

Analisis Sentimen pada Media Sosial Twitter terhadap Penanganan Bencana Banjir di Jawa Barat dengan Metode Jaringan Saraf Tiruan

Sentiment Analysis On Twitter Social Media On Flood Disaster Management In West Java With Neural Network Method

Aurell Layalia Safara Az-Zahra Gunawan¹, Jondri, S.Si., M.Si.²,
Dr. Kemas Muslim Lhaksamana, S.T., M.ISD.³

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

¹aurellsag@student.telkomuniversity.ac.id, ²jondri@telkomuniversity.ac.id,

³kemasmuslim@telkomuniversity.ac.id

Abstrak

Analisis sentimen adalah bidang studi yang menganalisis opini dan sentimen seseorang terhadap suatu masalah. Masalah yang diangkat dalam penelitian ini adalah bencana banjir yang melanda Jawa Barat pada bulan Januari tahun 2020. Penelitian ini dilakukan untuk menganalisis bagaimana opini masyarakat tentang penanganan bencana banjir di Jawa Barat pada bulan Januari tahun 2020 di media sosial Twitter. Analisis dilakukan dengan melakukan klasifikasi tweet yang berisi sentimen masyarakat terhadap penanganan bencana banjir. Metode klasifikasi yang digunakan dalam penelitian ini adalah Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation. Metode fitur ekstraksi yang digunakan dalam penelitian ini adalah Term Frequency – Inversed Document Frequency (TF-IDF). Metode validasi yang digunakan dalam penelitian ini adalah Confusion Matrix. Hasil pengujian pada model yang dibangun memberikan hasil akurasi yang cukup baik yaitu 73.83%.

Kata Kunci: Analisis Sentimen, Twitter, Banjir Jawa Barat, Term Frequency – Inversed Document Frequency (TF- IDF), Jaringan Saraf Tiruan (JST).

Abstract

Sentiment analysis is a field that analyzes a person's opinion and sentiment on a problem. The problem raised in this study is the flood disaster that hit West Java in January 2020. This research was conducted to analyze how public opinion about flood disaster management in West Java in January 2020 on Twitter social media. The analysis was carried out by classifying tweets containing public sentiment regarding flood disaster management. The classification method used in this research is an Artificial Neural Network (ANN) Multi Layer Perceptron (MLP) model with Backpropagation algorithm. The word weighting method used in this research is the Term Frequency - Inversed Document Frequency (TF-IDF). The validation method used in this research is the Confusion Matrix. The test results on the model built gave accurate results in a fairly good way, namely 73.83%.

Keywords: Sentiment Analysis, Twitter, West Java Flood, Term Frequency - Inversed Document Frequency (TF-IDF), Artificial Neural Networks (ANN).

1. Pendahuluan

Latar Belakang

Bencana banjir adalah bencana alam yang terjadi ketika aliran air yang berlebihan merendam daratan yang umumnya disebabkan oleh curah hujan yang meningkat, eksploitasi air tanah yang berlebihan, saluran air tersumbat, dan lainnya [6]. Bencana banjir merupakan salah satu bencana yang sering melanda berbagai daerah di Jawa Barat setiap tahunnya, salah satu bencana banjir yang terparah terjadi pada bulan Januari tahun 2020 yang menyebabkan berbagai permasalahan di pemukiman masyarakat di Jawa Barat.

Bencana banjir yang melanda berbagai daerah di Jawa Barat pada bulan Januari tahun 2020 ini menjadi perhatian berbagai kalangan masyarakat. Banyak masyarakat yang memanfaatkan media sosial sebagai sarana untuk berbagi informasi atau mengungkapkan pendapat tentang bencana banjir yang melanda suatu daerah di Jawa Barat. Salah satu media sosial yang digunakan adalah media sosial Twitter. Pengguna Twitter memiliki

reaksi yang beragam terhadap bencana banjir di Jawa Barat ini, pada bulan Januari tahun 2020 ada banyak sekali tweet yang diunggah mengenai banjir yang melanda berbagai daerah di Jawa Barat. Pendapat yang paling sering diungkapkan pengguna Twitter adalah tentang penanganan bencana banjir di daerah mereka.

Twitter adalah media sosial yang dapat digunakan pengguna untuk mengirim dan membaca pesan berbasis teks yang dikenal dengan sebutan kicauan atau tweet [17]. Menurut Dwi Andriansah selaku Country Industry Head Twitter Indonesia [7], pada laporan finansial Twitter kuartal ke-3 tahun 2019 dinyatakan bahwa pengguna aktif Twitter di Indonesia ada sebanyak 145 juta pengguna. Pengguna aktif Twitter di Indonesia tersebut merupakan salah satu pengguna aktif terbanyak di dunia.

Menurut Nayomi Kankanamge [9], memanfaatkan Twitter merupakan pendekatan yang menjanjikan untuk mencerminkan pengetahuan warga negara. Penelitian ini menunjukkan bagaimana tweet dapat digunakan untuk mengidentifikasi fluktuasi keparahan bencana dari waktu ke waktu dan memvalidasi penerapan pesan geo lokasi untuk membatasi zona bencana yang sangat berdampak. Nayomi Kankanamge mengatakan diharapkan pada penelitian di masa depan dapat menggunakan model yang lebih canggih untuk klasifikasi sentimen dan mendeteksi emosi yang halus, mendapatkan tweet yang informatif di Twitter selama peristiwa bencana dan memahami bagaimana analisis sentimen geo-mapped bekerja di berbagai bencana sangat diinginkan.

