

Analisis Sentimen dalam Aplikasi X terhadap Pengungsi Rohingya dengan LSTM

Andri Pratama¹, Miftahurrahma Rosyda^{2*}

^{1,2}Fakultas Teknologi Industri, Informatika, Universitas Ahmad Dahlan, Yogyakarta, Indonesia

E-mail: ¹andri2100018249@webmail.uad.ac.id, ^{2*}miftahurrahma.rosyda@tif.uad.ac.id

(* : corresponding author)

Abstrak

Pengungsi Rohingya di Indonesia menghadapi diskriminasi dan ujaran kebencian akibat disinformasi di media sosial. Polemik terhadap pengungsi Rohingya di Indonesia ini telah menciptakan beragam tanggapan di kalangan masyarakat. Pemahaman terhadap perspektif dan reaksi publik Indonesia terkait isu ini menjadi sangat krusial, terutama dalam menganalisis sentimen yang berkembang. Penelitian ini dilakukan untuk mengevaluasi pandangan dan sentimen masyarakat tentang pengungsi Rohingya di Indonesia dalam platform media sosial X. Data dikumpulkan melalui *web scraping* menggunakan *twikit* selama bulan Desember 2023 menghasilkan 17.613 *tweet*. Pelabelan dilakukan dengan model IndoBERTweet yang telah dilatih dengan *dataset* publik IndoNLU smsa_doc-sentiment-prosa. Model Long Short-Term Memory (LSTM) dijalankan dua metode penyeimbangan kelas, *Random Oversampling* dan *SMOTE*. Hasil menunjukkan bahwa dengan *Random Oversampling*, model mencapai *precision* 0.9139, skor *f1* 0.9379, *recall* 0.9632, dan akurasi 0.9069. Sementara itu, penggunaan *SMOTE* menghasilkan *precision* 0.9092, skor *f1* 0.9265, *recall* 0.9445, dan akurasi 0.8906. Temuan penelitian memperlihatkan sentimen yang dominan negatif, dari 17.613 data, 73.04% bersentimen negatif dan 26.96% bersentimen positif.

Kata kunci: analisis sentimen, long-short term memory, deep learning, pemrosesan bahasa alami, rohingya

Abstract

Rohingya refugees in Indonesia face discrimination and hate speech due to disinformation on social media. The polemic against Rohingya refugees in Indonesia has created various responses among the public. Understanding the perspectives and reactions of the Indonesian public regarding this issue is crucial, especially in analyzing the growing sentiment. This exploration aimed to investigate the views and sentiments of Indonesians towards the Rohingya refugees in Indonesia through X social media platforms. Data was collected through web scraping using twikit during December 2023, resulting in 17,613 tweets. Labeling was done using the trained IndoBERTweet model with the IndoNLU smsa_doc-sentiment-prosa dataset. The Long Short-Term Memory (LSTM) model applied with two class balancing methods, Random Oversampling and SMOTE. The results show that with Random Oversampling, the model achieves precision 0.9139, f1 score 0.9379, recall 0.9632, also accuracy 0.9069. Meanwhile, SMOTE resulted in precision 0.9092, f1 score 0.9265, recall 0.9445, also accuracy 0.8906. The probing result show that the sentiment is predominantly negative from 17,613 data, 73.04% have negative sentiment and 26.96% have positive sentiment.

Keywords: sentiment analysis, long-short term memory, deep learning, natural language processing, rohingya

1. PENDAHULUAN

Rohingya merupakan golongan etnis minoritas dari Rakhine, Myanmar. Mereka menghadapi diskriminasi sistematis dari pemerintah Myanmar. Kondisi ini mengakibatkan eksodus besar-besaran ke negara-negara tetangga, termasuk Indonesia [1]. Meskipun sebelumnya masyarakat Indonesia menerima pengungsi Rohingya dengan baik, belakangan ini terjadi perubahan sentimen yang signifikan, terutama dipicu oleh disinformasi dan ujaran kebencian yang meningkat [2] termasuk di media sosial seperti Instagram, X, TikTok dan lainnya. Hal ini membuat analisis sentimen mengenai Rohingya menjadi penting untuk mengetahui opini publik.

