

## ANALISIS SENTIMEN PENGGUNA TWITTER TERHADAP PERKEMBANGAN ARTIFICIAL INTELLIGENCE DENGAN PENERAPAN ALGORITMA LONG SHORT-TERM MEMORY (LSTM)

Ahmad Azrul<sup>1</sup>, Ade Irma Purnamasari<sup>2</sup>, Irfan Ali<sup>3</sup>

<sup>1,2</sup> Teknik Informatika, STMIK IKMI Cirebon

<sup>3</sup> Rekayasa Perangkat Lunak, STMIK IKMI Cirebon

Jl Perjuangan No.10B, Kecamatan Kesambi, Kota Cirebon, Indonesia

azrdneperone@gmail.com

### ABSTRAK

Dalam era transformasi teknologi yang pesat, perkembangan kecerdasan buatan (*Artificial Intelligence* atau AI) menjadi salah satu tonggak utama yang membentuk pandangan masyarakat global. Meskipun perkembangan *Artificial Intelligence* (AI) menghadirkan berbagai potensi positif, permasalahan yang muncul adalah kurangnya pemahaman mendalam terhadap pandangan dan sentimen pengguna Twitter terhadap evolusi teknologi ini. Tujuan penelitian ini adalah untuk menganalisis sentimen pengguna Twitter terhadap perkembangan AI dan mengevaluasi kinerja model LSTM dalam menginterpretasikan opini menggunakan algoritma *Long Short-Term Memory* (LSTM) dengan memanfaatkan metode *Knowledge Discovery in Database* (KDD). Hasil analisis sentimen dari 1513 *tweet* menunjukkan bahwa 999 di antaranya bersentimen positif dan 514 bersentimen negatif terhadap perkembangan AI. Melalui uji coba, model LSTM mencapai tingkat akurasi tertinggi sebesar 74,25% pada *epoch* 20 dengan 100 *neuron* pada lapisan LSTM. Pengujian dengan *split data* 80:20 menghasilkan kinerja terbaik dengan 1210 data *training* dan 303 data *testing*. Kesimpulan penelitian ini menegaskan bahwa penerapan LSTM pada data *tweet* AI memberikan kinerja yang baik, dengan rata-rata akurasi 74.25%, *precision* 80.29%, *recall* 81.09%, dan *f1-score* 80.69%. Hasil ini menunjukkan kemampuan model LSTM dalam efektif memahami perasaan pengguna Twitter terhadap perkembangan AI.

**Kata kunci :** Analisis Sentimen, *Artificial Intelligence*, Perkembangan AI, Twitter, *Long Short-Term Memory* (LSTM)

### 1. PENDAHULUAN

Dalam era digital yang terus berkembang pesat, teknologi *Artificial Intelligence* (AI) menjadi inovasi signifikan yang berpotensi mengubah berbagai aspek kehidupan manusia [1]. Kecerdasan Buatan (*Artificial Intelligence*) memungkinkan komputer mereplikasi aktivitas manusia dan melaksanakan tugas yang memerlukan kecerdasan manusia [2]. Perkembangan teknologi AI diantisipasi akan terus berdampak mendalam, terutama dalam sektor bisnis dan industri [3]. Twitter, sebagai platform media sosial aktif dengan pertumbuhan pengguna yang pesat, menjadi sumber data kaya untuk memahami pandangan dan perasaan pengguna terkait perkembangan AI. Penentuan permasalahan melibatkan upaya untuk memahami lebih mendalam bagaimana pandangan pengguna Twitter terhadap AI mengalami evolusi seiring berjalannya waktu. Beberapa pertanyaan yang perlu dijawab mencakup apakah sikap pengguna cenderung bersifat positif, negatif, atau netral terhadap perkembangan teknologi ini, serta faktor-faktor apa yang memiliki pengaruh terhadap perubahan sentimen mereka.

Dalam penelitian ini, pendekatan metodologi untuk menganalisis sentimen pengguna Twitter terhadap kemajuan *Artificial Intelligence* (AI) melibatkan penerapan algoritma *Long Short-Term Memory* (LSTM). Tujuan utamanya adalah mengevaluasi respons dan pandangan pengguna

terhadap perkembangan teknologi AI di platform media sosial Twitter [4]. Selain itu, penelitian ini berusaha menjelajahi faktor-faktor yang dapat mempengaruhi perubahan sentimen pengguna seiring waktu. Langkah selanjutnya adalah mengidentifikasi pola sentimen positif dan negatif dalam *tweet* yang terkait dengan AI, sambil menilai evaluasi kinerja algoritma LSTM untuk mendukung pengambilan keputusan yang lebih efektif dalam penerapan dan perkembangan teknologi ini.

Penentuan untuk menggunakan algoritma LSTM, yang termasuk dalam keluarga algoritma *neural network*, dipilih karena kemampuannya yang unik dalam mengolah data sekuensial sambil tetap memperhatikan informasi dalam jangka panjang. Proses analisis sentimen dimulai dengan mengumpulkan data dari platform Twitter menggunakan teknik *web scraping/crawling* dan API Twitter. Data yang terhimpun mencakup *tweet* yang mengandung opini dan pandangan pengguna terkait perkembangan AI. Data ini kemudian melewati serangkaian langkah pengolahan, seperti pembersihan data dan vektorisasi, untuk mempersiapkannya sebelum dianalisis oleh model LSTM. Model ini dilatih menggunakan teknik *supervised learning* untuk memahami pola sentimen dalam data sekuensial, menghasilkan *output* analisis sentimen yang dapat diinterpretasikan [5]. Penggunaan LSTM diharapkan dapat meningkatkan akurasi analisis sentimen,

terutama dalam menangani konteks jangka panjang dan ketergantungan temporal dalam data teks.

