

ANALISIS SENTIMEN PENGGUNA TWITTER TERHADAP BOIKOT PRODUK ISRAEL PADA MEREK DAGANG UNILEVER INDONESIA MENGGUNAKAN ALGORITMA LONG SHORT TERM MEMORY

Andika Muhammad Aditya, Syariful Alam, Mutiara Andayani Komara

Teknik Informatika, Sekolah Tinggi Teknologi (STT) Wastukencana Purwakarta
Jalan Cikopak No.53, Mulyamekar, Kec.Babakancikao, Kab. Purwakarta, Jawa Barat, Indonesia
andikaaditya1359@gmail.com

ABSTRAK

Unilever Indonesia menghadapi tantangan serius terkait tanggapan publik terhadap isu sosial dan politik yang mempengaruhi citra merek dan kinerja keuangan mereka, terutama setelah dugaan dukungan akan tindakan Israel terhadap Palestina yang memicu gerakan boikot di Indonesia. Penelitian ini melibatkan *algoritma Long Short Term Memory* (LSTM) untuk melakukan klasifikasi sentimen. Data teks diubah menjadi vektor numerik melalui teknik *Word2Vec* dengan dimensi 300. Ada dua skenario pengujian yang dilakukan untuk mencapai akurasi terbaik yaitu skenario pertama dengan 80% data latih dan 20% data uji, sementara skenario kedua menggunakan 90% data untuk pelatihan dan 10% untuk pengujian. Beberapa parameter diuji, termasuk jumlah *neuron* pada *layer* LSTM, *layer dropout*, jumlah *epoch*, ukuran *batch*, dan nilai *learning rate* untuk meminimalisir *overfitting* dan meningkatkan kinerja model. Hasil penelitian menunjukkan skenario kedua lebih baik, dengan akurasi pelatihan 95,17% dan akurasi pengujian 77,11%, dibandingkan skenario pertama dengan akurasi pelatihan 97,58% dan akurasi pengujian 74,40%, yang mengindikasikan *overfitting*. Skenario kedua juga menunjukkan bahwa model cenderung lebih akurat dalam klasifikasi sentimen negatif, dibandingkan dengan sentimen positif, dan sentimen netral.

Kata kunci : Boikot, dropout, learning rate, LSTM, Word2Vec

1. PENDAHULUAN

Dengan perkembangan teknologi yang pesat, jumlah pengguna internet meningkat pesat dan volume data online juga bertambah. Data teks yang melimpah mendorong penelitian di bidang *text mining* dan NLP (*Natural Language Processing*). Salah satu Teknik yang populer adalah analisis sentimen, dimana teknik ini meningkat karena banyak pihak membutuhkan opini public dan memiliki dataset dengan kalimat panjang sehingga memerlukan pendekatan khusus [1].

Twitter adalah salah satu platform media sosial yang paling banyak digunakan untuk menyampaikan opini dan reaksi terkait berbagai isu. Didirikan oleh Jack Dorsey, yang berfungsi sebagai media untuk mengirimkan pesan singkat yang disebut *tweet*. Melalui Twitter, masyarakat dapat secara luas mengungkapkan sentimen dan pandangan mereka dengan cepat dan efektif [2].

Konflik Israel-Palestina yang berlangsung sejak 1917 hingga kini, melibatkan penguasaan wilayah Palestina oleh Israel dan pengungsian warga Palestina [2]. Berbagai negara mengambil sikap pro-Israel atau pro-Palestina, termasuk gerakan BDS yang menekan Israel dari segi ekonomi, budaya, dan politik. Unilever yang berdiri sejak 1930 menjadi sorotan karena dugaan dukungan terhadap Israel sehingga menyoroti pentingnya tanggung jawab sosial perusahaan terhadap citra merek dan kinerja finansial. Oleh karena itu, isu ini menjadi perhatian utama dalam hubungan internasional yang kompleks dan terus berkembang.

Situasi konflik yang berkelanjutan antara Israel dan Palestina telah memicu tanggapan aktif dari masyarakat untuk melakukan aksi boikot terhadap produk-produk yang terkait dengan Israel, yang juga menyebar luas di media sosial, terutama Twitter [2].

Keterlibatan analisis sentimen digunakan untuk mengklasifikasi sentimen pengguna Twitter terkait dugaan boikot produk Unilever Indonesia menggunakan algoritma *Long Short Term Memory* (LSTM) menjadi kategori positif, negatif, dan netral. Informasi dari Twitter bisa memberikan wawasan tentang pendapat masyarakat mengenai suatu isu. Hal ini bisa dimanfaatkan untuk penelitian dalam menganalisis sentimen.

Pada penelitian Yahyadi & Latifah mengenai analisis sentimen mobil listrik di Indonesia menghasilkan akurasi pengujian 70% dengan parameter terbaik berupa *batch size* 128, *dropout rate* 0.2, *embed dim* 32, *epoch* 10, *hidden unit* 16, dan *learning rate* 0.0001 [3].

