem pembelajaran *online* [1]. Arah ini berlaku untuk seluruh tingkatan pendidikan, mulai dari PAUD hingga perguruan tinggi.

Salah satu hal yang menjadi perhatian ialah bagaimana sekolah daring dapat mempengaruhi kesehatan mental. Penelitian yang dilakukan terhadap mahasiswa Universitas Telkom dan UIN SGD Bandung menunjukkan bahwa sekitar 59,5% mahasiswa merasakan keberatan dengan tugas yang diberikan dosen selama kuliah *online* dan juga 60% mahasiswa merasakan kesulitan tidur yang diakibatkan kuliah *online* [2]. Hal tersebut membuktikan bahwa adanya pandemi ini dapat mempengaruhi kesehatan mental pada mahasiswa. Sayangnya, masih ada mahasiswa yang tidak mengetahui ilmu tentang kesehatan mental. Wang dkk. [3] menemukan bahwa lebih dari 40% mahasiswa membutuhkan pengetahuan psikologis, dan 87,2% mahasiswa merasa perlu memahami gejala umum kecemasan dan depresi.

Media sosial, terutama Twitter, telah menjadi salah satu media untuk menceritakan hal yang dialami atau dirasakan oleh pengguna aplikasi tersebut, tidak terkecuali mahasiswa. Terdapat penelitian sebelumnya tentang identifikasi kicauan yang berkaitan dengan depresi pada media sosial Twitter. Santos dkk. [4] melakukan penelitian untuk mendeteksi isu kesehatan mental di Brazil dengan menggunakan data yang didapatkan dari kicauan pengguna Twitter di Brazil yang sudah didiagnosa oleh praktisi kesehatan mental. Selain itu, Chomutare dkk. [5] menilai kinerja pengklasifikasi untuk mendeteksi pasien obesitas dan diabetes yang beresiko mengalami depresi. Dari penelitian-penelitian tersebut, didapatkan hipotesis bahwa kondisi kesehatan mental seseorang dapat dianalisis dari unggahan statusnya di media sosial.

Topik dan Batasannya

Tugas akhir ini difokuskan pada identifikasi kicauan yang berpotensi depresi dengan menggunakan data yang bersumber dari kicauan atau *tweets* mahasiswa tentang kuliah *online* selama pandemi COVID-19. Penelitian ini mengusulkan untuk menggunakan metode LSTM dan FastText. Pemilihan metode didasarkan pada beberapa penelitian terdahulu sebagai referensi. Pemilihan metode didasarkan pada beberapa penelitian terdahulu sebagai referensi. Pada penelitian mengenai identifikasi kicauan tentang kesehatan [6], dilakukan perbandingan terhadap beberapa metode yang berbeda dan menghasilkan LSTM sebagai metode terbaik dengan akurasi 0,815. Wang dkk. [7] juga membuktikan LSTM menghasilkan akurasi terbaik dibandingkan Naïve Bayes dan Extreme Learning Maching (ELM) dengan akurasi 0,859. Sedangkan, penelitian yang dilakukan Tsugawa dkk. [8] dalam mengenali depresi berdasarkan aktivitas Twitter hanya menghasilkan akurasi tertinggi sebesar 66% berdasarkan

fitur yang berbeda menggunakan metode SVM. Adapun penelitian Joulin dkk. terhadap prediksi tagar yang membuktikan bahwa pengklasifikasian teks dengan menggunakan FastText dapat dilakukan dengan cepat dibandingkan dengan model Tagspace. Alessa dkk. [11] juga menyebutkan hasil prediksi dari FastText dapat dikatakan efisien karena menghasilkan akurasi F-measure hingga 89,9% dengan menggunakan berbagai fitur yang berbeda.

Pada tugas akhir ini, topik yang akan dibahas adalah bagaimana performansi metode FastText dan LSTM dalam mendeteksi depresi melalui kicauan mengenai kuliah *online*. Batasan masalah pada tugas akhir ini adalah sebagai berikut: Pertama, kicauan atau *tweets* yang dapat diidentifikasi sebagai *tweet* depresi adalah kicauan yang mengandung kata yang berhubungan dengan gejala awal depresi. Kedua, pengguna Twitter yang kicauannya akan digunakan dalam penelitian adalah mahasiswa yang melaksanakan kuliah *online*. Ketiga, penelitian ini tidak menangani cuitan depresi yang implisit, seperti “Saya telah didiagnosa depresi”.

Tujuan

Tujuan dari penelitian ini adalah untuk mengidentifikasi adanya indikasi depresi pada *tweets* mahasiswa mengenai keluhan selama kuliah *online* disaat pandemi COVID-19. Penelitian ini diharapkan dapat membantu mahasiswa pengguna media sosial untuk dapat meningkatkan kesadaran akan adanya gejala depresi.

Organisasi Tulisan

Struktur penulisan dari tugas akhir ini disusun sebagai berikut: Bagian pertama berisi pendahuluan terkait tugas akhir ini. Bagian kedua menjelaskan studi yang terkait dengan tugas akhir ini. Bagian ketiga akan menjelaskan pemodelan dan performansi dari sistem yang dibangun. Bagian keempat menjelaskan hasil dan evaluasi hasil pengujian yang telah dilakukan pada bagian ketiga. Kemudian, pada bagian terakhir menjelaskan kesimpulan dan saran berdasarkan hasil pengujian yang dilakukan pada tugas akhir ini.

2. Studi Terkait

Kuliah *online* telah menjadi hal yang lumrah di Indonesia bagi sebagian besar mahasiswa sejak adanya pandemi COVID-19 mengacu pada aturan Kemdikbud tahun 2020 [9]. Kuliah *online* dilakukan untuk mencegah adanya kemungkinan penyebaran virus *corona* di area kampus. Adanya kuliah *online* ini tentu memberikan dampak positif maupun negatif. Salah satu dampak negatifnya adalah kuliah *online* yang mempengaruhi kesehatan mental mahasiswa [10].

Media sosial diketahui telah menjadi salah satu sarana untuk menyalurkan apa yang sedang dirasakan, termasuk masalah dalam kehidupan sehari-hari. Hal ini yang menjadikan latar belakang pada beberapa penelitian untuk mengidentifikasi depresi pada media sosial berdasarkan postingannya dengan menggunakan berbagai macam metode. Salah satu media sosial yang banyak digunakan untuk mendeteksi depresi adalah Twitter [11]–[13]. Penelitian oleh Deshpande dkk. [11] mengaplikasikan *Natural Language Preprocessing* dengan metode Naïve Bayes dan SVM pada postingan Twitter untuk melakukan analisis emosi yang berfokus pada depresi. Deshpande dkk. mengklasifikasi *tweet* sebagai netral atau negatif, berdasarkan daftar kata yang dikuratori untuk mendeteksi kecenderungan depresi. Penelitian [11] berhasil menerapkan AI emosi berbasis teks untuk mendeteksi depresi dengan akurasi sebesar 83%. Asad dkk. [12] menganalisa postingan pengguna Twitter untuk mengetahui tingkatan depresi menggunakan pembobotan TF-IDF dan klasifikasi Naïve Bayes. Dengan akurasi 74% dan precision 100%, model yang dibuat oleh Asad dkk. dapat bermanfaat untuk setiap individu yang sedang mengalami depresi melalui unggahan media sosial Twitter dan Facebook. Salah satu penelitian di India, melakukan analisa deteksi depresi pada *tweet* dari 100 pengikut sebuah forum MS India [13]. Selain menganalisa *tweets*, penelitian Kumar dkk. [13] mempertimbangkan 5 aspek, yaitu; word, timing, frequency, sentiment, contrast, dan klasifikasi yang digunakan sebagai perbandingan yaitu Naïve Bayes, Random Forest, Gradient Boosting, dan Ensemble Vote Classifier. Penelitian tersebut berhasil mendapatkan akurasi sebesar 85% untuk dapat memprediksi adanya depresi kecemasan pada 100 akun pengikut forum MS India.