Menurut Venkata K. Neppalli [13], melakukan analisis sentimen menggunakan media sosial seperti Twitter dapat membantu responden darurat mengembangkan kesadaran situasional yang lebih kuat dari zona bencana itu sendiri. Penelitian ini dapat menunjukkan bagaimana sentimen pengguna berubah berdasarkan lokasi mereka dan jarak mereka dari bencana, serta bagaimana perbedaan sentimen dalam tweet yang diunggah selama badai mempengaruhi kemampuan retweet dari sebuah tweet. Venkata K. Neppalli mengatakan diharapkan pada penelitian di masa depan dapat mengunduh tweet geo-tag lalu mengklasifikasikan teks untuk analisis sentimen dan menggunakan algoritma pembelajaran mesin yang cocok untuk analisis sentimen.

Topik dan Batasannya

Topik dalam penelitian ini adalah melakukan penerapan metode Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation untuk mengklasifikasi tweet tentang penanganan bencana banjir di Jawa Barat ke dalam kelas positif atau negatif dan untuk mengetahui pengaruh dari data shuffling, learning rate, node hidden layer dan drop out terhadap nilai akurasi dari model yang dibangun.

Batasan dalam penelitian ini adalah dataset terdiri dari 3.000 tweet yang berkaitan dengan penanganan bencana banjir di Jawa Barat dan diunggah pada bulan Januari tahun 2020 di media sosial Twitter. Pengujian model yang dibangun dilakukan dengan 10 skenario berbeda, setiap skenario memiliki data shuffling, learning rate, node hidden layer dan drop out yang berbeda.

Tujuan

Tujuan dalam penelitian ini adalah untuk mengklasifikasi tweet tentang penanganan bencana banjir di Jawa Barat ke dalam kelas positif atau negatif menggunakan metode Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation dan untuk mengetahui data shuffling, learning rate, node hidden layer dan drop out terbaik untuk digunakan pada model yang dibangun agar mendapatkan nilai akurasi yang maksimal.

Organisasi Tulisan

Pada bab 1 pendahuluan berisi tentang latar belakang, topik dan batasannya, tujuan dan organisasi tulisan. Pada bab 2 studi terkait berisi tentang studi yang berhubungan dengan tugas akhir ini. Pada bab 3 sistem yang dibangun berisi tentang penjelasan sistem yang dibangun dalam tugas akhir ini. Pada bab 4 evaluasi berisi tentang hasil pengujian dan analisis hasil pengujian. Pada bab 5 kesimpulan berisi tentang kesimpulan dari tugas akhir ini. Pada bagian penutup berisi tentang daftar pustaka dan lampiran.

2. Studi Terkait

Media sosial adalah teknologi mediasi digital interaktif yang memfasilitasi kreasi, berbagi dan bertukar informasi, ide, minat, karier dan bentuk ekspresi lain melalui komunitas dan jaringan virtual [16]. Saat ini media sosial sudah menjadi kebutuhan kehidupan sehari-hari masyarakat dunia. Ada berbagai macam aplikasi media sosial yang memiliki fungsi dan fitur yang berbeda untuk memenuhi kebutuhan para penggunanya. Twitter merupakan salah satu aplikasi media sosial yang sering digunakan oleh masyarakat. Twitter diluncurkan untuk

publik pada tahun 2006, sekarang pengguna Twitter sudah ratusan juta [17]. Berdasarkan hal tersebut, Twitter memiliki banyak sekali data tweet yang berisi berbagai macam opini penggunaannya di seluruh dunia. Dengan banyaknya data yang dimiliki Twitter, maka banyak peneliti yang memanfaatkan data ini untuk melakukan penelitian seperti analisis sentimen.

Pada penelitian tentang aplikasi linier regresi dengan algoritma jaringan syaraf tiruan untuk sentimen analisis [15], peneliti melakukan analisis sentimen terhadap tweet yang berhubungan dengan pemilihan calon pemimpin presiden dan pemilihan kepala daerah yang ada di Indonesia. Analisis sentimen ini bertujuan untuk mengklasifikasi sentimen secara otomatis ke tiga kelas sentimen yaitu positif, netral dan negatif menggunakan algoritma jaringan syaraf tiruan. Hasil yang didapat dari penelitian ini adalah hasil performansi algoritma yang didapatkan sebesar 53.33%.

Selanjutnya pada penelitian tentang analisis sentimen pada media sosial twitter untuk klasifikasi opini islam radikal menggunakan jaringan syaraf tiruan [18], peneliti melakukan analisis sentimen terhadap unggahan seseorang yang terduga sebagai konten yang mengandung ajaran radikalisme di media sosial Twitter. Analisis sentimen ini menggunakan algoritma jaringan syaraf tiruan dengan metode backpropagation. Hasil yang didapat dari penelitian ini adalah algoritma jaringan syaraf tiruan dengan akurasi sebesar 77%, presisi sebesar 77.1%, recall sebesar 73.4% dan f-measure 75.25%. Analisis hasil menggunakan graph wordcoudl dan graph network menunjukkan kecenderungan kategori sentimen opini tergantung pada semboyan dan kata yang digunakan pada tweet.

Terakhir pada penelitian tentang analisis sentimen mahasiswa terhadap fasilitas Universitas Telkom menggunakan metode jaringan syaraf tiruan dan TF-IDF [12], peneliti melakukan analisis sentimen terhadap kuesioner yang berisi sentimen mahasiswa tentang fasilitas di Universitas Telkom. Analisis sentimen ini menggunakan algoritma jaringan syaraf tiruan dengan metode multi-layer perceptron yang dikombinasikan dengan fitur ekstraksi untuk mendeteksi negasi dan pembobotan menggunakan TF-IDF. Hasil yang didapat dari penelitian ini adalah hasil akurasi aplikasi yang dibangun sebesar 91.23%.