Penelitian ini berfokus untuk mengkategorikan opini pengguna Twitter ke dalam tiga jenis kelompok yaitu meliputi kelompok positif, negatif, dan netral, dengan memanfaatkan algoritma *Long Short Term Memory* (LSTM). Algoritma ini dipilih karena mampu mengatasi masalah *vanishing gradient* pada data panjang, memiliki akurasi tinggi dalam pengolahan teks, serta lebih unggul dalam menangani *long-term dependency*, sehingga menawarkan solusi efektif untuk dataset besar dan kompleks yang memerlukan pemahaman mendalam mengenai konteks panjang

serta dapat memberikan hasil lebih baik dibandingkan algoritma lain [1].

Di penelitian ini, pengolahan data dilakukan dengan melibatkan *tools Google Colab* dengan bahasa pemrograman *Python* dan beberapa *library Python* terkait.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis sentimen yaitu sebuah teknologi baru yang dikembangkan untuk penelitian, membantu menilai topik dan opini. Teknologi ini sering digunakan untuk mengevaluasi dan menganalisis kepuasan pelanggan terhadap produk atau kebijakan, dengan menggunakan data latih dan data uji dalam model *machine learning* [3].

2.2. Twitter

Twitter adalah sebuah wadah media sosial untuk teman, keluarga, dan rekan kerja bertukar pesan cepat termasuk foto, video, tautan, dan teks, yang memungkinkan pengguna berbagi pendapat dan pengalaman mereka, sehingga data yang melimpah ini bisa dianalisis untuk mendapatkan wawasan baru yang berguna [3].

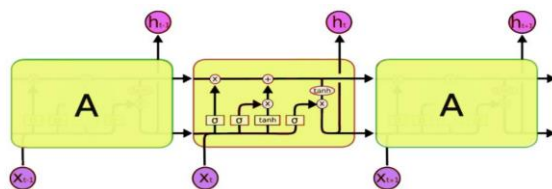
2.3. Boikot Produk

Boikot produk adalah tindakan konsumen yang tidak puas secara ideologis terhadap suatu organisasi atau negara, karena produk yang dihasilkan tidak sesuai dengan nilai atau pandangan mereka [4].

2.4. Long Short Term Memory (LSTM)

Pada tahun 1997, Hochreiter dan Schmidhuber memperkenalkan LSTM untuk mengatasi kelemahan RNN dalam mengelola memori jangka panjang dengan *memory cells* dan *gate units* untuk mempertahankan informasi relevan dari input awal hingga akhir sehingga memungkinkan pembelajaran yang lebih efektif dari data berurutan [5].

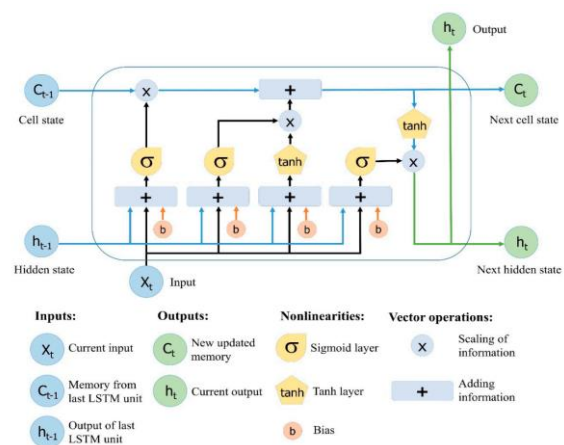
Memory Cells dari algoritma ini dapat dilihat pada Gambar 1.



Gambar 1. *Memory cells* pada LSTM [5]

Gambar 1 menggambarkan bagaimana proses kerja sel-sel memori pada setiap *neuron* LSTM beroperasi, dengan terdapat empat *gates units* yang mengontrol aktivasi masukan, yaitu *forget gates*, *input gates*, *cell states*, dan *output gates* [5].

Arsitektur jaringan LSTM dapat dilihat lebih jelasnya pada Gambar 2.



Gambar 2. Arsitektur jaringan LSTM [6]

Jaringan LSTM biasanya terdiri dari blok memori yang disebut *cell*, di mana terdapat dua keadaan yaitu *cell state* dan *hidden state*. yang kemudian diteruskan ke sel berikutnya. *Cell state* berfungsi sebagai jalur utama aliran data yang memungkinkan informasi mengalir hampir tanpa perubahan, sementara informasi dapat ditambah atau dikurangi melalui gerbang *sigmoid* yang mengontrol proses memori dan membantu mengatasi masalah ketergantungan jangka panjang. pada tahap awal pembuatan jaringan LSTM, fungsi *sigmoid* digunakan untuk menentukan informasi yang tidak diperlukan dan harus dihapus dari status *cell*, serta bagian mana dari *output* lama yang harus dihapus, melalui mekanisme yang disebut *forget gate*, di mana *forget gate* ini merupakan vektor dengan nilai antara 0 hingga 1 yang menggambarkan seberapa besar setiap elemen dalam status sel sebelumnya harus dihapus [6].

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

Pada persamaan (1) fungsi *sigmoid*, yang dinyatakan dengan σ , berperan dalam menentukan apakah informasi baru yang diperoleh dari *input* akan diterima atau diabaikan, berdasarkan nilai dari matriks bobot dan bias yang ada pada *forget gate* [6].