Figueredo dkk. [14], menggunakan metode preprocessing sebelum melakukan deteksi depresi pada dataset yang diambil dari media sosial reddit. Adapun penelitian mengenai deteksi emosi yang dilakukan oleh Riza dkk. [15], yang memanfaatkan word embedding sebelum melakukan klasifikasi. Penelitian [15] menunjukkan hasil akurasi sebesar 0.731458 menggunakan word embedding FastText dan Word2Vec. Uddin dkk. [16] melakukan analisis menggunakan pendekatan LSTM pada dataset berukuran kecil yaitu total 1968 *tweets*. Dari 10 percobaan yang dilakukan, penelitian [16] akurasi tertinggi yaitu sebesar 86,3% dengan penyetalan LSTM ukuran 128, batch size = 25, learning rate = 0,0001 dan epoch = 20.

Long-Short Term Memory

LSTM atau *Long Short-Term Memory* merupakan jenis *Recurrent Neural Network* (RNN) yang mampu mempelajari ketergantungan urutan dalam masalah prediksi urutan [17]. Untuk memperoleh informasi yang

melewati LSTM, digunakan tiga jenis *gate*, yaitu *forget gate*, *input gate*, dan *output gate* [18]. Semua jenis gerbang yang berbeda ini dapat menyampaikan informasi yang berguna tentang keadaan net atau jaring saat ini. Misalnya, *input gate* (*output gate*) dapat menggunakan masukan dari sel memori lain untuk memutuskan apakah akan menyimpan (mengakses) informasi tertentu dalam sel memorinya [19]. Untuk setiap *gate*, LSTM menggunakan persamaan sebagai berikut :

$$f_t = \sigma(W_f x_t + U_f h_{t-1} + b_f) \quad (1)$$

$$i_t = \sigma(W_i x_t + U_i h_{t-1} + b_i) \quad (2)$$

$$o_t = \sigma(W_o x_t + U_o h_{t-1} + b_o) \quad (3)$$

$$u_t = \tanh(W_u x_t + U_u h_{t-1} + b_u) \quad (4)$$

$$c_t = f_t \odot c_{t-1} + i_t \odot u_t \quad (5)$$

$$h_t = o_t \odot \tanh(c_t) \quad (6)$$

f_t merepresentasikan gerbang, i_t merepresentasikan input *gate*, dan output *gate* nya adalah o_t . $\odot$ mewakili perkalian elemen langkah demi langkah. Nilai W_f , W_i , W_o , dan W_u merupakan nilai weight untuk setiap *gate*. $U_{h_{t-1}}$ merupakan output dari blok LSTM sebelumnya. b merupakan bias untuk masing-masing gerbang. Fungsi sigmoid digunakan untuk setiap *gate*. u_t Merepresentasikan fungsi *tanh*. Informasi sel informasi ada di dalam c_t , dan h_t merepresentasikan hidden state. LSTM menyimpan informasi dari sel sebelumnya.

Metode LSTM telah digunakan pada beberapa penelitian untuk klasifikasi teks. Mumu dkk. [20] menggunakan LSTM untuk mendeteksi depresi dengan komposisi dataset 5053 berlabel depresi dan 2110 berlabel non-depresi dan menghasilkan akurasi sebesar 81.05%. Untuk penelitian sentiment analisis yang dilakukan oleh Miedema dkk. [21], nilai akurasinya adalah 0.8674 dengan loss 0.4366. Loss yang cukup tinggi ini disebabkan karena adanya overfitting pada saat training model. Selanjutnya, ada penelitian untuk mendeteksi adanya sarkasme yang dilakukan oleh Khotijah dkk. [22], dan berhasil mendapatkan akurasi sebesar 88.33%. Adapun penelitian mengenai sentiment analisis Bahasa Indonesia [23], yang mendapatkan akurasi sebesar 95.24 dengan menggunakan kombinasi RNN dan LSTM.

FastText

Word embedding merupakan proses untuk memberikan representasi pada setiap kata yang ada pada dataset. Jenis word embedding yang akan digunakan ialah FastText. FastText merupakan library yang dirancang oleh Facebook untuk membantu membangun solusi skalabel dalam representasi dan klasifikasi teks yang menggunakan vektor [24].

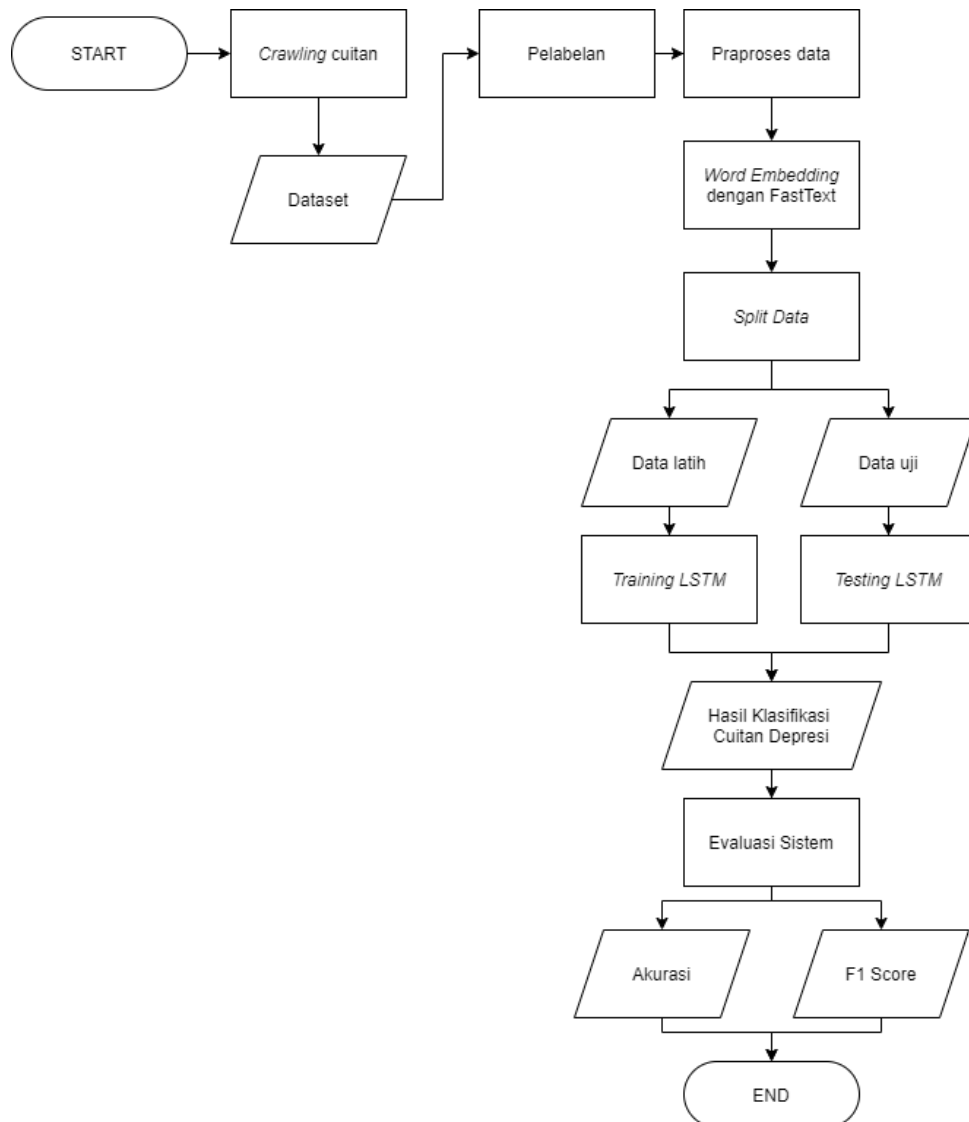
Salah satu kelebihan dari FastText diantaranya adalah waktu klasifikasi yang singkat dibanding metode lainnya. FastText merepresentasikan teks dengan vektor berdimensi rendah yang diperoleh dengan menjumlahkan vektor yang sesuai dengan kata-kata yang muncul dalam teks, lalu vektor berdimensi rendah dikaitkan dengan setiap kata dari kosakata [24]. Hal ini mengurangi kompleksitas waktu pada FastText, tanpa mengabaikan informasi tentang kata-kata yang dipelajari untuk satu kategori untuk digunakan oleh kategori lain.