Analisis Sentimen

Analisis Sentimen adalah bidang studi yang menganalisis opini, sentimen, evaluasi, penilaian, sikap, dan emosi orang terhadap entitas seperti produk, layanan, organisasi, individu, masalah, peristiwa, topik, dan atributnya [10]. Analisis Sentimen terbagi menjadi dua proses yaitu Sentiment Extraction dan Sentiment Classification. Sentiment Extraction adalah proses mengekstraksi aspek yang telah dievaluasi [10]. Sentiment Classification adalah proses menentukan opini tentang aspek yang berbeda itu bernilai positif, negatif atau netral [10].

Data Preprocessing

Data Preprocessing adalah proses yang dilakukan untuk menyiapkan data dalam format yang sesuai agar dapat digunakan sebagai dataset berkualitas sehingga proses ekstraksi pengetahuan dapat diterapkan [2]. Proses yang dilakukan dalam preprocessing adalah case folding, cleansing, tokenizing, normalizing, stopword removal dan stemming. Case Folding adalah proses mengubah semua huruf kapital menjadi huruf kecil. Cleansing adalah proses menghapus tab, new line, back slice, emoticon, non ASCII, mention, hashtag, link, URL, angka, tanda baca dan whitespace. Tokenizing adalah proses memotong kalimat menjadi beberapa susunan kata. Normalizing adalah proses mengubah kata gaul menjadi kata baku. Stopword Removal adalah proses menghapus kata sambung dan kata yang dianggap tidak memiliki makna. Stemming adalah proses mengubah kata dengan imbuhan menjadi kata dasar.

Term Frequency – Inverse Document Frequency (TF-IDF)

Term Frequency – Inverse Document Frequency (TF-IDF) adalah metode feature extraction dengan menghitung nilai Term Frequency dan menghitung kemunculan sebuah kata pada koleksi dokumen teks secara keseluruhan. [1]. Term Frequency adalah jumlah kemunculan sebuah kata di dalam sebuah dokumen tertentu [12], semakin sering kata yang muncul maka semakin besar nilai Term Frequency. Inverse Document Frequency adalah jumlah dokumen yang mengandung sebuah kata didasarkan pada seluruh dokumen yang ada pada dataset [12], semakin jarang kata yang muncul maka semakin besar nilai Inverse Document Frequency. Hasil dari feature extraction adalah perkalian dari nilai Term Frequency dan Inverse Document Frequency yang akan menghasilkan bobot lebih kecil jika kata yang muncul lebih sering dan akan menghasilkan bobot lebih besar jika kata yang muncul lebih jarang [14].

Persamaan untuk perhitungan Inverse Document Frequency adalah sebagai berikut [4]:

$$IDF = \log\left(\frac{N}{DF_i}\right) + 1$$

Keterangan:

N = Jumlah dokumen
 DF_i = Jumlah dokumen d yang mengandung kata i

Persamaan untuk perhitungan Term Frequency – Inverse Document Frequency adalah sebagai berikut [5]:

$$TF - IDF = TF(i, d) \times IDF(i)$$

Keterangan:

w_i = Kata ke i
 d = Dokumen
 $TF(i, d)$ = Jumlah kemunculan kata i pada dokumen d
 $IDF(i)$ = Inverse Document Frequency dari kata i

K-Fold Cross Validation

K-Fold Cross Validation adalah prosedur pembagian dataset menjadi data train dan data test untuk memvalidasi model yang telah dibangun. K-Fold Cross Validation bertujuan untuk menemukan kombinasi data terbaik untuk model yang telah dibangun.

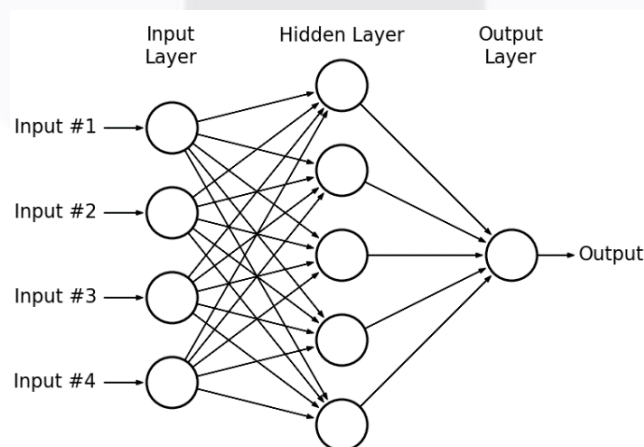
Jaringan Saraf Tiruan (JST)

Jaringan Saraf Tiruan (JST) adalah jaringan paradigma pemrosesan suatu informasi yang terinspirasi oleh sistem sel saraf manusia (neuron) [3]. Jaringan saraf tiruan terbagi menjadi dua jenis berdasarkan sifatnya yaitu supervised learning dan unsupervised learning. Pada supervised learning tahap pelatihan dilakukan dengan memasukkan data train ke dalam jaringan yang akan mengubah bobot yang menjadi penghubung antar node.

Lapisan penyusun jaringan saraf tiruan adalah lapisan input berisi unit input yang dapat menerima pola masukan data dari luar yang menggambarkan suatu permasalahan, lapisan hidden berisi unit hidden yang keluarannya tidak dapat secara langsung diamati dan lapisan output berisi unit output yang menghasilkan solusi yang menyelesaikan suatu permasalahan [3].

Multi Layer Perceptron (MLP)

Multi Layer Perceptron (MLP) adalah jaringan lapisan jamak yang terdiri dari satu atau lebih input layer, satu atau lebih hidden layer dan satu output layer. Setiap unit yang terdapat di lapisan input selalu terhubung dengan setiap unit yang terdapat pada layer hidden. Setiap unit yang terdapat di lapisan hidden selalu terhubung dengan setiap unit yang terdapat pada layer output.