Proses ini melibatkan dua tahap utama yang pertama yaitu lapisan *sigmoid* memutuskan informasi mana yang perlu diperbarui atau diabaikan, dengan hasil berupa nilai antara 0 dan 1 dan yang kedua yaitu fungsi *tanh* memberikan bobot pada nilai yang lolos dari tahap sebelumnya, dengan rentang antara -1 hingga 1, untuk menentukan tingkat kepentingannya. Kemudian, nilai-nilai ini dikalikan untuk memperbarui *cell state* baru. Selanjutnya, memori baru ini ditambahkan pada memori lama untuk menghasilkan *cell state* yang baru. Ketiga proses ini dapat dilihat pada persamaan (2), persamaan (3), dan persamaan (4) [6].

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

$$N_t = \tanh(W_n[h_{t-1}, x_t] + b_n) \quad (3)$$

$$C_t = (f_t \times C_{t-1} + i_t \times N_t) \quad (4)$$

Pada waktu t , nilai *cell state* saat ini dipengaruhi oleh nilai *cell state* sebelumnya melalui bobot dan bias. Nilai *output* akhir diperoleh dari keadaan *output cell state* yang telah difilter. Proses dimulai dengan menentukan bagian mana dari *cell state* yang akan digunakan untuk *output*, diikuti dengan mengalikan hasil ini dengan nilai baru yang dihasilkan oleh lapisan *tanh*. Selama proses ini, bobot dan bias dari *output gate* juga berperan dalam menentukan hasil akhir [6]. Proses ini dapat dilihat pada persamaan (5), dan persamaan (6)

$$O_t = \sigma(W_o[h_{t-1}, x_t] + b_o) \quad (5)$$

$$h_t = O_t \tanh(C_t) \quad (6)$$

Dalam LSTM, nilai *hidden state* (h_t) dihitung dengan mengalikan *output gates* (O_t) dengan fungsi *tangen hiperbolik* dari sel memori untuk memetakan nilai-nilai *input* ke dalam rentang -1 hingga 1 yang membantu dalam normalisasi data dan dapat mempercepat hasil yang optimal saat pelatihan, sehingga menghasilkan informasi terbaru yang telah diatur oleh *output gates* dan dimodifikasi oleh sel memori untuk digunakan dalam komputasi langkah waktu berikutnya atau tugas khusus lainnya [6].

2.5. Preprocessing Text

Pada tahap ini, bertujuan untuk mengubah data yang tidak terstruktur dan terdapat *noise* harus dibersihkan agar kualitasnya meningkat dan menghasilkan pengetahuan yang lebih baik [7].

Dalam penelitian ini, proses awal yang dilakukan adalah sebagai berikut.

a. Data Cleaning

Data cleaning adalah tahapan untuk menghapus elemen-elemen seperti tanda baca, angka, URL, *username*, *mention*, *hashtag*, dan tanda *retweet*, sehingga data tersebut dapat disaring dari karakter-karakter yang tidak relevan atau tidak diperlukan untuk analisis lebih lanjut [7].

b. Data Selection

Tahapan ini bertujuan untuk mengenali dan menentukan fitur atau kolom yang memberikan dampak positif terhadap hasil klasifikasi, serta menghapus fitur atau kolom yang tidak berkaitan dengan tujuan penelitian [8].

c. Case Folding

Case Folding adalah tahapan dalam mengubah semua huruf pada semua data pada dokumen menjadi huruf kecil [7].

d. Normalisasi Kata

Normalisasi Kata adalah tahapan mengubah kata-kata singkatan menjadi bentuk lengkapnya yang sesuai dengan Kamus Besar Bahasa Indonesia (KBBI). Contohnya, "utk" diubah menjadi "untuk" dan "yg" menjadi "yang". Dengan begitu, informasi dapat diproses lebih mudah dan jelas [7].

e. Translate Tweets Data

Pada tahap ini, data *tweets* berbahasa Indonesia diterjemahkan menjadi *tweets* berbahasa Inggris,

karena *VADER Sentiment*, alat analisis sentimen yang akan digunakan, didasarkan pada korpus bahasa Inggris [9].

f. Labelling Dengan VADER

VADER (Valence Aware Dictionary and sEntiment Reasoner) adalah sebuah alat untuk menganalisis sentimen yang dirancang khusus untuk media sosial. Alat ini menggunakan kamus leksikal yang telah dikurasi dan diberi skor sentimen secara manual, sehingga bisa memberikan nilai sentimen positif, negatif, atau netral untuk setiap tweet [10].

g. Tokenization

Tokenization adalah proses memecah sebuah kalimat atau string menjadi bagian-bagian kecil seperti kata-kata untuk mendapatkan informasi atau nilai dari masing-masing kata tersebut [7].

h. Stopword Removal

Pada langkah ini, penghapusan kata-kata yang biasa dipakai tetapi tidak jelas seperti "untuk", "dari", "yang", "di", "dari" dan sebagainya [7].

i. Stemming

Pada tahap ini, menghapus semua imbuhan seperti awalan, akhiran, dan konfiks yang terdapat dalam data tweet [7].

2.6. Feature Extraction

Fitur ekstraksi yang dilibatkan dalam penelitian ini adalah *Word2Vec* yang digunakan untuk mengubah kata menjadi vektor, di mana kata-kata dengan makna mirip akan memiliki representasi yang serupa. Model *skip-gram* dipilih karena efektif dalam memprediksi kata-kata di sekitar suatu kata dalam teks yang besar dan bervariasi [11].