Platform X menjadi salah satu media penyebaran informasi dalam bentuk teks mengenai isu Rohingya di Indonesia. Analisis sentimen melalui X melibatkan beragam komentar, penggunaan *hashtag*, dan bahasa sehari-hari, sehingga memungkinkan pemahaman mendalam dan langsung terhadap opini publik [3]. Analisis sentimen sendiri merupakan teknik *text mining* yang dapat mengidentifikasi opini positif atau negatif dari pengguna internet [4]. *Text mining* adalah

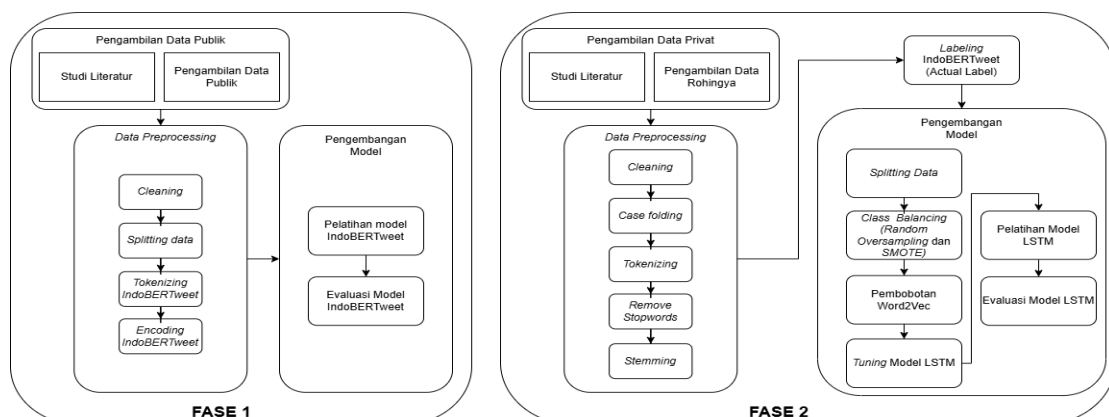
proses mengekstraksi pola atau informasi berguna dari data tekstual tidak terstruktur, seperti dokumen, untuk mengungkapkan pengetahuan baru yang belum diketahui atau terdokumentasi[5].

Analisis sentimen memakai LSTM merupakan bagian dari metode *text mining* yang menggunakan jaringan saraf tiruan dalam *deep learning* dirancang untuk mengolah data berurutan dan mengingat informasi jangka panjang [6]. Metode ini telah menunjukkan hasil yang menjanjikan dalam berbagai konteks penelitian. Mutmatimah dkk. [7] berhasil mencapai akurasi 92,51% dalam menganalisis sentimen pengguna aplikasi PeduliLindungi. Pendekatan serupa juga digunakan oleh Rahman dkk. [8] dalam analisis sentimen COVID-19 di Twitter, dengan menerapkan *word embedding* sebelum proses *training* LSTM, yang menghasilkan akurasi 81%. Penelitian tersebut membuktikan superioritas LSTM dibandingkan RNN dan Naive Bayes dengan selisih performa sekitar 8% [8]. Implementasi metode pembelajaran mesin untuk analisis sentimen juga telah dilakukan dengan berbagai pendekatan. Parameswari dan Prihandoko [9] menggunakan CNN untuk menganalisis sentimen lingkungan hidup di Kota Depok dengan akurasi 86%. Sementara itu, Maulana dkk. [10] menerapkan LSTM dengan *word embedding* untuk menganalisis sentimen terhadap implementasi kurikulum Merdeka, menghasilkan akurasi 81% dan mengembangkan *dashboard* visualisasi untuk keperluan *feedback*. Pradana dkk.[11] mengombinasikan LSTM dengan Word2Vec dalam analisis sentimen pemindahan ibu kota Indonesia, mencapai akurasi 95% dengan klasifikasi *biner*, menunjukkan efektivitas LSTM dalam menganalisis sentimen dengan dua kategori.