Pentingnya penelitian ini karena memberikan kontribusi pada pemahaman persepsi publik terhadap kemajuan AI, memberikan wawasan mendalam tentang respons pengguna terhadap perkembangan AI di Twitter melalui pendekatan analisis sentimen yang efisien dan andal.

## 2. TINJAUAN PUSTAKA

Tinjauan pustaka ini bertujuan untuk menyajikan kerangka konseptual yang mendukung pemahaman dan analisis terhadap perkembangan *Artificial Intelligence* (AI) serta respons pengguna Twitter terhadapnya.

### 2.1. Analisis Sentimen

Analisis sentimen merupakan bagian integral dari pengolahan teks yang terutama berfokus pada penilaian dokumen teks. Ini adalah suatu metode untuk mengolah data tekstual guna mendapatkan informasi yang terkandung dalam teks [6]. Selain itu, analisis sentimen dapat diartikan sebagai bidang ilmu yang mempelajari opini, sentimen, evaluasi, sikap, dan emosi seseorang yang diungkapkan secara tulisan dan mengevaluasi pendapat masyarakat melalui tulisan atau dari berbagai topik yang relevan [7].

### 2.2. Twitter

Twitter adalah salah satu platform media sosial yang memungkinkan pengguna untuk menyampaikan dan berbagi informasi tentang segala hal yang terjadi. Dengan API Twitter, Twitter juga memudahkan pengguna untuk mengakses data [4]. Setiap hari, pengguna Twitter mengirimkan posting, atau *tweet*, yang menyampaikan minat politik, ideologis, dan pendapat mereka serta fakta dan pendapat tentang produk atau layanan pemerintah yang mereka gunakan. Hal ini menghasilkan tingginya ketersediaan data di Twitter, yang menjadikannya tempat yang bagus untuk melakukan *opinion mining* atau analisis sentimen. Untuk studi politik, pemasaran, dan sosial, gudang data Twitter ini sangat berguna [8].

### 2.1. Text Mining

*Text mining* adalah metode pengumpulan dan pengestrakan data yang mirip dengan *Natural Language Processing* (NLP) [9]. *Text mining*, dikenal juga sebagai penambahan teks atau *text preprocessing*, merujuk pada proses pembersihan data atau area file yang tidak digunakan selama proses *training*, *validasi*, dan *pengujian* sehingga meningkatkan kualitas data, mengurangi kompleksitas dapat mengubah data tersebut menjadi informasi yang memiliki nilai tambah [10].

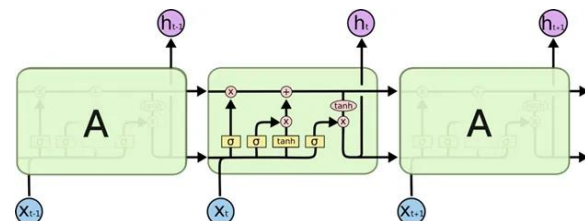
### 2.2. Feature Extraction

*Feature extraction* dalam konteks pengolahan bahasa alami, seperti analisis sentimen pada data teks, melibatkan transformasi kata-kata menjadi representasi vektor. Ekstraksi fitur ini terdiri dari dua tahap. Ekstraksi primer mengubah data ke format standar, dan kemudian diubah menjadi bentuk vektor [9]. Vektor merupakan kata terdistribusi yang kuat, dapat digunakan untuk prediksi suatu kata dan terjemahan [4].

### 2.3. Long Short-Term Memory (LSTM)

*Long Short-Term Memory* (LSTM) adalah Salah satu jenis arsitektur jaringan saraf tiruan (RNN) yang dapat mengatasi masalah *vanishing gradient* pada jaringan saraf biasa [11]. LSTM juga memiliki kemampuan untuk mengikat informasi dalam jangka panjang. Ini karena LSTM memiliki kelebihan dalam analisis sentimen, seperti kemampuan untuk mengingat informasi dari data yang telah diproses dan menggunakan kemampuan memori untuk memahami data *input* dengan lebih baik. [5].

LSTM memperkenalkan sebuah unit memori yang disebut "*cell*" dan menggunakan mekanisme gerbang, yang memberikan kemampuan sel untuk mengontrol aliran informasi [4]. Mekanisme ini terdiri dari tiga *gate* yaitu:

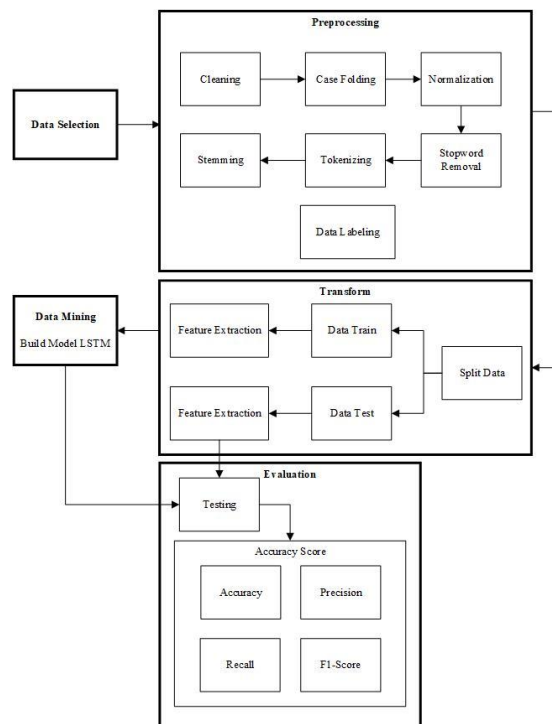


Gambar 1. Struktur LSTM

- Gate Forgotten*, juga dikenal sebagai *Gate Forgotten*, menentukan jumlah informasi yang harus dihapus dari sel memori.
- Gate Input*, juga dikenal sebagai *Gate Input*, menentukan jumlah informasi baru yang perlu disimpan dalam sel memori.
- Gate Output*, Menentukan jumlah informasi yang akan dikeluarkan sebagai *output* sel.