Joulin dkk. [25] menunjukkan bahwa FastText hanya membutuhkan waktu 1 menit 29 detik dengan nilai akurasi 41,1, sedangkan model Tagspace membutuhkan waktu 15 jam dengan nilai akurasi 35,6. Penelitian yang dilakukan oleh Alessa dkk. [26] melakukan klasifikasi teks dari kicauan yang berhubungan dengan flu. Dengan total 10.592 *tweets*, proses klasifikasi pada penelitian [26] hanya membutuhkan waktu 21,5 detik dengan akurasi F-measure sebesar 89,9%.

Dengan penjelasan tersebut, dan dari penelitian yang dilakukan Joulin dkk. [25] dan Alessa dkk. [26], diketahui bahwa FastText dapat melakukan klasifikasi pada dataset yang besar dengan waktu yang cukup singkat dibandingkan dengan metode yang sudah ada sebelumnya.

3. Perancangan Sistem

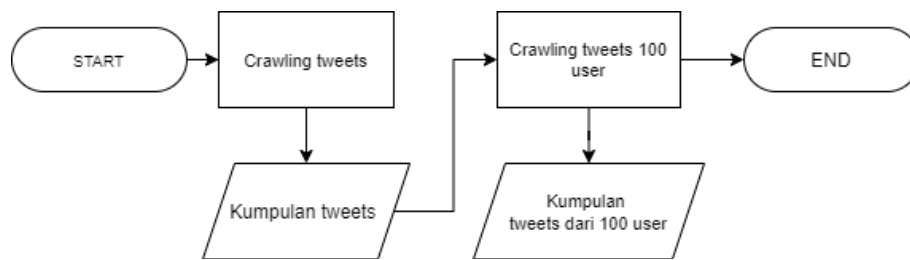
Proses yang akan dilakukan pada penelitian ini digambarkan dengan diagram alir berikut:



Gambar 3.1 Diagram alur proses penelitian.

Dataset

Metode yang digunakan untuk mendapatkan dataset yaitu crawling atau scrapping *tweets* dengan kata kunci “kuliah *online*” menggunakan *tool* *snsraper*. Kata kunci “kuliah *online*” dipilih dikarenakan pemilihan kata “kuliah” dan “*online*” lebih familiar dikalangan mahasiswa dibanding dengan istilah “kuliah daring”. Pertama, *tweets* yang diambil merupakan *tweets* yang dibuat pada bulan Oktober 2020 atau masa pertengahan semester pada umumnya. Dari *tweet* yang telah diambil dengan kata kunci “kuliah *online*”, dipilih 100 *user* yang menyebutkan bahwa mahasiswa tersebut mengeluhkan kuliah *online*, untuk kemudian diambil *tweetnya* setiap *user* selama tiga bulan mulai bulan Oktober hingga Desember tahun 2020. *Tweets* dari 100 *user* tidak semata-mata semuanya berkaitan dengan kuliah daring, tetapi tergantung apa yang dibicarakan oleh *user* selama tiga bulan masa kuliah daring. Total *tweets* yang didapatkan sekitar 11156 *tweets* dari 100 *user*.

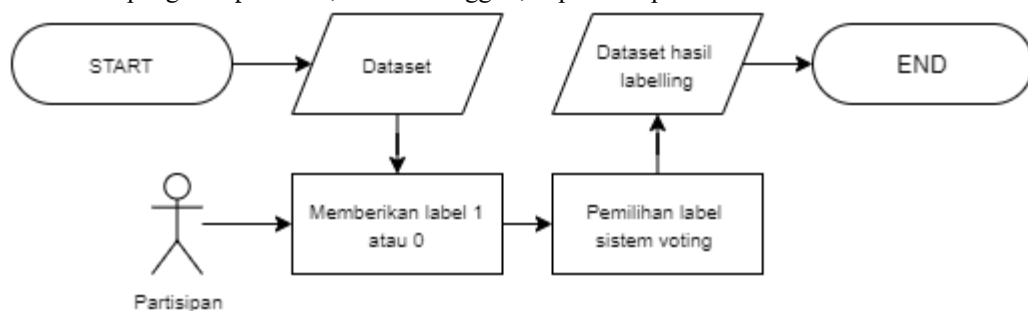


Gambar 3.2 Diagram alir pengambilan dataset.

Pelabelan

Labelling merupakan proses memberikan label pada setiap *tweet* yang ada pada dataset. Pemberian label berupa '1' untuk *tweet* depresi dan '0' untuk *tweet* non-depresi. Proses pelabelan dilakukan secara manual dengan membaca satu persatu data yang ada. Pelabelan dilakukan oleh tiga orang, dengan syarat yang mengacu pada gejala depresi pada remaja menurut buku DSM-5 yang telah dirangkum pada artikel [27]. Referensi gejala depresi tersebut ditulis menggunakan Bahasa Inggris, sehingga persyaratan diterjemahkan ke dalam Bahasa Indonesia. Syarat pemberian label yaitu adanya kemunculan kata atau ekspresi yang berkaitan dengan gejala depresi pada remaja sebagai berikut:

- Kemarahan atau permusuhan
Contoh : Ungkapan kata-kata kasar dan umpatan, 'benci', 'makin hari makin emosi'
- Perubahan kebiasaan makan atau tidur
Contoh : 'begadang', 'lupa makan', 'jam tidur berantakan'
- Kelelahan atau kekurangan energi
Contoh : 'capek', 'lemes', 'lelah'
- Keputusasaan
Contoh : 'tidak sanggup', 'aku menyerah', 'tidak kuat'
- Perasaan bersalah atau tidak berharga
Contoh : 'gak guna', 'insecure', 'kecewa sama diri sendiri'
- Prestasi sekolah yang buruk
Contoh : 'nilai turun', 'tugas keteteran', 'jadwal berantakan'
- Kurang motivasi
Contoh : 'males', 'udah gak ada semangat', 'gak nafsu nugas'
- Sulit berkonsentrasi
Contoh : 'mumet banget', 'overthinking', 'jadi ga konsen'
- Sering menangis
Contoh : 'pagi-pagi udah nangis', 'capek nangis', 'nangis tiap malem'
- Kegelisahan
Contoh : 'serangan panik', 'takut gagal', 'ga bisa tenang', 'anxiety'
- Agitasi
Contoh : 'pening banget', 'sakit kepala', 'mual'
- Sakit atau nyeri yang tidak dapat dijelaskan
Contoh : 'badan capek padahal gak ngapa-ngapain', 'sesek'
- Pikiran tentang kematian atau bunuh diri (dengan atau tanpa rencana)
Contoh : 'pengen cepet mati', 'mau meninggal', 'capek hidup'



Gambar 3.2 Diagram alir pemberian label pada dataset.

Berdasarkan hasil penggabungan pelabelan dari tiga orang, didapatkan total 608 *tweet* yang terindikasi depresi (1) dan 10.538 *tweet* untuk label non-depresi (0). Berikut adalah contoh *tweet* yang sudah diberi label :

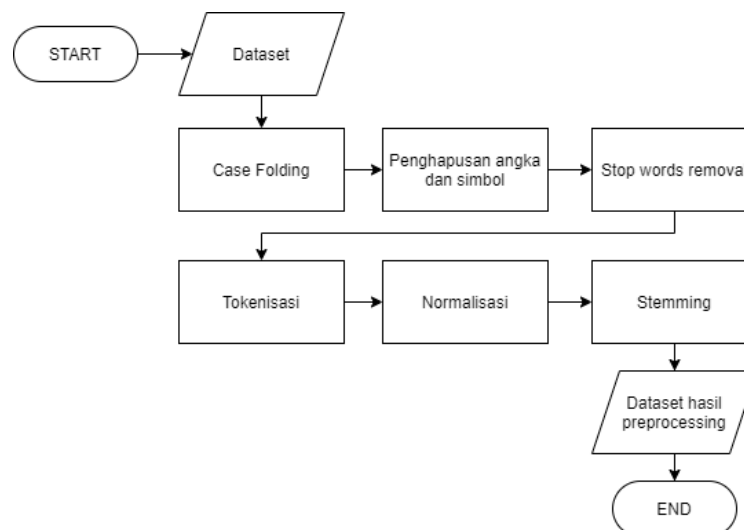
Tabel 1. Contoh *tweet* pada dataset.

No.	<i>Tweet</i>	Label
1.	Rasanya gada semangat hidup bgt ya allahhhhhhhhhh	1
2.	Tiap hari ngerasa cape pdhl gk lakuin aktifitas berat :'''	1
3.	kok orang-orang bisa ya balikan sama mantan..	0
4.	Maghriban dulu yuu	0

Preprocessing

Preprocessing merupakan langkah yang dilakukan untuk mempersiapkan data untuk meningkatkan hasil akurasi klasifikasi.

- Case Folding**
Case folding merupakan proses untuk menyamaratakan semua huruf yang ada pada dataset menjadi huruf kecil atau *lower case*.
- Penghapusan angka dan simbol**
Penghapusan angka dan simbol (*punctuation*) dilakukan agar dataset hanya terdiri dari huruf saja.
- Stop Words Removal**
Stop words removal merupakan proses untuk membuang kata-kata yang umum atau dianggap tidak penting dalam pemrosesan data. Misalnya, dalam Bahasa Indonesia ada kata-kata 'di', 'ke', 'yang', 'atau', 'dari', dan lainnya.
- Tokenisasi**
Tokenisasi merupakan proses pemecahan teks pada dataset menjadi bagian yang lebih kecil, baik dalam bentuk kata ataupun sebuah kalimat yang disebut token.
- Normalisasi**
Normalisasi merupakan proses untuk menyamaratakan kata yang berbeda tetapi memiliki arti yang sama, seperti 'aku', 'gue', 'gw', semuanya disamakan menjadi 'saya'.
- Stemming**
Stemming adalah proses untuk mengembalikan sebuah kata yang memiliki imbuhan menjadi kata dasar. Contohnya 'membanggakan' menjadi 'bangga', pertumbuhan menjadi 'tumbuh'. Proses ini menggunakan *library* sastrawi [28], yaitu library untuk stemming Bahasa Indonesia.



Gambar 3.3 Diagram alir preprocessing pada dataset.

Word Embedding FastText

Pada penelitian ini, FastText yang digunakan adalah library dari Gensim Fasttext. Gensim adalah library python bersifat open-source untuk merepresentasikan dokumen sebagai vektor semantik yang efisien [29]. Gensim ini menggunakan algoritma *unsupervised*, sehingga hanya perlu dokumen yang akan diubah menjadi

vector tanpa input dari manusia [29]. Untuk merepresentasikan dokumen pada penelitian ini, Gensim FastText hanya memerlukan 4 detik untuk merepresentasikan dokumen menjadi vektor.

Random Resampling

Dataset terdiri atas 608 *tweet* yang terindikasi depresi (1) dan 10.538 *tweet* untuk label non-depresi (0). Hal ini menyebabkan adanya ketidakseimbangan data atau *imbalanced dataset*. Untuk mengatasi hal ini, teknik yang dilakukan adalah teknik random resampling. *Random resampling* adalah salah satu teknik untuk mengatasi *imbalanced dataset* dengan cara mengambil sampel ulang secara acak [30]. Jenis random resampling yang digunakan ialah *over sampling*, yaitu memperbanyak sample dari kelas yang minoritas sehingga jumlah data pada kedua kelas seimbang. Hal ini dilakukan untuk mencegah adanya *overfitting*.

Tabel 2. Hasil teknik random sampling jenis oversampling.

Kelas	Sebelum Over Sampling	Setelah Over Sampling
Depresi (1)	608	10.538
Non-Depresi (0)	10.538	10.538
Total	11.156	21.076

Penyetelan (Tuning) Model LSTM

Model LSTM dibangun dengan learning rate=0,0001, *optimizer* Adam, loss function=binary crossentropy, dan aktivasi fungsi softmax. *Optimizer* Adam atau *Adaptive Moment Estimation* adalah algoritma untuk teknik optimasi penurunan gradien yang didasarkan pada estimasi momen adaptif urutan pertama dan urutan kedua [31]. Menurut Kingma dkk. [32], metode ini efisien secara komputasi, memiliki kebutuhan memori yang sedikit, dan cocok untuk masalah yang besar dalam hal data/parameter. *Learning rate* yang rendah diharapkan membantu model bekerja lebih optimal saat training. *Binary crossentropy* untuk *loss function* dipilih karena kelas yang akan diklasifikasi termasuk kelas biner atau hanya ada dua pilihan kelas. Tujuan dari *loss function* adalah untuk menghitung kuantitas *loss* yang harus diusahakan model untuk diminimalisasi [33].