2.7. Confusion Matrix

Confusion matrix adalah alat untuk menilai kinerja metode klasifikasi dengan membandingkan hasil klasifikasi sistem dengan klasifikasi sebenarnya dimana data uji yang dimasukkan ke dalam matriks ini menghasilkan nilai akurasi [12]. Tabel 1. menunjukkan struktur *confusion matrix* yang dapat dilihat berikut.

Tabel 1. Struktur *confusion matrix* [12]

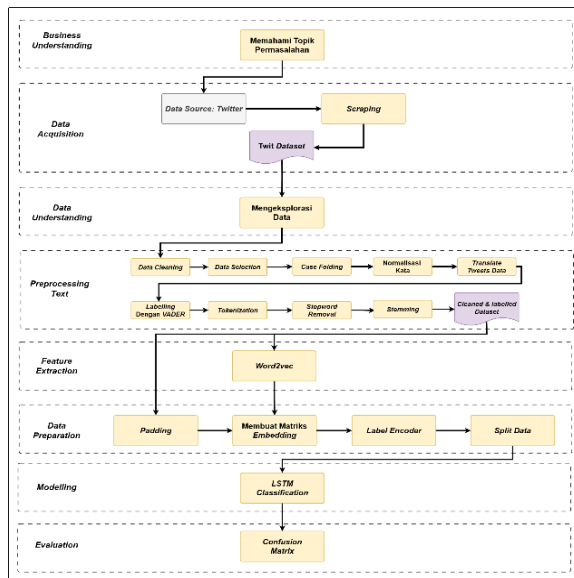
Confusion Matrix		Actual Value		
		Positive	Negative	Neutral
Predicted Value	Positive	TP	FP	FP
	Negative	FNe	TNe	FNe
	Neutral	FNt	FNt	TNt

Keterangan:

TP : True Positive
 TNt : True Neutral
 TNe : True Negative
 FP : False Positive
 FNt : False Neutral
 FNe : False Negative

3. METODE PENELITIAN

Penelitian ini dimulai dengan tahap *business understanding*, diikuti oleh *data acquisition*, *data understanding*, *text preprocessing*, dan *feature extraction* menggunakan *Word2Vec*, kemudian dilanjutkan dengan *data preparation*, pemodelan melibatkan algoritma *Long Short-Term Memory* (LSTM), dan diakhiri dengan tahap *evaluation*, sebagaimana ditunjukkan pada Gambar 4.



Gambar 3. Kerangka pemikiran

3.1. Business understanding

Pada tahap ini, melakukan analisis opini publik di Twitter tentang boikot produk Israel dari Unilever Indonesia untuk memahami masalah yang muncul dari perbedaan antara pandangan publik dan kebijakan perusahaan, serta menilai kejelasan fenomena ini guna menentukan fokus dan batasan penelitian.

3.2. Data acquisition

Pada langkah ini, melakukan pengambilan data dari Twitter untuk meneliti pandangan masyarakat terkait boikot produk Israel oleh Unilever Indonesia. Proses pengambilan data menggunakan teknik scraping dengan Google Colab dan pustaka *tweet-harvest*, dengan menggunakan kata kunci "Boikot Unilever lang:id". Data disimpan dalam format CSV untuk analisis yang akurat dan mempengaruhi hasil penelitian.

3.3. Data understanding

Pada langkah ini, langkah awal yang perlu dilakukan adalah menganalisis data untuk menemukan pola, konsistensi, keakuratan, dan kelengkapannya, serta mengidentifikasi anomali atau kesalahan, guna memastikan kualitas dan keandalan data sebelum digunakan dalam proses selanjutnya

3.4. Text preprocessing

Proses awal yang dilakukan yaitu pengolahan data teks yang melibatkan 9 tahap *preprocessing* yang meliputi *data cleaning*, *data selection*, *case folding*, normalisasi kata, *translate tweets data*, *labelling dengan VADER*, *tokenization*, *stopword removal*, dan *stemming*, sehingga teks menjadi lebih terstruktur dan siap untuk dilakukan pengolahan data lebih lanjut.

3.5. Feature extraction

Pada langkah ini, melakukan proses transformasi kata menjadi representasi vektor yang dilakukan dengan metode *Word2Vec* dengan model *skip-gram*.

3.6. Data preparation

Pada alur proses ini, langkah pertama yang dilakukan adalah melakukan proses *padding* untuk menyamakan panjang data input, kemudian melakukan *label encoding* untuk mengubah nilai teks menjadi nilai numerik pada kolom sentimen, serta membagi data menjadi *training* dan *testing* untuk memastikan konsistensi dan akurasi analisis sentimen yang akurat.

3.7. Modelling

Dalam tahap pemodelan, melibatkan algoritma *Long Short Term Memory* (LSTM) yang dimana pada algoritma ini urutan kata diproses oleh lapisan LSTM dengan pengaturan seperti jumlah *epoch*, ukuran *batch*, jumlah *neuron* di *hidden layer*, *dropout layer*, serta *learning rate*, menggunakan fungsi aktivasi *softmax* dan optimasi *Nadam*, dan model dilatih berulang hingga mencapai hasil terbaik.