## 3. METODE PENELITIAN

Penelitian ini menggunakan metode KDD (*Knowledge Discovery in Database*) karena memiliki keunggulan dalam proses pengidentifikasian pola yang terorganisir dari sekumpulan data yang kompleks, sehingga data menjadi mudah dipahami. Alur analisis data yang menerapkan metode KDD digambarkan di Gambar 2.



Gambar 2. Analisis Data

Berikut beberapa penjelasan mengenai tahapan analisis data yang direpresentasikan, Langkah-langkahnya adalah sebagai berikut.

### 3.1. Data Selection

*Data Selection* merupakan tahap awal dalam *Knowledge Discovery in Databases* (KDD), di mana dilakukan pengumpulan data dan pemilihan data dari berbagai kumpulan data operasional untuk keperluan analisis atau pengolahan lanjutan.

### 3.2. Preprocessing

*Preprocessing* adalah langkah pengelolaan struktur data sebelum dilakukan klasifikasi. Tahap *preprocessing* ini bertujuan untuk mengubah kata-kata yang semula bersifat tidak terstruktur menjadi bentuk standar, sehingga dapat membantu dalam ekstraksi informasi dari ulasan, dengan melakukan *preprocessing* ini, data *tweet* dapat diperoleh dalam keadaan yang lebih bersih sebelum diproses lebih lanjut pada langkah-langkah selanjutnya. Adapun tahapan dari perbaikan kesalahan penulisan sebagai berikut yaitu *cleaning*, *case folding*, *normalization*, *stopword removal*, *tokenization*, *stemming* dan *labeling*.

### 3.3. Transformation

Transformasi adalah istilah yang merujuk pada rangkaian proses yang digunakan untuk mengubah data menjadi bentuk yang dapat diolah pada tahap *data mining*. Sebagai bagian dari analisis sentimen menggunakan data *tweet* dalam lingkungan *Python*, langkah awal melibatkan pemuatan dan pembagian data *tweet* (*split data*) menjadi dua kelompok: data *training* dan data *testing*. Proses selanjutnya

melibatkan ekstraksi fitur, juga dikenal sebagai ekstraksi ciri, dari data teks, yang kemudian diubah menjadi fitur yang dapat dimanfaatkan oleh algoritma klasifikasi. Penelitian ini akan memanfaatkan *tokenizer* dan *padding* untuk membangun lapisan *embedding* ke dalam lapisan LSTM, serta menggunakan vektor pengenalan kata terlatih *fasttext* untuk *neural network*.

### 3.4. Data Mining

*Data mining* adalah proses untuk menemukan pola unik dalam kumpulan data yang besar dan kompleks, memiliki tujuan untuk menggali informasi dari sumber data atau kumpulan dokumen yang tidak terorganisir dengan memanfaatkan pola yang dapat memberikan bantuan dalam pengambilan keputusan dan memberikan wawasan yang signifikan. Dalam penelitian ini, model *neural network*, khususnya LSTM, digunakan sebagai alat untuk memahami struktur dan konteks sentimen dalam data. *Long Short-Term Memory* (LSTM) merupakan suatu tipe arsitektur jaringan saraf tiruan yang secara khusus dirancang untuk mengatasi permasalahan *vanishing gradient* yang muncul pada jaringan saraf konvensional dan memiliki kemampuan untuk menangani pengikatan informasi dalam jangka panjang. *Data mining* mencakup ekstraksi pengetahuan dan pola sentimen dari data *tweet*, yang secara substansial mendukung pengambilan keputusan terkait analisis sentimen pada model LSTM.

### 3.5. Evaluation

Evaluasi merupakan tahap krusial dalam pengembangan model atau sistem, di mana kinerja suatu solusi atau model diukur dan dinilai. Evaluasi dalam konteks *data mining* dan *machine learning* merujuk pada proses menilai dan mengukur kinerja suatu model atau algoritma terhadap suatu set data tertentu. Tujuan utama evaluasi adalah mengukur sejauh mana model mampu melakukan generalisasi dari data pelatihan ke data baru. Dalam tahap evaluasi penelitian ini, disediakan informasi rinci mengenai kinerja model klasifikasi pada data uji, termasuk nilai-nilai seperti akurasi, *precision*, *recall*, *f1-score*, dan dukungan, yang digunakan untuk mengevaluasi hasil *neural network*, terutama LSTM. Metrik-metrik ini mencerminkan kemampuan model dalam mengklasifikasikan sentimen dengan tepat, dan juga memberikan wawasan lebih mendalam tentang cara kerja model tersebut.

## 4. HASIL DAN PEMBAHASAN

Hasil dan pembahasan ini akan menjelaskan dari proses pengambilan data, termasuk data yang telah melalui tahap *preprocessing* dan menunjukkan keterkaitan kata pada setiap kelas hasil *labeling*. Selain itu, akan ditampilkan grafik dari hasil pelatihan data dan uji coba parameter untuk mencapai model yang optimal, serta hasil evaluasi yang diperoleh melalui representasi hasil *confusion matrix*.

#### 4.1. Data Selection

*Data selection* ini melibatkan proses autentikasi API Twitter untuk pengambilan data. Setelah terhubung, peneliti dapat pengambilan data dari platform media sosial Twitter dengan cara *crawling* data *tweet*. Pengambilan data *tweet* ini dilakukan dengan menggunakan metode *tweet* harvest dalam bahasa pemrograman *Python* di Google Colaboratory. Proses pengambilan data *tweet* ini dilakukan dengan menggunakan kata kunci "perkembangan AI" dan berlangsung dari tanggal 01 Januari 2019 hingga 20 November 2023, dengan jumlah maksimal data *tweet* yang diperoleh mencapai 1513 data *tweet*. tahap selanjutnya yaitu data *selection*. Pada data *selection* ini dilakukan pemilihan atribut yang akan digunakan dalam proses analisis sentimen dan hanya terdapat satu atribut yang dipakai untuk tahap *preprocessing*, yaitu atribut *full\_text*.