Gambar 1 Arsitektur Multi Layer Perceptron (MLP)

Salah satu algoritma pelatihan multi layer perception adalah algoritma backpropagation. Algoritma backpropagation dilakukan dalam tiga tahap yaitu propagasi maju, propagasi mundur dan perubahan bobot. Cara kerja pelatihan algoritma backpropagation adalah sebagai berikut [4]:

1. Propagasi Maju.

Persamaan untuk perhitungan hasil dari hidden layer adalah sebagai berikut:

$$Z_{net\ j} = v_{jo} + \sum_{i=1}^n x_i v_{ji}$$

$$z_j = (z_{net\ j}) = \max(0, z_{net\ j})$$

Persamaan untuk perhitungan hasil dari output layer adalah sebagai berikut:

$$y_{net\ k} = w_{ko} + \sum_{j=1}^p z_j w_{kj}$$

$$y_k = f(y_{net\ k}) = \frac{1}{1 + e^{-y_{net\ k}}}$$

2. Propagasi Mundur

Persamaan untuk perhitungan faktor unit output berdasarkan error pada setiap unit output adalah sebagai berikut:

$$\delta_k = (t_k - y_k) f'(y_{net\ k}) = (t_k - y_k) y_k (1 - y_k)$$

Persamaan untuk perhitungan suku perubahan bobot dengan laju percepatan adalah sebagai berikut:

$$\Delta w_{kj} = a \times \delta_k \times z_j$$

Persamaan untuk perhitungan unit tersembunyi berdasarkan pada error pada setiap unit tersembunyi adalah sebagai berikut:

$$\delta_{net\ j} = \sum_{k=1}^m \delta_k w_{kj}$$

Persamaan untuk perhitungan faktor unit tersembunyi adalah sebagai berikut:

$$\delta_j = \delta_{net\ j} f'(z_{net\ j}) = \delta_{net\ j} z_j$$

Persamaan untuk perhitungan suku perubahan bobot adalah sebagai berikut:

$$\Delta v_{ji} = a \times \delta_j \times x_i$$

3. Perubahan Bobot

Persamaan untuk perhitungan seluruh perubahan bobot yang menuju unit output adalah sebagai berikut:

$$w_{kj}(\text{baru}) = w_{kj}(\text{lama}) + \Delta w_{kj}$$

Persamaan untuk perhitungan seluruh perubahan bobot yang menuju unit tersembunyi adalah sebagai berikut:

$$v_{ji}(\text{baru}) = v_{ji}(\text{lama}) + \Delta v_{ji}$$

Confusion Matrix

Confusion Matrix adalah tata letak tabel khusus yang memungkinkan visualisasi kinerja suatu algoritma. Setiap baris matriks mewakili instance dalam kelas yang diprediksi dan setiap kolom mewakili instance dalam kelas yang sebenarnya atau sebaliknya [11].

Tabel 1 Visualisasi Confusion Matrix

	PREDIKSI	
REALITA	True Positive	False Negative
	False Positive	True Negative

Istilah pada Confusion Matrix adalah sebagai berikut [11]:

1. True Positive adalah jumlah data positif yang terklasifikasi dengan benar oleh sistem.
2. True Negative adalah jumlah data negatif yang terklasifikasi dengan benar oleh sistem.
3. False Positive adalah jumlah data positif namun terklasifikasi salah oleh sistem.
4. False Negative adalah jumlah data negatif namun terklasifikasi salah oleh sistem.

Accuracy adalah perhitungan keakuratan model dalam memberikan klasifikasi yang benar. Persamaan untuk perhitungan accuracy adalah sebagai berikut [11]:

$$Accuracy = \frac{True\ Positive + True\ Negative}{True\ Positive + True\ Negative + False\ Positive + False\ Negative} \times 100\%$$

Precision adalah perhitungan keakuratan antara target yang diminta dan hasil prediksi yang diberikan oleh model. Persamaan untuk perhitungan precision adalah sebagai berikut [11]:

$$Precision = \frac{True\ Positive}{True\ Positive + False\ Positive} \times 100\%$$

Recall adalah perhitungan keberhasilan model dalam menemukan kembali sebuah informasi. Persamaan untuk perhitungan recall adalah sebagai berikut [11]:

$$Recall = \frac{True\ Positive}{True\ Positive + False\ Negative} \times 100\%$$

F1-Score adalah perhitungan perbandingan rata-rata precision dan recall yang dibobotkan. Persamaan untuk perhitungan F1-Score adalah sebagai berikut [8]:

$$F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall} \times 100\%$$

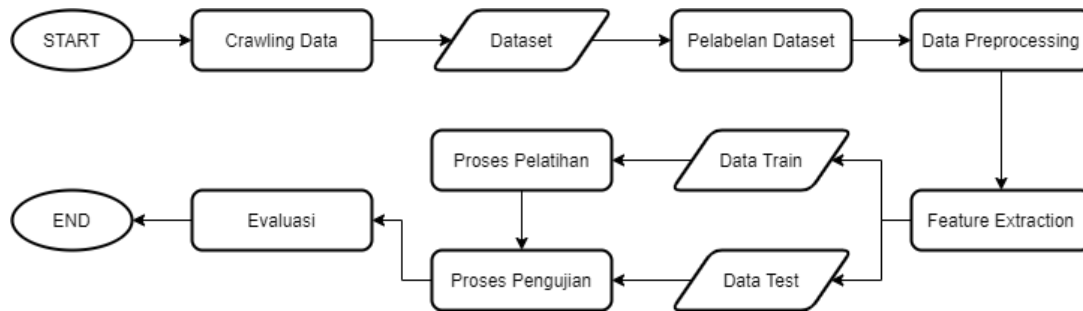
3. Sistem yang Dibangun

Deskripsi Sistem

Sistem yang dibangun dalam penelitian ini adalah pemodelan untuk analisis sentimen pada media sosial Twitter terhadap penanganan bencana banjir di Jawa Barat. Input dari sistem ini adalah dataset berisi tweet tentang penanganan bencana banjir di Jawa Barat. Metode yang digunakan untuk feature extraction adalah Term Frequency – Invers Document Frequency (TF-IDF). Term Frequency – Invers Document Frequency (TF-IDF) dipilih karena metode ini menghasilkan vektor dari perkalian antara frekuensi kemunculan term pada dokumen dan frekuensi dokumen terbalik. Metode yang digunakan untuk klasifikasi adalah Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation. Jaringan Saraf Tiruan (JST) dipilih karena metode ini dapat digunakan untuk melakukan klasifikasi, membuat fungsi prediksi yang kompleks dan meniru pemikiran manusia. Metode validasi yang digunakan dalam penelitian ini adalah Confusion Matrix. Confusion Matrix dipilih karena metode ini cocok untuk digunakan pada kasus klasifikasi biner yaitu positif dan negatif. Output dari sistem ini adalah hasil klasifikasi dari analisis sentimen menjadi 2 kelas yaitu positif dan negatif.

Alur Sistem

Alur sistem yang dibangun dalam penelitian ini digambarkan dengan diagram flow adalah sebagai berikut:



Gambar 2 Diagram Flow Sistem Yang Dibangun

Crawling Data

Proses crawling data dilakukan menggunakan program dengan bahasa pemrograman Python dan library twint. Data yang dikumpulkan adalah tweet di media sosial Twitter tentang penanganan bencana banjir di Jawa Barat. Jika menggunakan kata kunci 'penanganan banjir jawa barat' dan 'penanganan banjir jabar' data yang dikumpulkan hanya sebanyak 200 tweet. Karenanya, kata kunci yang digunakan 'banjir jawa barat' dan 'banjir jabar'. Data yang berhasil dikumpulkan sebanyak 6.000 tweet berbahasa Indonesia dengan periode waktu 1 Januari 2020 sampai 31 Januari 2020. Data yang didapatkan disimpan ke dalam file dengan format CSV.

Dataset

Dataset yang digunakan adalah tweet yang sudah dikumpulkan sebanyak 6.000 tweet. Tweet dengan kata kunci 'banjir jawa barat' dan 'banjir jabar' digabungkan ke dalam satu file dengan format CSV dan diurutkan berdasarkan waktu unggahnya. Karena menggunakan kata kunci yang kurang sesuai, dataset harus dianalisis secara manual satu per satu. Tweet yang tidak berkaitan dengan banjir jawa barat, tweet yang tidak menggunakan bahasa Indonesia dan tweet yang memiliki duplikat akan dihapus. Setelah melalui proses analisis secara manual, dataset yang didapatkan ada sebanyak 3.500 tweet.

Pelabelan Dataset

Proses pelabelan dataset dilakukan secara manual satu per satu. Karena kelas label yang digunakan hanya positif dan negatif, sehingga tweet yang netral akan dimasukkan ke dalam label positif. Ketentuan untuk label positif adalah tweet yang berisi kata-kata positif, memberitakan penanganan bencana banjir, memuji usaha penanganan bencana banjir, dan lainnya. Ketentuan untuk label negatif adalah tweet yang berisi kata-kata negatif, meminta pertolongan, memaki usaha penanganan bencana banjir dan lainnya.

Tabel 2 Contoh Pelabelan Data

Label Angka	Label Kelas	Tweet
0	Negatif	Banjir menyakiti banyak warga, ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir. Pun demikian dengan Jabar & Banten.
1	Positif	https://t.co/vCNx2x69oE Tim Social Disaster Rescue Sekolah Relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di Kawasan Jatiwaringin, Kota Bekasi, Jawa Barat #BanjirJakarta #banjirbekasi #banjir

Data Preprocessing

Proses data preprocessing dilakukan menggunakan program dengan bahasa pemrograman Python, library nltk dan library Sastrawi. Data preprocessing yang dilakukan adalah case folding, cleansing, tokenizing, normalizing, stopwords removal dan stemming. Hasil data preprocessing diperiksa secara manual untuk menghapus tweet bernilai null dan tweet yang memiliki duplikat. Setelah melalui proses data preprocessing, dataset yang didapatkan ada sebanyak 3.000 tweet.

1. Case Folding

Tahap pertama proses data preprocessing adalah case folding yang dilakukan untuk mengubah semua huruf kapital menjadi huruf kecil.

Tabel 3 Contoh Case Folding

Tweet Sebelum Case Folding	Tweet Sesudah Case Folding
Banjir menyakiti banyak warga, ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir. Pun demikian dengan Jabar & Banten.	banjir menyakiti banyak warga, ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir. pun demikian dengan jabar & banten.
https://t.co/vCNx2x69oe Tim Social Disaster Rescue Sekolah Relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di Kawasan Jatiwaringin, Kota Bekasi, Jawa Barat #BanjirJakarta #banjirbekasi #banjir	https://t.co/vcnx2x69oe tim social disaster rescue sekolah relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di kawasan jatiwaringin, kota bekasi, jawa barat #banjirjakarta #banjirbekasi #banjir

2. Cleansing

Tahap kedua proses data preprocessing adalah cleansing yang dilakukan untuk menghapus tab, new line, back slice, emoticon, non ASCII, mention, hashtag, link, URL, angka, tanda baca dan whitespace.