3.8. Evaluation

Pada tahap evaluasi model, digunakan *confusion matrix* sebagai evaluasi model klasifikasi yang bertujuan untuk membandingkan data aktual dengan hasil prediksi, sehingga dapat dihitung nilai akurasi, presisi, *recall*, dan *F1-Score* untuk menilai kinerja model secara menyeluruh.

4. HASIL DAN PEMBAHASAN

4.1. Business understanding

Penelitian ini akan berfokus pada analisis mendalam terhadap tanggapan masyarakat di media sosial khususnya Twitter, terkait dugaan boikot terhadap produk Unilever Indonesia. Tujuan utamanya adalah untuk mengidentifikasi dan mengevaluasi opini serta sentimen publik yang dapat dikategorikan menjadi kelas Positif yang berarti tweet yang mengandung seruan untuk melakukan boikot pada produk Unilever, kelas Negatif yang berarti tweet yang mengandung seruan untuk tidak melakukan boikot produk Unilever, dan Netral yang berarti tweet yang tidak mengandung seruan melakukan ataupun tidak melakukan boikot produk Unilever. Selain hal itu, penelitian ini juga menggunakan teknik analisis kuantitatif yang melibatkan algoritma *deep learning* LSTM untuk mengembangkan model klasifikasi.

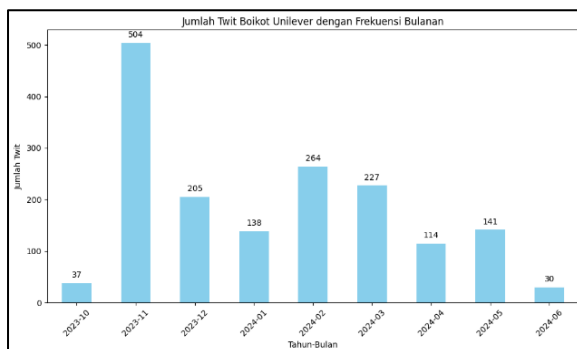
4.2. Data acquisition

Pada langkah ini, dilakukan proses pengumpulan data dengan teknik *scraping* yang dilakukan dua kali menggunakan *Google Collab* dan pustaka *tweet-harvest*, dengan kata kunci "Boikot Unilever lang:id", dan dipilihlah hasil *scraping* yang kedua yaitu pada tanggal 4 Juni 2024 diperoleh 1660 data dikarenakan jumlah data yang didapat lebih besar.

Gambar 4. Gambar dataset pada pengambilan data yang kedua

4.3. Data understanding

Dari analisis dataset Twitter selama rentang waktu 217 hari, ditemukan lonjakan tweet terkait boikot Unilever pada November 2023 akibat dugaan dukungannya terhadap agresi Israel. Rata-rata tweet pada 25 Januari 2024 menunjukkan interaksi rendah, dengan dominasi kata "Unilever" dan "boikot", yang menandakan perhatian publik tinggi.



Gambar 5. Gambar jumlah tweet boikot unilever dengan frekuensi bulanan

4.4. Text preprocessing

Tabel 2. menunjukkan data teks yang telah melalui proses pembersihan dan pemrosesan pada tahap *text preprocessing* sebelumnya.

Tabel 2. Hasil dari tahap *text preprocessing*

Label	Twit	Twit_Clean
Negatif	@dikaamfa @tanyakanrl Boikot ini tujuannya untuk intervensi kak Perusahaan gede sekelas Unilever gak akan semudah itu buat bangkrut Coba pemikirannya diperluas lagi	boikot tuju intervensi kakak usaha kelas unilever mudah bangkrut coba pikir luas
Positif	Pasti seru paradenya sebab Indonesia kan punya begitu	seru parade indonesia

Label	Twit	Twit_Clean
	banyak ragam busana daerah ~ Senin Malam Jatinangor Sebulan Unilever Boikot Durian Esqa Tidur Kebangun Apel Pagi	ragam busana daerah senin malam jatinangor bulan unilever boikot durian tidur
Netral	Produk P&G hingga Unilever Diskon Besar-besaran Imbas Boikot Produk Israel? https://t.co/52KayFBVob	produk unilever diskon besar akibat boikot produk israel

Tabel 2. menunjukkan hasil dari data yang telah diproses melalui beberapa langkah, termasuk tahap-tahap seperti *data cleaning*, *data selection*, *case folding*, normalisasi kata, *Translate Tweets Data*, *labelling* dengan *VADER*, *tokenization*, *stopword removal*, dan *stemming*.

4.5. Feature extraction

Model *Word2Vec* dengan algoritma *skip-gram* menggunakan beberapa parameter penting: "MIN_COUNT" mengatur frekuensi minimum kemunculan kata dalam korpus, "WINDOW" menentukan jarak maksimum antara kata target dan kata-kata konteks (5 kata), "EPOCH" menunjukkan jumlah iterasi pelatihan (15 epoch), "vector_size" menentukan panjang vektor *embedding* (300 dimensi), dan "SG" (Bernilai 1) menunjukkan penggunaan algoritma *skip-gram*.

4.6. Data preparation

Langkah awal pelatihan model meliputi persiapan data dengan *padding*, *matriks embedding*, dan *label encoding*. Ini mencakup tokenisasi teks, pembuatan indeks kata, dan penyesuaian panjang urutan yang dapat dilihat pada Tabel 3.