Terkait pengamatan sikap publik terhadap kehadiran pengungsi Rohingya di Indonesia, penelitian Ananda dan Suryono [12] menggunakan algoritma SVM dan Naive Bayes dengan klasifikasi *biner* menghasilkan akurasi masing-masing 76% dan 70%. Meskipun penelitian tersebut telah membuka jalan dalam analisis sentimen terkait isu Rohingya, masih terdapat ruang untuk peningkatan akurasi mengingat kompleksitas dan sensitivitas topik ini.

Penelitian ini bertujuan mengisi kesenjangan dalam analisis sentimen terkait isu Rohingya di Indonesia, dimana belum ada penelitian yang menggunakan LSTM untuk menganalisis sentimen di platform X (Twitter). LSTM dipilih karena telah terbukti unggul dalam klasifikasi sentimen dibandingkan metode klasik seperti SVM dan Naive Bayes. Dengan menggunakan data dari *web-scraping* platform X selama bulan Desember 2023 dengan kata kunci "rohingya", penelitian ini akan mengimplementasikan LSTM untuk mengklasifikasikan sentimen positif dan negatif masyarakat Indonesia. Melalui pendekatan ini, diharapkan tercapai pemahaman yang lebih intensif mengenai opini publik terhadap pengungsi Rohingya dan memperlihatkan efektivitas *neural network* dalam kajian sentimen media sosial.

2. METODE PENELITIAN



Gambar 1. Rangkaian Metode Penelitian

Rangkaian metode pada penelitian yang akan dilaksanakan dibagi menjadi 2 fase seperti pada Gambar 1. Fase 1 fokus pada melatih model IndoBERTweet, dimana model ini digunakan

sebagai pondasi utama pelabelan data tweet Rohingya untuk pemodelan LSTM, sementara itu fase 2 fokus pada model LSTM yang menggunakan data Rohingya.

2.1. Fase 1

a. Pengumpulan Data

Tahap pengumpulan data diawali dengan melakukan studi literatur secara komprehensif untuk memahami konsep analisis sentimen, model BERT, serta kajian terdahulu yang berkaitan. *Dataset* yang digunakan adalah IndoNLU smsa_doc-sentiment-prosa, sebuah *dataset* terstandar untuk analisis sentimen berbahasa Indonesia berupa teks yang dikumpulkan dari berbagai macam sumber *online* dengan label sentimen positif, negatif, serta netral yang dilakukan oleh beberapa ahli bahasa Indonesia[13]. *Dataset* dari IndoNLU pernah digunakan pada penelitian oleh Roni dkk. [14] dengan hasil akurasi 90%. Pada penelitian ini digunakan sentimen positif dan negatif dari *dataset* IndoNLU smsa_doc-sentiment-prosa.

b. Data Preprocessing

Dataset IndoNLU smsa_doc-sentiment-prosa sudah dilakukan *preprocessing* berupa *case folding* sebelumnya. Pada tahap ini, *preprocessing* yang dilakukan tokenizing dan *encoding* khusus dari IndoBERTweet. Dengan demikian *preprocessing* yang dilakukan adalah *cleaning* untuk menghapus duplikat dan data kosong, kemudian dilakukan pembagian(*splitting*), data dibagi menjadi 15% *testing data*, 15% *validation data*, dan 70% *training data*. Terakhir setiap data dilakukan *tokenizing* dan *encoding* IndoBERTweet.