Tabel 4 Contoh Cleansing

Tweet Sebelum Cleansing	Tweet Sesudah Cleansing
banjir menyakiti banyak warga, ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir. pun demikian dengan jabar & banten.	banjir menyakiti banyak warga ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir pun demikian dengan jabar amp banten
https://t.co/vcnx2x69oe tim social disaster rescue sekolah relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di kawasan jatiwaringin, kota bekasi, jawa barat #banjirjakarta #banjirbekasi #banjir	tim social disaster rescue sekolah relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di kawasan jatiwaringin kota bekasi jawa barat

3. Tokenizing

Tahap ketiga proses data preprocessing adalah tokenizing yang dilakukan untuk memotong kalimat menjadi beberapa kata.

Tabel 5 Contoh Tokenizing

Tweet Sebelum Cleansing	Tweet Sesudah Cleansing
banjir menyakiti banyak warga ditambah gubernurnya dan pemerintahnya tidak bisa kerja dalam mengatasi banjir pun demikian dengan jabar amp banten	['banjir', 'menyakiti', 'banyak', 'warga', 'ditambah', 'gubernurnya', 'dan', 'pemerintahnya', 'tidak', 'bisa', 'kerja', 'dalam', 'mengatasi', 'banjir', 'pun', 'demikian', 'dengan', 'jabar', 'amp', 'banten']
tim social disaster rescue sekolah relawan langsung melakukan penyisiran untuk mencari warga yang memerlukan evakuasi di kawasan jatiwaringin kota bekasi jawa barat	['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'melakukan', 'penyisiran', 'untuk', 'mencari', 'warga', 'yang', 'memerlukan', 'evakuasi', 'di', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']

4. Normalizing

Tahap keempat proses data preprocessing adalah normalizing yang dilakukan untuk mengubah kata gaul menjadi kata baku, mengubah penulisan kata menjadi kata yang benar dan memberikan spasi pada kata yang tergabung. Daftar kata yang digunakan untuk normalisasi dibuat secara manual dengan menyesuaikan dataset.

Tabel 6 Contoh Normalizing

Tweet Sebelum Normalizing	Tweet Sesudah Normalizing
['banjir', 'menyakiti', 'banyak', 'warga', 'ditambah', 'gubernur', 'dan', 'pemerintah', 'tidak', 'bisa', 'kerja', 'dalam', 'mengatasi', 'banjir', 'pun', 'demikian', 'dengan', 'jabar', 'amp', 'banten']	['banjir', 'menyakiti', 'banyak', 'warga', 'ditambah', 'gubernur', 'dan', 'pemerintah', 'tidak', 'bisa', 'kerja', 'dalam', 'mengatasi', 'banjir', 'pun', 'demikian', 'dengan', 'jabar', 'amp', 'banten']
['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'melakukan', 'penyisiran', 'untuk', 'mencari', 'warga', 'yang', 'memerlukan', 'evakuasi', 'di', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']	['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'melakukan', 'penyisiran', 'untuk', 'mencari', 'warga', 'yang', 'memerlukan', 'evakuasi', 'di', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']

5. Stopword Removal

Tahap kelima proses data preprocessing adalah stopwords removal yang dilakukan untuk menghapus kata yang dianggap tidak memiliki makna. Daftar kata yang digunakan untuk stopwords removal dibuat secara manual dengan menyesuaikan dataset.

Tabel 7 Contoh Stopword Removal

Tweet Sebelum Stopword Removal	Tweet Sesudah Stopword Removal
['banjir', 'menyakiti', 'banyak', 'warga', 'ditambah', 'gubernur', 'dan', 'pemerintah', 'tidak', 'bisa', 'kerja', 'dalam', 'mengatasi', 'banjir', 'pun', 'demikian', 'dengan', 'jabar', 'amp', 'banten']	['banjir', 'menyakiti', 'warga', 'ditambah', 'gubernur', 'pemerintah', 'kerja', 'mengatasi', 'banjir', 'jabar', 'banten']
['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'melakukan', 'penyisiran', 'untuk', 'mencari', 'warga', 'yang', 'memerlukan', 'evakuasi', 'di', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']	['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'penyisiran', 'mencari', 'warga', 'evakuasi', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']

6. Stemming

Tahap keenam proses data preprocessing adalah stemming yang dilakukan untuk mengubah kata dengan imbuhan menjadi kata dasar.

Tabel 8 Contoh Stemming

Tweet Sebelum Stemming	Tweet Sesudah Stemming
['banjir', 'menyakiti', 'warga', 'ditambah', 'gubernur', 'pemerintah', 'kerja', 'mengatasi', 'banjir', 'jabar', 'banten']	['banjir', 'sakit', 'warga', 'tambah', 'gubernur', 'perintah', 'kerja', 'atas', 'banjir', 'jabar', 'banten']
['tim', 'social', 'disaster', 'rescue', 'sekolah', 'relawan', 'langsung', 'penyisiran', 'mencari', 'warga', 'evakuasi', 'kawasan', 'jatiwaringin', 'kota', 'bekasi', 'jawa', 'barat']	['tim', 'social', 'disaster', 'rescue', 'sekolah', 'rawan', 'langsung', 'sisir', 'cari', 'warga', 'evakuasi', 'kawasan', 'jatiwaringin', 'kota', 'bekas', 'jawa', 'barat']

Feature Extraction

Proses feature extraction dilakukan menggunakan program dengan bahasa pemrograman Python dan library sklearn feature extraction text. Metode yang digunakan adalah Term Frequency – Inverse Document Frequency (TF-IDF). Perhitungan term yang digunakan adalah trigram dan max features yang digunakan bernilai 10.000. Untuk menentukan max features dilakukan percobaan perhitungan terlebih dahulu, hasil dari percobaan perhitungan ini menghasilkan features sebanyak 27.000. Karena features yang didapatkan terlalu banyak, maka ditetapkan max features bernilai 10.000. Hasil Term Frequency – Inverse Document Frequency (TF-IDF) adalah vektor dari dataset.