Adapun kelas EarlyStopping digunakan untuk menghentikan proses training jika tidak ada peningkatan akurasi. Pengulangan training model.fit() akan memeriksa di akhir setiap epoch apakah kerugian tidak lagi berkurang, dengan mempertimbangkan patience. Patience adalah jumlah epoch yang akan dihentikan jika tidak ada peningkatan pada proses *training* [34].

4. Evaluasi

Hasil Pengujian

Dataset yang digunakan dibagi menjadi dua bagian untuk data testing dan data training, dengan perbandingan data testing 20:80 data training. Pembagian data testing dan data training dipilih karena umumnya 20:80 yang dipakai, dan sebelumnya sudah mencoba 40:60, 30:70, dan 10:90, tetapi perbandingan tersebut tidak memberikan hasil yang lebih besar dibandingkan menggunakan 20:80. Total dataset yang digunakan untuk training yaitu 16.861 *tweets*, sedangkan untuk testing totalnya adalah 4215 *tweets*. Setelah membagi dataset, dilakukan pembagian data training dengan metode 10-cross validation yang merupakan bagian dari modul *python Scikit-learn* [35]. Cross validation ini bertujuan untuk menghindari kegagalan prediksi pada sampel data yang belum teruji, atau yang biasa disebut dengan *overfitting* [36]. Sebelum model LSTM dibangun, dilakukan proses pembentukan model vector untuk setiap token yang ada pada dataset menggunakan library Gensim FastText [29]. Hipotesa untuk penelitian ini ialah ukuran batch size dan epoch mempengaruhi hasil akurasi. Percobaan penelitian ini dilakukan berkali-kali untuk mendapatkan hasil yang terbaik.

Tabel 3. Hasil dari percobaan model LSTM.

No.	LSTM Size	Epoch	Batch Size	Vektor FastText	Dropout	True Negative	False Positive	False Negative	True Positive	Accuracy	F1
1.	64	10	128	256	0,2	39,15%	10,42%	38,96%	11,48%	0,5062	0,3175
2.	64	50	128	64	0,3	40,57%	8,99%	40,85%	9,58%	0,5015	0,2731
3.	128	3	5	64	0,2	29,82%	19,74%	30,65%	19,78%	0,4961	0,4399
4.	128	3	8	32	0,1	25,88%	23,73%	25,41%	24,98%	0,5086	0,5041
5.	128	3	8	128	0,1	34,85%	14,75%	36,32%	12,79%	0,4892	0,3336
6.	128	5	8	64	0,1	39,55%	10,06%	40,57%	9,82%	0,4937	0,2795
7.	128	5	16	64	0,2	38,05%	11,55%	39,17%	11,22%	0,4927	0,3067

8.	128	5	32	64	0,2	33,83%	15,73%	34,66%	15,78%	0,4961	0,3851
9.	128	5	64	64	0,2	37,79%	11,81%	38,79%	11,60%	0,4939	0,3143
10.	128	10	8	32	0,1	37,53%	12,08%	37,48%	12,91%	0,5043	0,3424
11.	128	10	8	64	0,2	34,57%	15,04%	34,80%	15,59%	0,5015	0,3848
12.	128	10	16	64	0,2	37,11%	12,45%	37,25%	13,19%	0,5029	0,3436
13.	128	10	32	128	0,1	43,75%	5,81%	44,32%	6,12%	0,4986	0,1963
14.	128	15	8	64	0,2	42,37%	7,19%	42,47%	7,97%	0,5034	0,2431
15.	128	20	100	64	0,2	41,07%	8,49%	41,23%	9,20%	0,5034	0,2702
16.	128	100	128	32	0,1	39,60%	10,01%	40,59%	9,80%	0,4939	0,2791

Pemilihan LSTM Size 128 berdasarkan referensi penelitian Uddin dkk. [16] yang menunjukkan LSTM Size 128 memberikan hasil akurasi paling tinggi. Pemilihan LSTM Size 64 untuk melihat apakah LSTM Size yang lebih kecil dapat bekerja lebih baik di penelitian ini. Hasilnya, nilai LSTM Size 128 memberikan nilai lebih tinggi dibanding 64, sehingga percobaan lainnya dilanjutkan dengan LSTM Size 128. Selanjutnya, untuk pemilihan epoch cenderung didominasi oleh nilai yang kecil, dikarenakan nilai epoch diatas 50 menyebabkan program terhenti sebelum eksekusi program selesai. Nilai epoch yang dicoba mulai dari 3 dan maksimalnya 50. Penelitian [16] berhasil mendapatkan akurasi tertinggi dengan epoch 20, sehingga penelitian ini mencoba nilai epoch yang sama. Penurunan nilai epoch dilakukan untuk melihat adanya peningkatan untuk akurasi dan F1. Pada penelitian ini, nilai akurasi dan F1 tertinggi justru didapatkan dengan nilai epoch 3. Untuk pemilihan batch size, menurut artikel [37], nilai batch size yang digunakan pada umumnya mencakup nilai 32, 64, dan 128. Pada penelitian ini, variasi nilai batch size ditambahkan untuk melihat adanya peningkatan nilai akurasi dan F1. Tidak jauh berbeda dengan epoch, nilai batch size yang kecil pun memberikan nilai akurasi dan F1 yang lebih tinggi dibanding nilai batch size yang lebih besar. Untuk penentuan parameter vector size=256 dan dropout=0.2, penelitian ini menggunakan referensi dari program sentimen analisis yang juga menggunakan Gensim FastText [38]. Penentuan vector size yang lainnya mengikuti cakupan batch size pada umumnya (32, 64, dan 128), mengingat referensi pertama bernilai 256. Variasi penentuan nilai dropout 0.1 dan 0.3 dipilih berdasarkan apakah kedua nilai tersebut dapat memberikan nilai akurasi dan F1 yang lebih baik dibanding nilai dropout 0.2. Percobaan ini sempat mencoba menggunakan nilai epoch=100 dan batch_size=128, tetapi percobaan terhenti di epoch ke 26, karena fungsi EarlyStopping yang menghentikan percobaan jika tidak ada peningkatan akurasi.

Untuk mengetahui bagaimana hasil prediksi terhadap data yang sebenarnya, penelitian ini menggunakan *confusion matrix* terhadap hasil prediksi. Prediksi didapatkan dari fungsi 'model.prediction_class()' agar dapat mengelompokkan hasil prediksi sesuai kelas label yang ada, yaitu 0 dan 1. Persamaan untuk menghitung *accuracy*, *recall*, *precision*, dan F1 adalah sebagai berikut :

- *Accuracy* merupakan nilai yang menunjukkan kedekatan antara data yang diprediksi dengan data yang sebenarnya.

$$Accuracy = \frac{True\ Positive + True\ Negative}{(True\ Positive + False\ Positive + False\ Negative + True\ Negative)} \quad (7)$$

- *Precision* merupakan perbandingan dari data yang diprediksikan positif sebenarnya dengan data yang diprediksikan positif.

$$Precision = \frac{True\ Positive}{(True\ Positive + False\ Positive)} \quad (8)$$

- *Recall* merupakan perbandingan dari data yang diprediksikan positif sebenarnya dengan data yang positif sebenarnya.

$$Recall = \frac{True\ Positive}{(True\ Positive + False\ Negative)} \quad (9)$$

- F1 merupakan nilai rata-rata harmonik dari Precision dan Recall.

$$F1 = \frac{True\ Positive}{(True\ Positive + \frac{1}{2}(False\ Positive + False\ Negative))} \quad (10)$$

Setelah dilakukannya percobaan menjalankan model yang dibuat, dapat dilihat bahwa hasil *confusion*

matrix setiap percobaan cenderung tidak jauh berbeda. Hal ini dapat dilihat dari jumlah kelas *true* dan *false* yang selalu hampir seimbang, sehingga hasil akurasi yang didapatkan berada pada rentang 0,4 hingga 0,5. Pada Tabel 3, Hasil akurasi dan F1 tertinggi didapatkan dari percobaan nomor 10, yaitu sebesar 0,5086 dan 0,5041 dengan penyetelan parameter vector word embedding=32, dimensi LSTM=128, batch size=8, epoch=3, dropout=0,1. Model menghasilkan nilai akurasi yang baik saat training, tetapi saat testing tidak menunjukkan hasil yang cukup baik.