c. Pengembangan Model IndoBERTweet

IndoBERTweet merupakan model *BERT (Bidirectional Encoder Representations from Transformers)* yang telah dioptimalkan untuk bahasa Indonesia dengan fokus utama pada *tweet* [15]. Model ini menggunakan arsitektur Transformer untuk memahami konteks dalam kalimat dan memberikan representasi yang untuk setiap kata. IndoBERTweet dilatih dengan 409 juta *token* yang diambil dari Twitter/X dimana data ini lebih banyak 2 kali dibandingkan dengan jumlah data yang digunakan untuk pengembangan IndoBERT [16]. Untuk mengukur hasil evaluasi IndoBERTweet metrik yang digunakan dalam evaluasi model adalah *F1 score*, *accuracy*, *recall*, *precision* serta visualisasi *confusion matrix*. Dalam evaluasi model klasifikasi, *True Positive* (TP) serta *True Negative* (TN) menggambarkan sampel yang diprediksi benar sebagai positif dan negatif. *False Positive* (FP) serta *False Negative* (FN) yaitu sampel yang salah diprediksi. Sementara itu, *confusion matrix* menampilkan distribusi FN, TN, TP, dan FP. Adapun rumus metrik yaitu [6]:

$$Accuracy = \frac{TN+TP}{FN+TN+TP+FP} \quad (1)$$

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

$$Recall = \frac{TP}{FN+TP} \quad (3)$$

$$F1\ Score = 2 \times \frac{Precision * Recall}{Recall + Precision} \quad (4)$$

2.2. Fase 2

a. Pengumpulan Data

Pengumpulan data diawali dengan studi literatur untuk memahami konteks data dalam hal ini Rohingya. Data diambil dari platform X dengan metode *web-scraping*. Alat yang digunakan yaitu *twikit* dengan bahasa Python. Pengambilan data dilakukan pada rentang waktu bulan Desember 2023 dengan *keyword* ‘*rohingya*’ untuk bahasa Indonesia. Jumlah data yang berhasil dikumpulkan sebanyak 17.613 *tweet*.

b. Data Preprocessing

Di fase 2, proses *preprocessing* yang dilakukan adalah *cleaning* (membersihkan teks dari *noise* seperti URL, *hashtag*, *mention*, *emoji*, dan karakter khusus), *case folding* (standarisasi huruf kecil), *tokenizing* (memenggal menjadi kata-kata terpisah), *remove stopwords* (menghapus kata tanpa makna), dan *stemming* (meyesuaikan kata menjadi bentuk dasar). Tahap *preprocessing* diperlukan mengingat data twitter memiliki banyak *noise* yang dapat mempengaruhi hasil

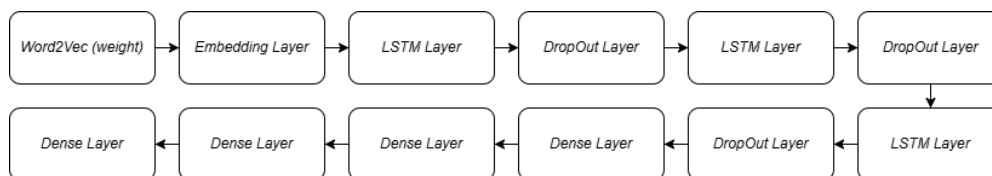
labeling yang nantinya akan dilakukan.

c. *Labeling*

Proses *labeling* dilakukan menggunakan model pra-terlatih besar untuk dijadikan sebagai *annotator* [17]. Model yang digunakan yaitu model IndoBERTweet yang telah dilatih pada fase 1 dengan *dataset* IndoNLU smsa_doc-sentiment-prosa. Model IndoBERTweet yang telah dilatih kemudian digunakan sebagai *annotator* pelabelan data *tweet* Rohingya dalam pemodelan LSTM. Data *tweet* Rohingya yang diperoleh akan melalui proses *cleaning* kemudian digunakan sebagai *input* ke model IndoBERTweet yang telah dilatih untuk diprediksi sehingga memperoleh *actual label* dari data *tweet* tersebut. Hasil dari pelabelan ini berupa label dari *tweet* Rohingya dengan sentimen positif dan negatif.