Tabel 9 Contoh Feature Extraction

Kata	TF		DF	IDF	IDF+1	TF x (IDF+1)	
	D1	D2				D1	D2
Banjir	2	0	1	0.301	1.301	2.602	0
Evakuasi	0	1	1	0	1	0	1
Warga	1	1	2	0	1	1	1
Gubernur	1	0	1	0	1	1	0
Jabar	1	0	1	0	1	1	0

Proses Pelatihan

Proses pelatihan dataset dilakukan menggunakan program dengan bahasa pemrograman Python, library sklearn model selection, library keras dan library tensorflow. Metode yang digunakan adalah K-Fold Cross Validation dan Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation. Data train yang digunakan berisi 2.400 tweet dengan 71.50% tweet untuk label positif dan 28.50% tweet untuk label negatif. Proses pelatihan menghasilkan nilai Train Loss dan Train Accuracy.

Proses Pengujian

Proses pelatihan dataset dilakukan menggunakan program dengan bahasa pemrograman Python, library sklearn model selection, library keras dan library tensorflow. Metode yang digunakan adalah K-Fold Cross Validation dan Jaringan Saraf Tiruan (JST) model Multi Layer Perceptron (MLP) dengan algoritma Backpropagation. Data test yang digunakan berisi 600 tweet dengan 73.33% tweet untuk label positif dan 26.67% tweet untuk label negatif. Proses pelatihan menghasilkan nilai Test Loss dan Test Accuracy.

Evaluasi

Proses evaluasi dataset dilakukan menggunakan program dengan bahasa pemrograman Python dan library sklearn metrics. Metode yang digunakan adalah Confusion Matrix. Evaluasi menghasilkan nilai dari Accuracy, Precision, Recall, F1-Score dan Confusion Matrix.

4. Evaluasi

Skenario Pengujian

Model yang dibangun memiliki 10 skenario berbeda, setiap skenario memiliki data shuffling, learning rate, node hidden layer dan drop out yang berbeda. Model yang digunakan adalah Sequential. Input layer memiliki 3.000 node dengan 10.000 features. Hidden layer memiliki 64 node atau 128 node dengan activation relu. Output layer memiliki 1 node dengan activation sigmoid. Model menggunakan optimizer adam dengan learning rate yang berbeda, loss binary cross entropy dan metrics accuracy. Untuk setiap model dilakukan 5 kali fold, untuk setiap fold dilakukan 5 kali epoch.

Tabel 10 Skenario Pengujian

Model	Skenario			
	Data Shuffling	Learning Rate	Node Hidden Layer	Drop Out
ANN_01	X	0.001	64	X
ANN_02	X	0.005	64	X
ANN_03	X	0.001	64	0.2
ANN_04	O	0.001	64	X
ANN_05	O	0.001	64	0.2
ANN_06	O	0.005	64	X
ANN_07	O	0.1	64	X
ANN_08	X	0.001	128	X
ANN_09	O	0.001	128	X
ANN_10	O	0.001	128	0.2

Hasil Pengujian

Tabel 11 Hasil Rata-Rata Pengujian

Model	Hasil Rata-Rata					
	Train Accuracy	Test Accuracy	Precision	Recall	F1-Score	Accuracy
ANN_01	99.16%	72.66%	77.38%	88.63%	82.62%	72.66%
ANN_02	99.33%	73.50%	77.82%	89.31%	83.17%	73.50%
ANN_03	99.00%	73.33%	77.77%	89.09%	83.05%	73.33%
ANN_04	99.08%	73.16%	76.90%	90.13%	82.99%	73.16%
ANN_05	99.00%	73.33%	78.34%	87.95%	82.86%	73.33%
ANN_06	99.29%	71.83%	75.04%	90.37%	82.00%	71.83%
ANN_07	99.29%	67.83%	73.06%	84.54%	78.38%	67.83%
ANN_08	99.29%	73.83%	77.79%	90.00%	83.45%	73.83%
ANN_09	99.41%	69.16%	77.37%	85.11%	81.06%	71.49%
ANN_10	99.33%	71.49%	76.87%	86.92%	81.59%	71.50%

Analisis Hasil Pengujian

Dari tabel 10 dan tabel 11 dapat disimpulkan bahwa skenario kedelapan merupakan skenario terbaik dengan memberikan hasil akurasi yaitu 73.83% dengan tanpa data shuffling, learning rate bernilai 0.001, node hidden layer sebanyak 128, dan tanpa drop out. Pada data test terdapat 430 tweet untuk label positif dan 170 tweet untuk label negatif. Setelah proses pengujian, didapatkan hasil klasifikasi dari model yaitu 509 tweet untuk label positif dan 91 tweet untuk label negatif.

Skenario kesembilan memiliki nilai Train Accuracy tertinggi dengan persentase 99.41%. Skenario kedelapan memiliki nilai Test Accuracy tertinggi dengan persentase 73.83% dan nilai F1-Score tertinggi dengan persentase 83.45%. Skenario kelima memiliki nilai Precision tertinggi dengan persentase 78.34%. Skenario keenam memiliki nilai Recall tertinggi dengan persentase 90.37%.

Dari seluruh hasil pengujian dapat diketahui bahwa data shuffling, learning rate, node hidden layer dan drop out berpengaruh terhadap nilai akurasi model yang dibangun. Hasil akurasi setiap model dipengaruhi oleh jumlah data yang terbatas, pelabelan data yang kurang akurat, pelabelan data yang kurang seimbang atau tweet dalam dataset terlalu bebas dan unik.