#### 4.2. Preprocessing

*Preprocessing* merupakan tahapan yang bertujuan untuk mengeliminasi gangguan atau data yang tidak relevan, seperti emoji, URL, angka, dan juga untuk menjadikan format kata seragam. Dengan melakukan *preprocessing* ini, data *tweet* dapat diperoleh dalam keadaan yang lebih bersih sebelum diproses lebih lanjut pada langkah-langkah selanjutnya.

##### 4.2.1. Cleaning

*Cleaning* adalah tindakan membersihkan kata-kata yang tidak relevan, seperti nama pengguna, tanda pagar, URL, tautan, tanda baca, angka, ruang kosong, dan emoji. Analisis Sentimen, *Artificial Intelligence*, Perkembangan AI, Twitter, *Long Short-Term Memory* (LSTM)

Tabel 1. Cleaning

| Sebelum                                                                                                                                                                                                     | Sesudah                                                                                                                                                                                      |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Lohhhh harus gitu biar keren, ikutin perkembangan zaman, dan canggih pokoknya teknologi terkini uh keren face recognition uh apatuh wih AI iya AI big data future economics taikucing 2045 uh keren @KAI121 | Lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren face recognition uh apatuh wih AI iya AI big data future economics taikucing uh keren |

##### 4.2.2. Case Folding

*Case Folding* adalah proses mengubah huruf besar atau kapital menjadi huruf kecil, sehingga teks *tweet* dapat memiliki format yang seragam.

Tabel 2. Case Folding

| Sebelum                                                                                                                           | Sesudah                                                                                                                           |
|-----------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------|
| Lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren face recognition uh apatuh | lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren face recognition uh apatuh |

| Sebelum                                                    | Sesudah                                                    |
|------------------------------------------------------------|------------------------------------------------------------|
| wih AI iya AI big data future economics taikucing uh keren | wih ai iya ai big data future economics taikucing uh keren |

##### 4.2.3. Normalization

*Normalization* adalah tahap di mana mengubah dari kata yang tidak baku menjadi kata yang baku sesuai dengan KBBI.

Tabel 3. Normalization

| Sebelum                                                                                                                                                                                      | Sesudah                                                                                                                                                                                       |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren face recognition uh apatuh wih ai iya ai big data future economics taikucing uh keren | lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren wajah recognition uh apatuh wih ai iya ai big data future economics taikucing uh keren |

##### 4.2.4. Stopword Removal

*Stopword Removal* sebuah teknik untuk mengurangi jumlah kata yang terdiri dari kata hubung tetapi tidak memiliki makna signifikan atau memberikan kontribusi yang kurang penting.

Tabel 4. Stopword Removal

| Sebelum                                                                                                                                                                                       | Sesudah                                                                                                                                                                             |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| lohhhh harus gitu biar keren ikutin perkembangan zaman dan canggih pokoknya teknologi terkini uh keren wajah recognition uh apatuh wih ai iya ai big data future economics taikucing uh keren | lohhhh gitu biar keren ikutin perkembangan zaman canggih pokoknya teknologi terkini uh keren wajah recognition uh apatuh wih ai iya ai big data future economics taikucing uh keren |

##### 4.2.5. Tokenization

*Tokenization* adalah proses yang menggunakan karakter pemisah kata untuk memecah teks menjadi satuan kata dalam *tweet*.

Tabel 5. Tokenization

| Sebelum                                                                                                                                                                             | Sesudah                                                                                                                                                                                                                                                                  |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| lohhhh gitu biar keren ikutin perkembangan zaman canggih pokoknya teknologi terkini uh keren wajah recognition uh apatuh wih ai iya ai big data future economics taikucing uh keren | ['lohhhh', 'gitu', 'biar', 'keren', 'ikutin', 'perkembangan', 'zaman', 'canggih', 'pokoknya', 'teknologi', 'terkini', 'uh', 'keren', 'wajah', 'recognition', 'uh', 'apatuh', 'wih', 'ai', 'iya', 'ai', 'big', 'data', 'future', 'economics', 'taikucing', 'uh', 'keren'] |

##### 4.2.6. Stemming

*Stemming* adalah langkah di mana kata-kata yang memiliki imbuhan diubah menjadi bentuk kata dasar dengan menghilangkan afiksnya.

Tabel 6. *Stemming*

| Sebelum                                                                                                                                                                                                                                                                  | Sesudah                                                                                                                                                                                          |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| ['lohhhh', 'gitu', 'biar', 'keren', 'ikutin', 'perkembangan', 'zaman', 'canggih', 'pokoknya', 'teknologi', 'terkini', 'uh', 'keren', 'wajah', 'recognition', 'uh', 'apatuh', 'wih', 'ai', 'iya', 'ai', 'big', 'data', 'future', 'economics', 'taikucing', 'uh', 'keren'] | lohhhh gitu biar keren<br>ikutin kembang zaman<br>canggih pokok<br>teknologi kini uh<br>keren wajah<br>recognition uh apatuh<br>wih ai iya ai big data<br>future economics<br>taikucing uh keren |

#### 4.2.7. Labeling

*Labeling* adalah langkah di mana kalimat diberikan label sebagai negatif atau positif, proses ini menggunakan *lexicon based*. Langkah yang dilakukan adalah menghitung sentimen yang terkandung dalam setiap *tweet* berdasarkan hasil perhitungan skor sentimen.