Analisis Hasil Pengujian

Dari Tabel 3, dapat dilihat bahwa hasil akurasi tidak ada yang melebihi 0,6, dan perbandingan jumlah data *false* dan *true* hampir selalu seimbang. Meskipun pada percobaan ke-10 jumlah prediksi antara kelas 1 dan 0 seimbang, tetapi jumlah *false negative* dan *false positive*-nya pun mendekati setengah dari jumlah prediksi. Hal ini dapat dilihat dari nilai F1 yang tidak terlalu jauh dengan nilai akurasi. Berikut adalah contoh hasil prediksi pada percobaan ke-10:

Tabel 4. Hasil dari percobaan model LSTM.

Label	Data	Prediksi
1	['sakit', 'mental', 'saya', 'rumah']	0
1	['sudah', 'kali', 'saya', 'tekan', 'mental', 'kaya', 'roller', 'coaster']	0
1	['sedih', 'sangat', 'saya']	0
0	['ingat', 'ibu', 'tidak', 'kurus']	0
0	['orang', 'main', 'among', 'us', 'bodoh', 'kenapa', 'kru', 'tuduh', 'belum', 'kalau', 'masuk', 'ruang', 'seperti', 'detik', 'tiba tiba', 'sudah', 'kalah', 'jujur', 'malas', 'sangat', 'main', 'banyak', 'cheaters', 'gin']	0
0	['lo', 'tampan', 'sangat', 'saya', 'pusing', 'changbin']	0
1	['tugas', 'daring', 'buat', 'pusing']	0
1	['apa', 'sudah', 'enek', 'sangat', 'kuliah', 'daring', 'yaallah', 'tolong', 'sangat', 'corona', 'sudah']	1
0	['album', 'be', 'seperti', 'barang', 'mahal', 'beli', 'pakai', 'uang']	1
1	['capek', 'sangat', 'hanya', 'ingin', 'hidup', 'my', 'own', 'dengan', 'joyfulness', 'tetapi', 'kenapa', 'hidup', 'isi', 'sedih']	1
1	['corona', 'sudah', 'sudah', 'gitu', 'capek', 'sumpah', 'capek']	1
1	['tugas', 'hari', 'tidak ada', 'kerja', 'buntu', 'sangat', 'otak', 'butuh', 'refreshing']	1
0	['mbak', 'walaupun', 'foto', 'profil', 'pakai', 'masker', 'lihat', 'cantik']	1
0	['baik', 'gitu', 'nadin', 'besok', 'turun', 'jalan', 'lihat']	1
0	['sehat', 'masyarakat', 'abai', 'suara', 'rakyat', 'dengar']	1

Pada Tabel 3, dapat dilihat bahwa kombinasi ukuran epoch dan batch size yang bereda memberikan pengaruh terhadap besarnya akurasi. Ukuran batch size dan epoch yang dapat mempengaruhi nilai akurasi juga disebutkan pada penelitian yang dilakukan Uddin dkk. [16]. Penelitian tersebut menyimpulkan bahwa epoch yang dibutuhkan sangat bergantung pada batch size, dan LSTM yang tinggi tidak terlalu membantu untuk memberikan hasil akurasi yang tinggi. Mumu dkk. [20] pun tidak menyebutkan LSTM berpengaruh terhadap akurasi, tetapi menyebutkan bahwa batch size, epoch, dan penyetelan atau *tuning parameter* mempengaruhi perbedaan akurasi.

Tabel 5. Perbandingan dengan penelitian terdahulu.

Penelitian	Dataset	Tuning Parameter	Hasil
Penelitian yang diusulkan (FastText+LSTM)	10.538 label depresi, 10.538 label non- depresi - Hasil oversampling	LSTM size=128 Batch size=8 Epoch=3	Akurasi=0,5086 F1=0,5041
Deteksi Emosi	1304 dataset, dengan 6	LSTM size=50	Akurasi=0,731458

(FastText+LSTM) [15]	kelas 250 label senang, 250 label sedih, 200 label ketakutan, 200 label jijik, 204 label marah, 201 label kaget	Dropout=0,5	F1=0,731458
Deteksi Sarkasme Bahasa Indonesia (Paragraph2Vec+LSTM) [22]	3554 <i>tweet</i> non-sarkasme, 1000 <i>tweet</i> sarkasme	LSTM size=50 Batch size=32 Epochs=20-1000	Akurasi=0,8833 F1=0,8703
Sentimen Analisis Bahasa Indonesia (Word2Vec+ LSTM) [23]	12500 positif 12500 negatif	LSTM size= 64 Batch size = 24 Epoch = 10000	Akurasi=0,9611 F1=0,9573

Pada Tabel 5, menunjukkan hasil akurasi dan F1 yang berbeda-beda berdasarkan kondisi dataset, dan penyetelan atau tuning model LSTM. Penelitian untuk deteksi emosi dengan metode FastText dan LSTM [15], berhasil mendapatkan akurasi dan F1 sebesar 0,731458 dengan distribusi jumlah kelas yang cukup seimbang. Penelitian [22], jika dilihat dari jumlah dataset, dapat dibilang sebagai dataset yang tidak seimbang atau imbalanced. Tetapi, jika dibandingkan dengan kondisi data pada penelitian ini, yaitu 608 untuk label depresi dan 10.538 label non depresi, dataset pada penelitian tersebut tidak mengalami imbalanced yang terlalu jauh antar kelas, dan percobaan dilakukan dengan rentang epoch 20-1000, sehingga berhasil mendapatkan akurasi dan F1 diatas 80%. Persamaan dataset dan penelitian [23] yaitu kedua kelas sama-sama seimbang, tetapi dataset penelitian ini merupakan hasil teknik oversampling.

Dibandingkan dengan penelitian yang dilakukan oleh Uddin dkk. [16], akurasi tertinggi pada penelitian tersebut adalah 0,863, sedangkan untuk penelitian yang diusulkan, akurasi tertinggi yang didapatkan adalah 0,5086. Penelitian [16] pada awalnya mengalami imbalanced data, yaitu 2930 untuk *tweet* non depresi dan 984 *tweet* berlabel depresi, sehingga dilakukan resampling dengan teknik undersampling, sehingga data yang digunakan menjadi 984 *tweet* berlabel non depresi dan 984 *tweet* berlabel depresi. Penelitian yang diusulkan pun mengalami imbalanced data dengan perbandingan yang cukup jauh, yaitu 608 *tweet* berlabel depresi (1) dan 10.538 *tweet* untuk label non-depresi (0). Dengan teknik random oversampling, *tweet* berlabel depresi diperbanyak hingga sesuai dengan jumlah label non-depresi.