d. Pengembangan Model LSTM

Tahap pengembangan model pada fase kedua dimulai dengan *splitting data* menjadi data *training*, *validation*, dan *testing*. *Splitting data* ini dilakukan dengan rasio 70% *training*, 15% *validation*, dan 15% *testing*. Selanjutnya dilakukan *class balancing* menggunakan teknik *Random Oversampling* dan *SMOTE* untuk mengatasi ketidakseimbangan kelas dalam *dataset*. *Random Oversampling* menambahkan salinan data dari kelas minoritas untuk menyeimbangkan *dataset* [18], sedangkan metode *SMOTE* menciptakan data tambahan bagi kelas minoritas dengan menginterpolasi titik-titik data yang ada [18]. Proses dilanjutkan dengan pembobotan kata menggunakan Word2Vec untuk mendapatkan representasi vektor kata [19]. Adapun arsitektur model LSTM dan Word2Vec di Gambar 2.



Gambar 2. Struktur Model

Word2Vec akan menghasilkan *embedding*, hasil *embedding* ini digunakan sebagai matriks *embedding* pada *layer embedding* model LSTM. Word2Vec yang digunakan pada penelitian ini yaitu CBOW (*Continuous Bag Of Word*). Pada penelitian oleh Af'idah dkk. [20] CBOW mampu memberikan hasil akurasi yang lebih tinggi daripada variasi Word2Vec lain yaitu *Skip-gram*. Parameter yang digunakan untuk Word2Vec dan model LSTM ditentukan dengan melakukan *auto tuning* menggunakan *keras tuner*. *Keras tuner* adalah *library* untuk mengoptimalkan *hyperparameter* secara otomatis, sehingga meningkatkan akurasi, kecepatan, dan generalisasi model [21]. Dalam penelitian ini, *keras tuner* digunakan untuk menemukan *hyperparameter* terbaik pada Word2Vec dan model LSTM, yang kemudian diterapkan saat pelatihan model LSTM. Model LSTM berguna dalam menangkap konteks dari teks karena kemampuannya mengingat informasi dalam jangka waktu yang panjang. Pada LSTM, terdapat tiga gerbang utama yang mengatur aliran informasi, *forget gate*, *input gate*, dan *output gate*. Adapun perhitungan dari

masing-masing gerbang dapat dijabarkan dalam rumus berikut [6]:

a. *Forget gate*

$$f_t = \sigma(W_{xf}x^{(t)} + W_{hf}h^{(t-1)} + b_f) \quad (5)$$

f_t = forget gate saat t

$x^{(t)}$ = vektor input saat t

$h^{(t-1)}$ = hidden state saat (t - 1)

W_{xf} = bobot dalam forget gate, (h x x)

W_{hf} = bobot dalam forget gate, (h x h)

b_f = bias dalam forget gate

b. *Input gate*

$$i_t = \sigma(W_{xi}x^{(t)} + W_{hi}h^{(t-1)} + b_i) \quad (6)$$

i_t = input gate saat t

$x^{(t)}$ = vektor input saat t

$h^{(t-1)}$ = hidden state saat (t - 1)

W_{xi} = bobot dalam input gate, (h x x)

W_{hi} = bobot dalam input gate, (h x h)

c. *Output gate*

$$o_t = \sigma(W_{xo}x^{(t)} + W_{ho}h^{(t-1)} + b_o) \quad (7)$$

o_t = output state saat t

$x^{(t)}$ = vektor input saat t

$h^{(t-1)}$ = hidden state saat (t - 1)

W_{xo} = bobot dalam output gate, (h x x)

W_{ho} = bobot dalam output gate, (h x h)

b_o = bias dalam output gate (h)

Proses dilanjutkan dengan pelatihan model LSTM menggunakan data yang telah diproses, diikuti dengan evaluasi untuk mengukur performa model dalam menangani *tweet* tentang Rohingya dengan metrik yang sama seperti fase pertama.