Tabel 3. Hasil dari tahap *padding*, *matriks embedding*, dan *label encoding*

Label	Twit_Clean	Array_Padding	Lbel_Kode
Negatif	boikot tuju intervensi kakak usaha kelas unilever mudah bangkrut coba pikir luas	Array([1, 215, ..., 0, 0])	0
Positif	seru parade indonesia ragam busana daerah senin malam jatinangor bulan unilever boikot durian tidur	Array([57, 1153, ..., 0, 0])	2
Netral	produk unilever diskon besar akibat boikot produk israel	Array([3, 2, ..., 0, 0])	1

Setelah proses sebelumnya telah selesai dilakukan, maka pada tahap ini yaitu melakukan pembagian dataset menjadi 2 bagian yaitu dengan

proporsi 80%:20% sebagai skenario pertama dan 90:10% sebagai skenario kedua.

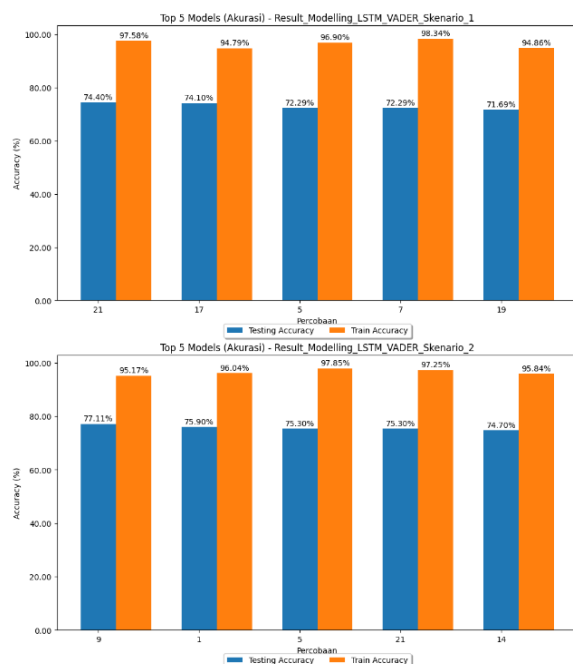
4.7. Modelling

Penelitian ini memanfaatkan algoritma LSTM dan Word2Vec model Skip-gram untuk analisis sentimen. Proses *modelling* dengan LSTM dilakukan dengan menetapkan beberapa parameter, seperti jumlah *neuron* pada lapisan LSTM, *dropout rate*, *learning rate*, jumlah *epoch*, ukuran *batch*, fungsi aktivasi *Softmax*, serta *optimizer Nadam*. Dan juga beberapa parameter ini menghasilkan 32 kombinasi percobaan model yang akan diseleksi mana yang optimal dan tertinggi. Konfigurasi parameter model LSTM dapat dilihat pada Tabel 4.

Tabel 4. Konfigurasi parameter model LSTM

Parameter	Value
jumlah <i>neuron</i> pada lapisan LSTM	{75, 100}
<i>dropout rate</i>	{0.3, 0.6}
<i>learning rate</i>	{0.001, 0.0001}
jumlah <i>epoch</i>	{10, 30}
ukuran <i>batch</i>	{32, 64}
fungsi aktivasi	Softmax
<i>optimizer</i>	Nadam

Dari Tabel 4, diperoleh hasil untuk dua skenario yang terlihat pada Gambar 6 serta Gambar 7.



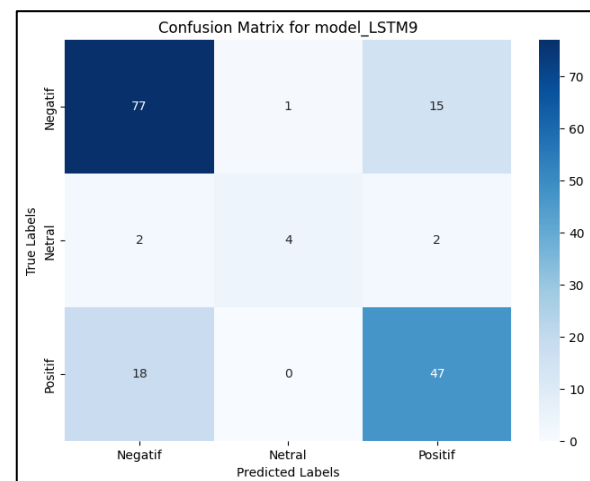
Gambar 6. Hasil 32 kombinasi percobaan model pada skenario pertama dan skenario kedua

Dari Gambar 6 adalah hasil 32 kali percobaan dalam melatih model LSTM dengan nilai parameter yang berbeda pada setiap percobaannya. Pada skenario pertama, akurasi tertinggi pada pengujian dicapai pada percobaan model ke-21, dengan akurasi pelatihan mencapai 97.58% dan pengujian 74.40%. Model ini

menggunakan 100 *neuron* LSTM, *dropout* 0.3, 30 *epoch*, *batch size* 32, dan *learning rate* 0.001. Skenario kedua, di mana akurasi pengujian tertinggi dicapai pada percobaan model ke-9 dengan akurasi pelatihan 95.17% dan pengujian 77.11%, memanfaatkan 75 *neuron* LSTM, *dropout* 0.6, 10 *epoch*, *batch size* 32, dan *learning rate* 0.001. Perbandingan ini menunjukkan bahwa dalam skenario pertama, model membutuhkan kapasitas dan waktu pelatihan yang lebih besar dengan lebih banyak *neuron* LSTM dan lebih banyak *epoch* untuk mencapai performa optimal. Sebaliknya, dalam skenario kedua, model menggunakan pengaturan yang lebih efisien dengan lebih sedikit *neuron* LSTM, *dropout* lebih tinggi, dan jumlah *epoch* yang lebih rendah.