Tabel 7. *Labeling*

| Text                                                                                                                                                                                                      | Positif | Negatif | Total | Label   |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|---------|-------|---------|
| lohhhh gitu biar<br>keren ikutin<br>kembang zaman<br>canggih pokok<br>teknologi kini uh<br>keren wajah<br>recognition uh<br>apatuh wih ai iya<br>ai big data future<br>economics<br>taikucing uh<br>keren | 10      | -9      | 1     | Positif |

#### 4.3. Transformation

Transformasi adalah suatu proses perubahan atau konversi bentuk atau struktur data menjadi bentuk atau struktur yang berbeda. Dalam konteks analisis data atau pemrosesan bahasa alami, transformasi seringkali merujuk pada serangkaian langkah-langkah yang diterapkan pada data mentah untuk mengubahnya menjadi bentuk yang lebih sesuai atau memiliki makna yang lebih tinggi untuk analisis atau pemodelan yang lebih lanjut.

##### 4.3.1. Split Data

Pembagian data menjadi data *training* dan data *testing* dilakukan dengan menggunakan fungsi *train\_test\_split*, di mana parameter *random\_state=42* digunakan untuk memastikan reproduktibilitas hasil. *Data training*, yang mencakup 80% dari total data sebanyak 1513 *tweet*, digunakan untuk melatih model dan membentuknya. Sebaliknya, data *testing* sejumlah 20% dari total data, yaitu 303 *tweet*, digunakan untuk menguji performa model yang telah dilatih.

#### 4.4. Feature Extraction

*Feature Extraction* melibatkan langkah-langkah mengidentifikasi dan mengekstrak daftar kata-kata dari data teks. Dalam penelitian ini, metode *Feature Extraction* melibatkan penggunaan Tokenizer untuk algoritma klasifikasi LSTM, serta pemanfaatan vektor

*embedding* kata yang telah di *pre-trained* dari *fasttext* untuk *Neural Network*. Pertama, parameter *num\_words* disesuaikan dengan jumlah 5000 kata, yang diperoleh dengan pembulatan dari jumlah kata umum dalam dataset sebanyak 4360 kata. Selanjutnya, Tokenizer digunakan untuk mengonversi teks *tweet* menjadi urutan bilangan bulat yang merepresentasikan indeks setiap kata. Hasil dari pemberian indeks pada data *tweet* dapat dilihat dalam Tabel 9.

Tabel 8. Hasil Pemberian Indeks pada Data Tweet

| Indeks Setiap Kata                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| [[1136, 1727, 1728, 205, 853, 46, 2, 47, 23, 114, 1136, 1729, 1730, 1],<br>[98, 2, 18, 59, 294, 335, 295, 1731, 296, 28, 65, 2, 1, 39, 5, 156, 1732, 20],<br>[115, 657, 541, 2, 1733, 1, 657, 190, 542, 854, 1734, 124, 855, 5, 168, 856, 1137],<br>...<br>[122, 1, 4352, 4353, 431, 2, 1, 93, 519, 104, 35, 512, 43, 2, 4, 1087, 1575, 1521, 20, 1]<br>[4354, 4355, 128, 1289, 679, 975, 129, 981, 4356, 1094, 591, 2, 3, 1, 325, 4357, 4358, 126],<br>[2, 1, 1, 304, 1, 4359, 129]]. |

Langkah selanjutnya adalah mentransformasi urutan kata menjadi serangkaian angka dengan panjang maksimal 200 (*max\_len*). Jika panjang urutan kata melebihi 200, akan dipotong, dan jika kurang, akan diisi dengan nilai nol. Proses ini dikenal sebagai *padding*, yang memastikan bahwa semua urutan kata memiliki panjang yang sama untuk konsistensi dalam pelatihan model dan pembentukan *word embedding*. Dalam penelitian ini, panjang *input\_length* ditetapkan sebanyak 32. Dengan demikian, vektor *pad sequence* yang dihasilkan dapat dilihat dalam Tabel 10.

Tabel 9. Hasil *Pad\_Sequence*

| Vektor Pad_Sequence                                                                                                                                                                                                           |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| array([[ 0, 0, 0, ..., 1729, 1730, 1],<br>[ 0, 0, 0, ..., 156, 1732, 20],<br>[ 0, 0, 0, ..., 168, 856, 1137],<br>...<br>[ 0, 0, 0, ..., 1521, 20, 1],<br>[ 0, 0, 0, ..., 4357, 4358, 126],<br>[ 0, 0, 0, ..., 1, 4359, 129]]) |

Pada tahap ini, terjadi pengenalan pada lapisan pertama yang dikenal sebagai lapisan *word embedding* atau *layer embedding*. Lapisan ini menerima sejumlah parameter yang sebelumnya telah diperoleh, seperti *input\_dim* sebanyak 4360 kata, *output\_dim* yang bervariasi antara 100 hingga 300 *neuron*, dan *input\_length* sebanyak 30. Hasilnya, vektor *word embedding* dengan *output\_dim* 32 direpresentasikan dalam Tabel 11. Vektor ini akan selanjutnya digunakan dalam lapisan LSTM pada tahapan selanjutnya.