3. HASIL DAN PEMBAHASAN

Hasil serta pembahasan dibagi menjadi fase 1 dan fase 2. Fase 1 fokus pada pembuatan model IndoBERTweet yang digunakan sebagai *annotator* untuk memberikan label. Fase 2 fokus pada pengembangan model LSTM pada *dataset* Rohingya.

3.1 Fase 1

- Data untuk pelatihan IndoBERTweet merupakan data *open source* IndoNLU smsa_doc-sentiment-prosa. Pada dataset ini diambil sentimen positif dan negatif dengan total data sebanyak 11.324 baris data dengan rincian 4.005 data sentimen positif dan 7.319 data sentimen negatif.
- Data IndoNLU smsa_doc-sentiment-prosa kemudian dilakukan *preprocessing* berupa *cleaning*. *Dataset* ini telah memiliki label yang diberikan oleh ahli bahasa Indonesia [13]. Label sentimen diubah menjadi representasi angka, 0 sebagai sentimen positif dan 1 sebagai sentimen positif seperti pada Tabel 1.

Tabel 1. Dataset IndoNLU smsa_doc-sentiment-prosa

Original_text	Cleaned_text	Label	Label_num
di restoran ini makanan yang disediakan banyak sekali dan harga cukup terjangkau , suasana yang enak dan nyaman	di restoran ini makanan yang disediakan banyak sekali dan harga cukup terjangkau suasana yang enak dan nyaman	Positif	0
masa ya harus datang ke gerai indosat dan komplain marah-marah di sana	masa ya harus datang ke gerai indosat dan komplain marah-marah di sana	Negatif	1

Data yang sudah melalui *preprocessing* kemudian dibagi menjadi 15% *testing data*, 15% *validation data*, dan 70% *training data*. Setelah dibagi, data diproses oleh *tokenizer* dan *encode* dari IndoBERTweet.

- c. Terakhir adalah proses pelatihan model IndoBERTweet. Pelatihan dilakukan dalam 3 *epochs*, dengan ukuran *batch* untuk training yaitu 32 dan ukuran *batch* untuk validasi yaitu 64. Hasil dari pelatihan pada Tabel 2.

Tabel 2. *Training* IndoBERTweet

<i>Epoch</i>	<i>Training Loss</i>	<i>Validation Loss</i>
1	0.1499	0.1395
2	0.1316	0.1431
3	0.0366	0.1921

Setelah pelatihan, evaluasi model IndoBERTweet dilakukan menggunakan *testing data*, dan hasilnya disajikan dalam Tabel 3, yang menunjukkan rata-rata performa serta performa masing-masing kelas.

Tabel 3. Hasil Evaluasi IndoBERTweet

	Recall	Precision	Skor f1	Support
Positif	0.9353	0.9790	0.9567	1098
Negatif	0.9634	0.8908	0.9257	601
Accuracy			0.9453	1699
Rata-rata makro	0.9494	0.9349	0.9412	1699
Rata-rata tertimbang	0.9453	0.9478	0.9457	1699

Dari Tabel 3 dapat dilihat bahwa model IndoBERTweet mempunyai performa yang cukup baik di masing-masing kelas maupun secara keseluruhan data. Pada rata-rata makro model mampu memberikan 93.49% *precision*, 94.94% *recall*, 94.12% *f1 score* dan 94.53% *accuracy*. Dengan hasil ini maka model cukup baik untuk digunakan pada tahap *labeling* di fase 2 dimana model ini akan digunakan sebagai *annotator* dalam pelabelan *tweet* Rohingya. Performa yang sangat baik pada kedua kelas serta metrik yang seimbang menunjukkan bahwa *class balancing* tidak diperlukan pada fase 1 ini, karena model sudah dapat menangani distribusi data yang ada dengan efektif.