Perbandingan banyaknya data yang tidak seimbang data pada penelitian yang diusulkan dengan penelitian-penelitian yang telah disebutkan diatas, dapat menjadi salah satu penyebab nilai akurasi dan F1 yang tidak dapat mencapai diatas 0,6. Penyetelan *hyperparameter*, terutama untuk batch size dan epoch pada model LSTM pun menjadi penyebab akurasi saat testing kurang baik jika dibandingkan dengan nilai akurasi saat training.

5. Kesimpulan

Penelitian ini bertujuan untuk mengidentifikasi depresi melalui media sosial Twitter pada mahasiswa yang menjalankan kuliah *online* saat pandemi Covid-19. Penelitian ini menggunakan dataset yang diambil dari Twitter dengan 10.538 berlabel depresi dan 10.538 berlabel non depresi hasil *random resampling* teknik *oversampling*. Dataset menggunakan teknik *oversampling* untuk mengatasi masalah imbalanced data. Percobaan menjalankan model LSTM pada penelitian ini menghasilkan akurasi dan F1 tertinggi yaitu 0,5086 dan 0,5041 dengan penyetelan vector FastText=32, dimensi LSTM=128, batch size=8, dan epoch=3. Di antara semua percobaan, hanya ada satu percobaan yang nilai akurasi dan F1 yang seimbang, menunjukkan bahwa jumlah nilai benar atau true pada kedua kelas depresi dan non-depresi seimbang. Pada penelitian ini, terdapat masalah yang dihadapi, yaitu ketidakteraturan pada data hasil oversampling dan pemilihan nilai epoch yang besar (di atas 50) yang membuat program terhenti sebelum program selesai dijalankan, sehingga percobaan pada penelitian ini menggunakan nilai epoch di bawah 50 dan nilai akurasi dan F1 yang terbesar justru didapatkan dari jumlah epoch yang kecil disertai batch size yang kecil. Adapun kelemahan lainnya, yaitu proses preprocessing data yang belum sempurna, terutama pada normalisasi, karena terdapat banyak kata yang sulit untuk dinormalisasikan. Dari penelitian dan perbandingan dengan penelitian lain yang sudah dilakukan, dapat disimpulkan bahwa kondisi data, pemilihan *hyperparameter* terutama *batch size* dan epoch mempengaruhi hasil akurasi yang didapatkan.

Untuk penelitian mengidentifikasi depresi ke depannya, dapat diperhatikan kondisi data yang akan digunakan juga penyetelan atau tuning parameter saat menjalankan model. Untuk proses pelabelan, disarankan ada diskusi dengan psikolog atau tenaga kesehatan mental yang profesional, sehingga dapat membantu proses pemilihan label cuitan depresi. Saat crawling tweets, dapat digunakan beberapa kata kunci yang berhubungan dengan depresi. Penelitian ini hanya menggunakan dua label, yaitu 1 untuk label depresi dan 0 untuk label non depresi,

sehingga disarankan untuk menambahkan kelas neutral untuk penelitian yang akan datang.

Daftar Pustaka

- [1] A. P. Kasih, "Mendikbud: Perguruan Tinggi di Semua Zona Dilarang Kuliah Tatap Muka," *Kompas.com*. <https://www.kompas.com/edu/read/2020/06/16/103917571/mendikbud-perguruan-tinggi-di-semua-zona-dilarang-kuliah-tatap-muka> (diakses Okt 06, 2020).
- [2] A. K. Watnaya, M. Hifzul Muiz, N. Sumarni, A. S. Mansyur, dan Q. Y. Zaqiah, "PENGARUH TEKNOLOGI PEMBELAJARAN KULIAH ONLINE DI ERA COVID-19 DAN DAMPAKNYA TERHADAP MENTAL MAHASISWA," *EduTeach J. Edukasi Dan Teknol. Pembelajaran*, vol. 1, no. 2, Art. no. 2, Jun 2020, doi: 10.37859/eduteach.v1i2.1987.
- [3] Z.-H. Wang dkk., "Prevalence of anxiety and depression symptom, and the demands for psychological knowledge and interventions in college students during COVID-19 epidemic: A large cross-sectional study," *J. Affect. Disord.*, vol. 275, hlm. 188–193, Okt 2020, doi: 10.1016/j.jad.2020.06.034.
- [4] W. Santos, A. Funabashi, dan I. Paraboni, "Searching Brazilian Twitter for Signs of Mental Health Issues," dalam *Proceedings of The 12th Language Resources and Evaluation Conference*, Marseille, France, Mei 2020, hlm. 6111–6117. Diakses: Sep 17, 2020. [Daring]. Tersedia pada: <https://www.aclweb.org/anthology/2020.lrec-1.750>
- [5] T. Chomutare, "Text Classification to Automatically Identify Online Patients Vulnerable to Depression," dalam *Pervasive Computing Paradigms for Mental Health*, vol. 100, P. Cipresso, A. Matic, dan G. Lopez, Ed. Cham: Springer International Publishing, 2014, hlm. 125–130. doi: 10.1007/978-3-319-11564-1_13.
- [6] K. Jiang, S. Feng, Q. Song, R. A. Calix, M. Gupta, dan G. R. Bernard, "Identifying tweets of personal health experience through word embedding and LSTM neural network," *BMC Bioinformatics*, vol. 19, no. S8, hlm. 210, Jun 2018, doi: 10.1186/s12859-018-2198-y.
- [7] J.-H. Wang, T.-W. Liu, X. Luo, dan L. Wang, "An LSTM Approach to Short Text Sentiment Classification with Word Embeddings," dalam *Proceedings of the 30th Conference on Computational Linguistics and Speech Processing (ROCLING 2018)*, Hsinchu, Taiwan, Okt 2018, hlm. 214–223. Diakses: Okt 04, 2020. [Daring]. Tersedia pada: <https://www.aclweb.org/anthology/O18-1021>
- [8] S. Tsugawa, Y. Kikuchi, F. Kishino, K. Nakajima, Y. Itoh, dan H. Ohsaki, "Recognizing Depression from Twitter Activity," dalam *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems - CHI '15*, Seoul, Republic of Korea, 2015, hlm. 3187–3196. doi: 10.1145/2702123.2702280.
- [9] "SKB Pembelajaran Tahun Ajaran Baru di Masa Pandemi COVID-19 - Regulasi," *covid19.go.id*, 2020. <https://covid19.go.id/p/regulasi/skb-pembelajaran-tahun-ajaran-baru-di-masa-pandemi-covid-19> (diakses Agu 05, 2021).
- [10] F. A. Rochimah, "DAMPAK KULIAH DARING TERHADAP KESEHATAN MENTAL MAHASISWA DITINJAU DARI ASPEK PSIKOLOGI," hlm. 7, 2020.
- [11] M. Deshpande dan V. Rao, "Depression detection using emotion artificial intelligence," dalam *2017 International Conference on Intelligent Sustainable Systems (ICISS)*, Palladam, Des 2017, hlm. 858–862. doi: 10.1109/ISS1.2017.8389299.
- [12] N. A. Asad, Md. A. Mahmud Pranto, S. Afreen, dan Md. M. Islam, "Depression Detection by Analyzing Social Media Posts of User," dalam *2019 IEEE International Conference on Signal Processing, Information, Communication & Systems (SPICSCON)*, Dhaka, Bangladesh, Nov 2019, hlm. 13–17. doi: 10.1109/SPICSCON48833.2019.9065101.
- [13] A. Kumar, A. Sharma, dan A. Arora, "Anxious Depression Prediction in Real-time Social Data," *SSRN Electron. J.*, 2019, doi: 10.2139/ssrn.3383359.
- [14] J. Figueredo dan R. Calumby, "On Text Preprocessing for Early Detection of Depression on Social Media," *An. Princ. Simpósio Bras. Comput. Apl. À Saúde SBCAS*, hlm. 84–95, Sep 2020, doi: 10.5753/sbcas.2020.11504.
- [15] M. A. Riza dan N. Charibaldi, "Emotion Detection in Twitter Social Media Using Long Short-Term Memory (LSTM) and Fast Text," *Int. J. Artif. Intell. Robot. IJAIR*, vol. 3, no. 1, hlm. 15–26, Mei 2021, doi: 10.25139/ijair.v3i1.3827.
- [16] A. H. Uddin, D. Bapery, dan A. S. M. Arif, "Depression Analysis from Social Media Data in Bangla Language using Long Short Term Memory (LSTM) Recurrent Neural Network Technique," dalam *2019 International Conference on Computer, Communication, Chemical, Materials and Electronic Engineering (IC4ME2)*, Rajshahi, Bangladesh, Jul 2019, hlm. 1–4. doi: 10.1109/IC4ME247184.2019.9036528.
- [17] J. Brownlee, "A Gentle Introduction to Long Short-Term Memory Networks by the Experts," *Machine Learning Mastery*, Mei 23, 2017. <https://machinelearningmastery.com/gentle-introduction-long-short-term-memory-networks-experts/> (diakses Okt 27, 2020).