5. Kesimpulan

Berdasarkan hasil pengujian dan analisis hasil pengujian yang telah dilakukan, maka dapat disimpulkan bahwa:

1. Skenario kedelapan merupakan skenario terbaik dengan memberikan hasil akurasi yaitu 73.83% dengan tanpa data shuffling, learning rate bernilai 0.001, node hidden layer sebanyak 128 dan tanpa drop out.
2. Penggunaan data shuffling berpengaruh terhadap nilai akurasi, pada skenario kesatu tanpa data shuffling mendapat nilai akurasi 72.66% dan pada skenario keempat dengan data shuffling mendapat nilai akurasi 73.16%.
3. Nilai learning rate berpengaruh terhadap nilai akurasi, pada skenario kesatu dengan learning rate bernilai 0.001 mendapat nilai akurasi 72.66% dan pada skenario kedua dengan learning rate bernilai 0.005 mendapat nilai akurasi 73.50%.
4. Jumlah node hidden layer berpengaruh terhadap nilai akurasi, pada skenario kesatu dengan node hidden layer sebanyak 64 mendapat nilai akurasi 72.66% dan pada skenario kedelapan dengan node hidden layer sebanyak 128 mendapat nilai akurasi 73.83%.
5. Penggunaan drop out berpengaruh terhadap nilai akurasi, pada skenario kesatu tanpa drop out mendapatkan nilai akurasi 72.66% dan pada skenario ketiga dengan drop out mendapatkan nilai akurasi 73.33%.

Referensi

- [1] Achmad, A., Ilham, A. A., & Herman. (2012). Implementasi Algoritma Term Frequency – Inverse Document Frequency dan Vector Space Model untuk Klasifikasi Dokumen Naskah Dinas. *Forum Pendidikan Tinggi Teknik Elektro Indonesia (FORTEI) 2012*, 88.
- [2] Acuna, E. (2010). Preprocessing in Data Mining. 2.
- [3] Agustin, M. (2012). *PENGUNAAN JARINGAN SYARAF TIRUAN BACKPROPAGATION UNTUK SELEKSI PENERIMAAN MAHASISWA BARU PADA JURUSAN TEKNIK KOMPUTER DI POLITEKNIK NEGERI SRIWIJAYA*. Semarang: Universitas Diponegoro.
- [4] Amalia, C., & Sibaroni, Y. (2020). *ANALISIS SENTIMEN DATA TWEET MENGGUNAKAN MODEL JARINGAN SARAF TIRUAN DENGAN PEMBOBOTAN DELTA TF-IDF*. Bandung: Telkom University.
- [5] Assuja, M. A., & Saniati. (2016). ANALISIS SENTIMEN TWEET MENGGUNAKAN BACKPROPAGATION NEURAL NETWORK. *Jurnal TEKNOINFO*, 24.
- [6] *Banjir*. (2007). Diambil kembali dari wikipedia.org: <https://id.wikipedia.org/wiki/Banjir>
- [7] Clinton, B. (2019, Oktober 30). *Pengguna Aktif Harian Twitter Indonesia Diklaim Terbanyak*. Diambil kembali dari kompas.com: <https://tekno.kompas.com/read/2019/10/30/16062477/pengguna-aktif-harian-twitter-indonesia-diklaim-terbanyak>
- [8] *Confusion matrix*. (2010). Diambil kembali dari wikipedia: https://en.wikipedia.org/wiki/Confusion_matrix
- [9] Kankanamge, N., Yigitcanlar, T., Goonetilleke, A., & Kamruzzaman, M. (2019). Determining disaster severity through social media analysis: Testing the methodology with South East Queensland Flood tweets. *International Journal of Disaster Risk Reduction*, 13.
- [10] Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. Morgan & Claypool Publishers.
- [11] Menarianti, I. (2015). KLASIFIKASI DATA MINING DALAM MENENTUKAN PEMBERIAN KREDIT BAGI NASABAH KOPERASI. *Jurnal Ilmiah Teknosains*.
- [12] Muzakki, M. F., Jondri, & Umbara, R. F. (2019). *ANALISIS SENTIMEN MAHASISWA TERHADAP FASILITAS UNIVERSITAS TELKOM MENGGUNAKAN METODE JARINGAN SARAF TIRUAN DAN TF-IDF*. Bandung: Universitas Telkom.
- [13] Neppalli, V. K., Caragea, C., Squicciarini, A., Tapia, A., & Stehle, S. (2017). Sentiment analysis during Hurricane Sandy in emergency response. *International Journal of Disaster Risk Reduction*, 10.
- [14] Septian, J. A., Fahrudin, T. M., & Nugroho, A. (2019). Analisis Sentimen Pengguna Twitter Terhadap Polemik Persepakbolaan Indonesia Menggunakan Pembobotan TF-IDF dan K-Nearest Neighbor. *Journal of Intelligent Systems and Computation*, 44.
- [15] Siregar, A. M., & Hasan, T. A. (2018). APLIKASI LINIER REGRESI DENGAN ALGORITMA JARINGAN SYARAF TIRUAN UNTUK SENTIMEN ANALISIS. *Jurnal Ilmu Komputer dan Teknologi Informasi*, 43.
- [16] *Social Media*. (2006). Diambil kembali dari wikipedia.org: https://en.wikipedia.org/wiki/Social_media
- [17] *Twitter*. (2007). Diambil kembali dari wikipedia.org: <https://en.wikipedia.org/wiki/Twitter>
- [18] Zannah, R. (2019). *ANALISIS SENTIMEN PADA MEDIA SOSIAL TWITTER UNTUK KLASIFIKASI OPINI ISLAM RADIKAL MENGGUNAKAN JARINGAN SARAF TIRUAN*. Surabaya: Universitas Islam Negeri Sunan Ampel.