4.8. Evaluation

Setelah menyelesaikan pembuatan model LSTM, langkah berikutnya adalah mengevaluasi performanya dengan memanfaatkan *confusion matrix* untuk menghitung akurasi, presisi, *recall*, dan *F1-Score*. Fokus evaluasi ini dilakukan pada percobaan ke-9 dalam skenario kedua, karena pada skenario tersebut diperoleh hasil yang paling optimal dan terbaik. Berikut adalah hasil dari *confusion matrix* dari model LSTM pada percobaan ke-9 untuk skenario kedua.



Gambar 7. Confusion matrix dari model LSTM pada percobaan ke-9 untuk skenario kedua

Dari Gambar 7 nilai-nilai yang terdapat pada *confusion matrix* akan dilakukan proses perhitungan untuk mengetahui performa dari model dengan mengambil salah satu dari tiga kelompok sentimen, dengan diambil kelompok sentimen positif, berikut adalah proses dalam pengukuran model yang meliputi akurasi, presisi, *recall*, dan *F1-Score* pada kelompok sentimen positif.

1) Akurasi

$$\begin{aligned}
 Accuracy &= \frac{TP + TNe + TNe}{Total} \\
 &= \frac{47 + 4 + 77}{166} \\
 Accuracy &= 0.7711 = 77.11\%.
 \end{aligned}$$

2) Presisi

$$\begin{aligned}
 \text{Precision} &= \frac{TP}{TP + FP} \\
 &= \frac{47}{47 + 2 + 15} \\
 &= \frac{47}{64} \\
 &= 0.7344 \\
 \text{Precision} &= 73.44\%.
 \end{aligned}$$

3) Recall

$$\begin{aligned}
 \text{Recall} &= \frac{TP}{TP + FNe + FNt} \\
 &= \frac{47}{47 + 18 + 0} \\
 &= \frac{47}{65} \\
 &= 0.7231 \\
 \text{Recall} &= 72.31\%.
 \end{aligned}$$

4) F1-Score

$$\begin{aligned}
 F1-S &= \left(2 \times \left(\frac{(\text{Precision} \times \text{Recall})}{(\text{Precision} + \text{Recall})} \right) \right) \\
 &= \left(2 \times \left(\frac{(73.44 \times 72.31)}{(73.44 + 72.31)} \right) \right) \\
 &= \left(2 \times \left(\frac{0.5310}{1.4575} \right) \right) \\
 &= (2 \times 0.3643) \\
 F1-S &= 0.7287 = 72.87\%.
 \end{aligned}$$

Selain dari perhitungan pada kelompok positif juga dapat dilihat nilai pada kelompok negatif dan netral yang ditunjukkan pada Tabel 5 dan Tabel 6.

Tabel 5. Akurasi dan presisi dari model skenario kedua percobaan ke-9

Skenario Kedua Percobaan Ke-9				
Akurasi		Presisi		
Pelatihan	Pengujian	Negatif	Netral	Positif
95.17%	77.11%	79.38%	80%	73.44%

Tabel 6. Recall dan f1-score dari model skenario kedua percobaan ke-9

Skenario Kedua Percobaan Ke-9					
Recall			F1-Score		
Negatif	Netral	Positif	Negatif	Netral	Positif
82.80%	50%	72.31%	81.05%	61.54%	72.87%

Pada Tabel 5 dan Tabel 6 secara keseluruhan hasil untuk model pada skenario kedua percobaan ke-9 memiliki performa model yang lebih baik pada kelompok negatif, pada kelompok positif mencapai hasil yang seimbang meskipun sedikit lebih rendah dibandingkan kelompok negatif, sementara performa pada kelompok netral kurang baik dan kurang optimal.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian, algoritma LSTM dengan *Word embedding Word2Vec* diuji dalam dua skenario berbeda, dengan skenario pertama dengan 80% data latih dan 20% data uji yang menghasilkan akurasi pelatihan 97,58% dan pengujian 74,40%, sedangkan skenario kedua menggunakan 90% data

latih dan 10% data uji, menunjukkan akurasi pelatihan 95,17% dan pengujian 77,11%. Secara keseluruhan, model yang terbaik terdapat pada skenario kedua percobaan ke-9 dan terlihat bahwa model cenderung lebih akurat dalam mengenali kelompok negatif, yang menunjukkan *tweet* yang mengandung seruan untuk tidak melakukan boikot produk Unilever, dibandingkan dengan kelompok positif, yang menunjukkan *tweet* yang mengandung seruan untuk melakukan boikot, dan kelompok netral, yang menunjukkan *tweet* yang tidak mengandung seruan untuk melakukan ataupun tidak melakukan boikot produk Unilever.