- [18] P. Goyal, S. Pandey, dan K. Jain, *Deep Learning for Natural Language Processing*. Berkeley, CA: Apress, 2018. doi: 10.1007/978-1-4842-3685-7.
- [19] S. Hochreiter dan J. Schmidhuber, "Long Short-Term Memory," *Neural Comput.*, vol. 9, no. 8, hlm. 1735–1780, Nov 1997, doi: 10.1162/neco.1997.9.8.1735.
- [20] T. F. Mumu, I. J. Munni, dan A. K. Das, "Depressed People Detection from Bangla Social Media Status using LSTM and CNN Approach," *J. Eng. Adv.*, vol. 2, no. 01, Art. no. 01, Mar 2021, doi: 10.38032/jea.2021.01.006.
- [21] F. Miedema dan S. Bhulai, "Sentiment Analysis with Long Short-Term Memory networks," 2018. <https://www.semanticscholar.org/paper/Sentiment-Analysis-with-Long-Short-Term-Memory-Miedema-Bhulai/5f7c2ea6d1974a91b99dfab62121176d2e874015> (diakses Jul 15, 2021).
- [22] S. Khotijah, J. Tirtawangsa, dan A. A. Suryani, "Using LSTM for Context Based Approach of Sarcasm Detection in Twitter," dalam *Proceedings of the 11th International Conference on Advances in Information Technology*, Bangkok Thailand, Jul 2020, hlm. 1–7. doi: 10.1145/3406601.3406624.
- [23] L. Kurniasari dan A. Setyanto, "SENTIMENT ANALYSIS USING RECURRENT NEURAL NETWORK-LSTM IN BAHASA INDONESIA," vol. 15, hlm. 15, 2020.
- [24] P. Bojanowski, E. Grave, A. Joulin, dan T. Mikolov, "fastText," *Facebook Research*, Agu 18, 2016. <https://research.fb.com/blog/2016/08/fasttext/> (diakses Okt 26, 2020).
- [25] A. Joulin, E. Grave, P. Bojanowski, dan T. Mikolov, "Bag of Tricks for Efficient Text Classification," *ArXiv160701759 Cs*, Agu 2016, Diakses: Okt 15, 2020. [Daring]. Tersedia pada: <http://arxiv.org/abs/1607.01759>
- [26] A. Alessa, M. Faezipour, dan Z. Alhassan, "Text Classification of Flu-Related Tweets Using FastText with Sentiment and Keyword Features," dalam *2018 IEEE International Conference on Healthcare Informatics (ICHI)*, New York, NY, Jun 2018, hlm. 366–367. doi: 10.1109/ICHI.2018.00058.
- [27] J. Truschel, "Depression Definition and DSM-5 Diagnostic Criteria," *Psycom.net - Mental Health Treatment Resource Since 1986*, 2020. <https://www.psycom.net/depression-definition-dsm-5-diagnostic-criteria/> (diakses Okt 24, 2020).
- [28] H. A. Robbani, *Sastrawi Python*. 2021. Diakses: Agu 03, 2021. [Daring]. Tersedia pada: <https://github.com/har07/PySastrawi>
- [29] R. Rehurek dan P. Sojka, "Software Framework for Topic Modelling with Large Corpora," Mei 2010. <https://radimrehurek.com/gensim/models/fasttext.html> (diakses Jul 10, 2021).
- [30] J. Brownlee, "Random Oversampling and Undersampling for Imbalanced Classification," *Machine Learning Mastery*, Jan 14, 2020. <https://machinelearningmastery.com/random-oversampling-and-undersampling-for-imbalanced-classification/> (diakses Jun 08, 2021).
- [31] K. Team, "Keras documentation: Adam." <https://keras.io/api/optimizers/adam/> (diakses Jul 19, 2021).
- [32] D. P. Kingma dan J. Ba, "Adam: A Method for Stochastic Optimization," *ArXiv14126980 Cs*, Jan 2017, Diakses: Jul 19, 2021. [Daring]. Tersedia pada: <http://arxiv.org/abs/1412.6980>
- [33] K. Team, "Keras documentation: Losses." <https://keras.io/api/losses/> (diakses Jul 19, 2021).
- [34] K. Team, "Keras documentation: EarlyStopping." https://keras.io/api/callbacks/early_stopping/ (diakses Jul 19, 2021).
- [35] F. Pedregosa dkk., "Scikit-learn: Machine Learning in Python," *J. Mach. Learn. Res.*, vol. 12, no. 85, hlm. 2825–2830, 2011.
- [36] F. Pedregosa dkk., "3.1. Cross-validation: evaluating estimator performance — scikit-learn 0.23.2 documentation." https://scikit-learn.org/stable/modules/cross_validation.html (diakses Nov 24, 2020).
- [37] J. Brownlee, "Difference Between a Batch and an Epoch in a Neural Network," *Machine Learning Mastery*, Jul 19, 2018. <https://machinelearningmastery.com/difference-between-a-batch-and-an-epoch/> (diakses Agu 17, 2021).
- [38] J. Mandav, *Neural-Network*. 2021. Diakses: Agu 17, 2021. [Daring]. Tersedia pada: https://github.com/jatinmandav/Neural-Networks/blob/88d3f07d332427ffcd3321f5841be4827f7f90f/Sentiment-Analysis/fastText/sentiment_analysis_fasttext.ipynb