Tabel 10. Hasil Word Embedding Layer

| Vektor Word Embedding                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| [[ -0.00472056 -0.0022824 -0.00128613 ...<br>0.01099967 -0.00395288 0.00674827]<br>[ 0.00672713 0.00950765 -0.03559843 ... 0.0190312<br>-0.01346661 0.01296655]<br>[ -0.0578333 0.03627495 0.03307132 ...<br>0.00680525 0.06648805 -0.00392976]<br>...<br>[ -0.03204157 0.0307051 -0.04918642 ... -<br>0.04162159 -0.04931731 -0.01754837]<br>[ 0.01868479 0.00382652 -0.00450106 ...<br>0.00576706 -0.01479268 0.00787947]<br>[ -0.03885082 0.00092512 0.02564813 ...<br>0.01225079 -0.01797752 -0.0170056 ]] |

#### 4.5. Data Mining

Salah satu pendekatan *data mining* yang diterapkan dalam penelitian ini adalah menggunakan model *Long Short-Term Memory* (LSTM). Model *Long Short-Term Memory* (LSTM) telah dirancang untuk mencapai tingkat akurasi optimal. Setelah menerima *output* dari lapisan *embedding*, langkah berikutnya melibatkan pelatihan model LSTM. Lapisan LSTM ini menggunakan 100 *neuron*, dan seluruh jaringan terhubung ke setiap *neuron* melalui lapisan *Fully Connected* dengan 1 unit, sesuai dengan jumlah kelas dalam penelitian ini. Tahap terakhir melibatkan penggunaan fungsi aktivasi *sigmoid*. Setelah seluruh model dibangun, konfigurasi dilakukan dengan menggunakan optimasi Adam dan *Binary Cross Entropy* untuk mengevaluasi nilai *loss* dari model yang telah dibentuk. Berdasarkan penjelasan parameter yang digunakan, dapat diterangkan bahwa arsitektur jaringan pada model yang terbentuk adalah sebagai berikut.

WARNING:absl:"lr" is deprecated in Keras optimizer, please use "learning\_rate"

Model: "sequential\_9"

| Layer (type)                         | Output Shape    | Param # |
|--------------------------------------|-----------------|---------|
| embedding_9 (Embedding)              | (None, 200, 32) | 139552  |
| spatial_dropout_9 (SpatialDropout1D) | (None, 200, 32) | 0       |
| lstm_9 (LSTM)                        | (None, 100)     | 53200   |
| dropout_9 (Dropout)                  | (None, 100)     | 0       |
| dense_9 (Dense)                      | (None, 1)       | 101     |

=====  
Total params: 192853 (753.33 KB)  
Trainable params: 192853 (753.33 KB)  
Non-trainable params: 0 (0.00 Byte)

Gambar 3. Total Parameter model LSTM

#### 4.6. Evaluation

*Data Selection* merupakan tahap awal dalam *Knowledge Discovery in Databases* (KDD), di mana dilakukan pengumpulan data dan pemilihan data dari berbagai kumpulan data operasional untuk keperluan analisis atau pengolahan lanjutan.

#### 4.7. Testing

Langkah terakhir sebelum dilakukan evaluasi adalah melakukan pengujian terhadap model yang telah dibangun. Pada tahap pengujian pertama, dilakukan analisis perbandingan antara jumlah data *training* dan data *testing*. Ada tiga model dengan rasio pembagian dataset yang dievaluasi dari jumlah total data *tweet* yang digunakan dalam penelitian ini adalah sebanyak 1513 data. Rincian pembagian data *training* dan data *testing* berdasarkan *tweet* yang digunakan dapat dilihat pada Tabel 12.

Tabel 11. Data Training dan Data Testing

| Perbandingan | Data Tweet    | Split Data | Jumlah |
|--------------|---------------|------------|--------|
| 90:10        | Data Training | 90%        | 1361   |
|              | Data Testing  | 10%        | 152    |
| 80:20        | Data Training | 80%        | 1210   |
|              | Data Testing  | 20%        | 303    |
| 70:30        | Data Training | 70%        | 1059   |
|              | Data Testing  | 30%        | 454    |

Pengujian dilanjutkan dengan menggunakan model LSTM untuk melakukan pelatihan dan pengujian, dengan evaluasi kinerja menggunakan metrik. Hasil pengujian dari beberapa model dengan variasi data *training* dan *testing* menunjukkan bahwa model dengan rasio 80:20, yang terdiri dari 1210 data *training* dan 303 data *testing*, memberikan tingkat akurasi tertinggi sebesar 73.92%, dengan *precision* 85.71%, *recall* 77.67%, dan *f1-score* 81.49%. Rincian perbandingan jumlah data *training* dan *testing* dapat ditemukan dalam Tabel 13.

Tabel 12. Hasil Jumlah Data Training dan testing

| Model | Acc    | Prec   | Recall | F1-Score |
|-------|--------|--------|--------|----------|
| 90:10 | 73.68% | 86.59% | 75.67% | 80.76%   |
| 80:20 | 73.92% | 85.71% | 77.67% | 81.49%   |
| 70:30 | 71.36% | 82.87% | 75.15% | 78.82%   |

Pada uji coba kedua, dilakukan pengujian terhadap jumlah *neuron* pada lapisan model LSTM dengan variasi 100, 200, dan 300 *neuron*. Pengujian ini bertujuan untuk menemukan jumlah *neuron* yang paling optimal dalam konstruksi model, dengan fokus mencapai akurasi tertinggi sebagaimana tercermin dalam hasil eksperimen. Hasil pengujian menunjukkan bahwa hanya jumlah *neuron* sebanyak 100 yang memberikan akurasi tertinggi, mencapai 73.92%. Selain itu, *precision* sebesar 85.71%, *recall* 77.67%, dan *f1-score* 81.49%. Rincian hasil pengujian jumlah *neuron* dapat ditemukan pada Tabel 14.