Untuk penelitian selanjutnya, disarankan untuk terus menyempurnakan metode dan teknik yang digunakan, seperti meningkatkan algoritma deteksi akun palsu, memperbaiki model penerjemahan, dan melakukan pelabelan dengan algoritma *clustering* serta validasi yang tepat. Selain itu, penting untuk mengoptimalkan *hyperparameter* model dengan mencoba berbagai kombinasi, termasuk menambahkan jenis *layer* tambahan dalam algoritma LSTM, dan menggunakan metode optimasi seperti *Gridsearch* untuk menemukan kombinasi parameter yang paling optimal, yang diharapkan dapat meningkatkan performa model secara signifikan.

DAFTAR PUSTAKA

- [1] W. Widayat, "Analisis Sentimen Movie Review menggunakan Word2Vec dan metode LSTM Deep Learning," *J. Media Inform. Budidarma*, vol. 5, no. 3, hal. 1018–1026, 2021, doi: <http://dx.doi.org/10.30865/mib.v5i3.3111>.
- [2] T. S. Ade, N. A. Lestari, dan P. A. Nina, "Analisis Sentimen Publik Pada Twitter Terhadap Boikot Produk Israel Menggunakan Metode Naïve Bayes," *J. Ilm. Mhs.*, vol. 2, no. 1, hal. 26–35, 2024, doi: <https://doi.org/10.59603/niantanasikka.v2i1.240>.
- [3] A. Yahyadi dan F. Latifah, "ANALISIS SENTIMEN TWITTER TERHADAP KEBIJAKAN PPKM DI TENGAH PANDEMI COVID-19 MENGGUNAKAN MODE LSTM," *J. Inf. Syst. Applied, Manag. Account. Res.*, vol. 6, no. 2, hal. 464–471, 2022, doi: <https://doi.org/10.52362/jisamar.v6i2.791>.
- [4] C.-C. Wang, S.-C. Chang, dan P.-Y. Chen, "The brand sustainability obstacle: Viewpoint incompatibility and consumer boycott," *Sustainability*, vol. 13, no. 9, hal. 1–23, 2021, doi: 10.3390/su13095174.
- [5] I. L. Rais dan J. Jondri, "Klasifikasi Data Kuesioner dengan Metode Recurrent Neural Network," *eProceedings Eng.*, vol. 7, no. 1, hal. 2817–2826, 2020, [Daring]. Tersedia pada: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/11458>
- [6] X.-H. Le, H. V. Ho, G. Lee, dan S. Jung, "Application of Long Short-Term Memory

- (LSTM) neural network for flood forecasting,” *Water*, vol. 11, no. 7, hal. 1–19, 2019, doi: <https://doi.org/10.3390/w11071387>.
- [7] M. Furqan, S. Sriani, dan S. M. Sari, “Analisis Sentimen Menggunakan K-Nearest Neighbor Terhadap New Normal Masa Covid-19 Di Indonesia,” *J. Teknol. Inf. Techno.Com*, vol. 21, no. 1, hal. 52–61, 2022, doi: [10.33633/tc.v21i1.5446](https://doi.org/10.33633/tc.v21i1.5446).
- [8] M. R. Givari, M. R. Sulaeman, dan Y. Umaidah, “Perbandingan Algoritma SVM, Random Forest Dan XGBoost Untuk Penentuan Persetujuan Pengajuan Kredit,” *J. Nuansa Inform.*, vol. 16, no. 1, hal. 141–149, 2022, doi: [10.25134/nuansa.v16i1.5406](https://doi.org/10.25134/nuansa.v16i1.5406).
- [9] F. Amaliah dan I. K. D. Nuryana, “Perbandingan Akurasi Metode Lexicon Based Dan Naive Bayes Classifier Pada Analisis Sentimen Pendapat Masyarakat Terhadap Aplikasi Investasi Pada Media Twitter,” *J. Informatics Comput. Sci.*, vol. 3, no. 03, hal. 384–393, 2022, doi: <https://doi.org/10.26740/jinacs.v3n03.p384-393>.
- [10] A. D. Pratama dan H. Hendry, “Analisa Sentimen Masyarakat Terhadap Penggunaan Chatgpt Menggunakan Metode Support Vector Machine (Svm),” *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.)*, vol. 9, no. 1, hal. 327–338, 2024, doi: [10.29100/jipi.v9i1.4285](https://doi.org/10.29100/jipi.v9i1.4285).
- [11] E. Suryati, S. Styawati, dan A. A. Aldino, “Analisis Sentimen Transportasi Online Menggunakan Ekstraksi Fitur Model Word2vec Text Embedding Dan Algoritma Support Vector Machine (SVM),” *J. Teknol. dan Sist. Inf.*, vol. 4, no. 1, hal. 96–106, 2023, doi: <https://doi.org/10.33365/jtsi.v4i1.2445>.
- [12] D. N. F. Fitriana dan Y. Sibaroni, “Sentiment Analysis on KAI Twitter Post Using Multiclass Support Vector Machine (SVM),” *RESTI J. Syst. Eng. Inf. Technol.*, vol. 4, no. 5, hal. 846–853, 2020, doi: <https://doi.org/10.29207/resti.v4i5.2231>.