Tabel 13. Hasil Pengujian Jumlah Neuron

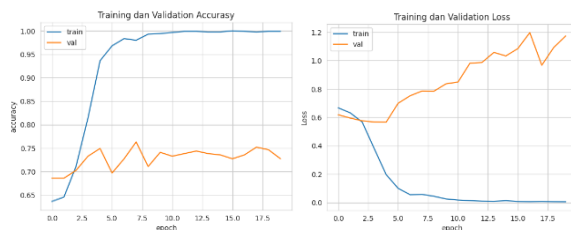
| Jumlah Neuron | Acc    | Prec   | Recall | F1-Score |
|---------------|--------|--------|--------|----------|
| 100           | 73.92% | 85.71% | 77.67% | 81.49%   |
| 200           | 72.93% | 85.22% | 76.88% | 81%      |
| 300           | 72.93% | 83.25% | 77.88% | 80.47%   |

Pada uji coba terakhir, dilakukan pengujian terhadap jumlah *epoch* pada model LSTM dengan Dropout, dengan variasi nilai *epoch* 10, 20, dan 30 untuk mengevaluasi performa model. Pemilihan jumlah *epoch* yang tepat sangat krusial untuk menghindari *overfitting* dan memastikan generalisasi model yang efektif terhadap data baru. Hasil menunjukkan bahwa pada *epoch* 20, model mencapai tingkat akurasi yang baik, mencapai 74.25%. Selain itu, *precision* sebesar 80.29%, *recall* 81.09%, dan *f1-score* 80.69%. dari hasil berikut bahwa dalam 20 iterasi melalui dataset pelatihan, model dapat memahami pola dan fitur dengan memadai untuk memberikan prediksi yang akurat.

Tabel 14. Hasil Pengujian Jumlah Epoch

| Jumlah Epoch | Accuracy | Precision | Recall | F1-Score |
|--------------|----------|-----------|--------|----------|
| 10           | 73.92%   | 85.71%    | 77.67% | 81.49%   |
| 20           | 74.25%   | 80.29%    | 81.09% | 80.69%   |
| 30           | 74.25%   | 80%       | 80.95% | 81.15%   |

Dari hasil uji parameter menunjukkan bahwa model optimal ditemukan dengan pembagian data *training* sebesar 80% dan data *testing* sebesar 20%, menggunakan 100 neuron dan 20 *epoch*. Performa model akan dievaluasi berdasarkan akurasi pada data *training* dan *validation*. Gambar 4 menampilkan hasil pelatihan untuk pembangunan model.



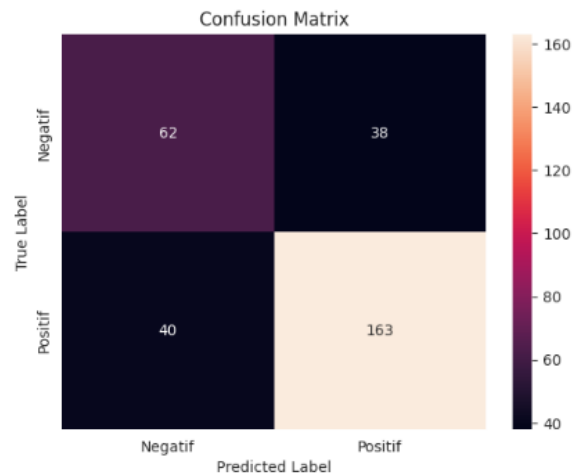
Gambar 4. (a) Accuracy Training, (b) Loss Training

Berdasarkan gambar 4, *accuracy training* mencapai sebesar 99.88% dengan *loss training* sebesar 72.72%. Perbandingan ini menunjukkan kinerja yang baik, yang ditandai oleh kesamaan akurasi antara data *training* dan *validation*. Hal ini mengindikasikan bahwa model tidak cenderung mengalami *overfitting*.

#### 4.8. Accuracy Score

Pada tahap akhir penelitian, metrik evaluasi adalah akurasi yang mengukur kemampuan model

klasifikasi dalam memprediksi secara benar pada keseluruhan dataset. Model diuji dengan 1513 data *tweet*, terdiri dari 1210 *tweet* positif dan 303 *tweet* negatif. Pengujian dilakukan dengan parameter terbaik, yaitu pembagian data *training* 80% dan data *testing* 20%, jumlah neuron 100, *learning rate* 0.001, 20 *epochs*, dan *batch size* 32.



Gambar 5. Confusion Matrix

Evaluasi dari *confusion matrix* menunjukkan akurasi sebesar 74.25%, *precision* 80.29%, *recall* 81.09%, dan *f1-score* 80.69%. Analisis *confusion matrix* ini menunjukkan bahwa model mampu memprediksi dengan baik untuk kedua kelas sentimen, dengan nilai yang memuaskan untuk akurasi, *precision*, *recall*, dan *f1-score*. Keseluruhan, model LSTM ini berhasil memberikan kinerja yang baik dalam memprediksi sentimen pada data *tweet* yang diuji.

Visualisasi untuk sentimen kelas positif dan negatif dapat direpresentasikan kedalam *wordcloud* pada gambar 6 dan 7.



Gambar 6. Wordcloud Positif

Pada gambar 6 menghasilkan beberapa kata-kata untuk *tweet* positif mengenai perkembangan ai, diantaranya yaitu: “kembang”, “teknologi”, “cerdas”, “manfaat” dan lain-lain.



Gambar 7. *Wordcloud* Negatif

Pada gambar 7 menghasilkan beberapa kata-kata untuk *tweet* negatif mengenai perkembangan ai, diantaranya yaitu: “salah”, “hidup”, “pakai”, “khawatir” dan lain-lain.

## 5. KESIMPULAN DAN SARAN

Bahwa penerapan metode *Knowledge Discovery in Database* (KDD) dengan algoritma *Long Short-Term Memory* (LSTM) berhasil menganalisis sentimen pengguna Twitter terhadap perkembangan *Artificial Intelligence* (AI). Analisis terhadap 1513 *tweet* menunjukkan bahwa sebanyak 999 *tweet* bersentimen positif dan 514 *tweet* bersentimen negatif terhadap kemajuan AI. Pengujian menunjukkan bahwa model LSTM mencapai akurasi tertinggi sebesar 74.25% pada *epoch* 20 dengan 100 *neuron* pada *layer* LSTM. Eksperimen dengan pembagian data 80:20 antara data latih dan data uji menghasilkan kinerja terbaik dengan 1210 data latih dan 303 data uji. Kesimpulan akhir menekankan bahwa penerapan algoritma LSTM pada data *tweet* perkembangan AI memberikan kinerja yang baik, dengan akurasi 74.25%, *precision* 80.29%, *recall* 81.09%, dan *f1-score* 80.69%.

Saran untuk penelitian selanjutnya melibatkan peningkatan validitas hasil analisis sentimen dengan mengumpulkan jumlah data *tweet* yang lebih besar. Fokus yang lebih mendalam pada tahap *preprocessing*, terutama dalam hal *labeling*, dianggap penting untuk meningkatkan akurasi hasil. Metode manual untuk *labeling* dapat dieksplorasi sebagai opsi alternatif untuk meningkatkan kualitas label, terutama saat menghadapi kasus khusus yang sulit diinterpretasikan oleh algoritma. Selain itu, disarankan untuk melakukan pengujian tambahan dengan variasi metode atau parameter lainnya agar dapat memperoleh pemahaman yang lebih komprehensif mengenai kinerja model dalam menganalisis sentimen pengguna Twitter terhadap kemajuan AI.

## DAFTAR PUSTAKA

- [1] S. Masrichah, "Ancaman dan Peluang *Artificial Intelligence* (AI)," *Khatulistiwa: Jurnal Pendidikan dan Sosial Humaniora*, vol. 3, no. 3, pp. 83–101, 2023, [Online]. Available:

- <https://journal.amikveteran.ac.id/index.php/Khatulistiwa/article/view/1860>
- [2] Y. S. Pongtambing *et al.*, “Peluang dan Tantangan Kecerdasan Buatan Bagi Generasi Muda,” *Bakti Sekawan: Jurnal Pengabdian Masyarakat*, vol. 3, no. 1, pp. 23–28, 2023, doi: 10.35746/bakwan.v3i1.362.
- [3] E. Indrayuni and A. Nurhadi, “OPTIMASI NAIVE BAYES BERBASIS PSO UNTUK ANALISA SENTIMEN PERKEMBANGAN ARTIFICIAL INTELLIGENCE DI TWITTER,” *INTI Nusa Mandiri*, vol. 18, no. 1, pp. 65–70, 2023, [Online]. Available: <https://ejournal.nusamandiri.ac.id/index.php/inti/article/view/4282>
- [4] S. F. Handayani, R. W. Pratiwi, and M. Putriyani, “Analisis Sentimen Pada *Data Ulasan Twitter Dengan Menggunakan Long Short Term Memory*,” *POLITEKNIK HARAPAN BERSAMA TEGAL*, pp. 1–29, 2021, [Online]. Available: <http://eprints.poltektegal.ac.id/979/>
- [5] M. H. Al-Areef and K. Saputra, “Analisis Sentimen Pengguna Twitter Mengenai Calon Presiden Indonesia Tahun 2024 Menggunakan Algoritma LSTM,” *Jurnal SAINTIKOM (Jurnal Sains ...)*, vol. 22, no. 2, pp. 270–279, 2023, [Online]. Available: <http://ojs.trigunadharma.ac.id/index.php/jis/article/view/8680>
- [6] M. Diki Hendriyanto, A. A. Ridha, and U. Enri, “Analisis Sentimen Ulasan Aplikasi Mola Pada Google Play Store Menggunakan Algoritma Support Vector Machine Sentiment Analysis of Mola Application Reviews on Google Play Store Using Support Vector Machine Algorithm,” *Journal of Information Technology and Computer Science (INTECOMS)*, vol. 5, no. 1, pp. 1–7, 2022.
- [7] Y. yuli Astari, Afyati, and S. W. Rozaqi, “Analisis Sentimen Multi-Class pada Sosial Media menggunakan metode *Long Short-Term Memory (LSTM)*,” *Jurnal Linguistik Komputasional*, vol. 4, no. 1, pp. 8–12, 2021, [Online]. Available: <http://inacl.id/journal/index.php/jlk/article/view/43>
- [8] F. Yuspriyadi, “KLASIFIKASI SENTIMEN TWITTER MENGGUNAKAN LSTM,” *METHODIKA: Jurnal Teknik Informatika dan ...*, vol. 9, no. 1, pp. 4–8, 2023, [Online]. Available: <https://ejurnal.methodist.ac.id/index.php/methodika/article/view/1720>
- [9] J. Nurvania, J. Jondri, and ..., “Analisis Sentimen Pada Ulasan di TripAdvisor Menggunakan Metode *Long Short-Term Memory (LSTM)*,” *eProceedings ...*, vol. 8, no. 4, pp. 1–12, 2021, [Online]. Available: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/15240>

- [10] A. A. Aljabar and A. A. A. Karim, “ANALISIS SENTIMEN MENGGUNAKAN ALGORITMA LSTM PADA MEDIA SOSIAL,” *Jurnal Publikasi Ilmu Komputer dan Multimedia*, vol. 1, no. 3, pp. 181–187, 2022, doi: 10.55606/jupikom.v1i3.517.
- [11] A. Yahyadi and F. Latifah, “Analisis Sentimen Twitter Terhadap Kebijakan PPKM di Tengah Pandemi COVID-19 Menggunakan Mode LSTM,” *Journal of Information System, Applied, Management, Accounting and Research*, vol. 6, no. 2, pp. 464–471, 2022, doi: 10.52362/jisamar.v6i2.791.