3.2 Fase 2

- a. Pengumpulan data Rohingya pada platform X dilakukan dengan *twikit* menghasilkan data sebanyak 17.613 *tweet*. Data diambil selama bulan Desember 2023 dengan keyword 'rohingya' dan parameter bahasa Indonesia.
- b. *Data preprocessing* yang dilakukan berupa pembersihan data (*cleaning*), perubahan huruf kecil (*case folding*), *tokenizing*, *removing stopwords*, serta *stemming*. Pembersihan data dan perubahan huruf kecil dilakukan untuk menghilangkan *noise* serta standarisasi data, kemudian diubah dalam bentuk *token* yang terlihat di Tabel 4.

Tabel 4. Hasil *Cleaning*, *Case Folding*, *Tokenizing*

<i>Original_text</i>	<i>Cleaning & case folding</i>	<i>Tokenizing</i>
@tanyakanrl Rohingya jg hrs rajin di up biar ada antisipasi dr pemerintah.	rohingya jg hrs rajin di up biar ada antisipasi dr pemerintah	['rohingya', 'jg', 'hrs', 'rajin', 'di', 'up', 'biar', 'ada', 'antisipasi', 'dr', 'pemerintah']
Desember mereka ada rencana masuk lagi, tolong diawasi. Mereka kemungkinan besar memang sindikat perdagangan manusia nawarin pelayaran ke Rohingya @Polresacehjava @PolsekIdiRayeuk @polresatim	desember mereka ada rencana masuk lagi tolong diawasi mereka kemungkinan besar memang sindikat	['desember', 'mereka', 'ada', 'rencana', 'masuk', 'lagi', 'tolong', 'diawasi', 'mereka', 'kemungkinan', 'besar', 'memang', 'sindikat']

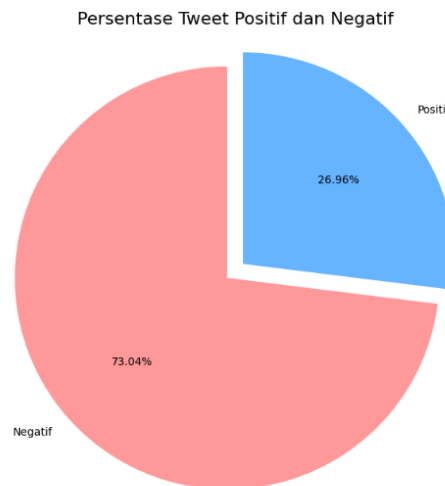
@DivHumas_Polri https://t.co/vf4np0XpqP	@_TNIAL_	perdagangan manusia nawarin pelayaran ke Rohingya	'perdagangan', 'manusia', 'nawarin', 'pelayaran', 'ke', 'rohingya']
--	----------	---	---

Data yang sudah berubah menjadi bentuk *token* kemudian dihapus bagian *stopwords* lalu kata-kata dalam data diubah menjadi bentuk dasarnya dengan *stemming* seperti di Tabel 5.

Tabel 5. Hasil *Removing Stopwords* dan *Stemming*

<i>Removing stopwords</i>	<i>Stemming</i>
['rohingya', 'hrs', 'rajin', 'biar', 'antisipasi', 'pemerintah']	['rohingya', 'hrs', 'rajin', 'biar', 'antisipasi', 'perintah']
['desember', 'rencana', 'masuk', 'diawasi', 'kemungkinan', 'besar', 'memang', 'sindikata', 'perdagangan', 'manusia', 'nawarin', 'pelayaran', 'rohingya']	['desember', 'rencana', 'masuk', 'awas', 'mungkin', 'besar', 'memang', 'sindikat', 'dagang', 'manusia', 'nawarin', 'layar', 'rohingya']

- c. Pelabelan dilakukan pada data yang sudah dibersihkan (*cleaned*) yang dijadikan *input* pada model IndoBERTweet di fase 1 untuk mendapatkan label *ground truth* dengan otomatis. Label yang digunakan adalah representasi angka, 0 pada sentimen positif serta 1 pada sentimen negatif.



Gambar 3. Hasil Pelabelan Sentimen dengan IndoBERTweet

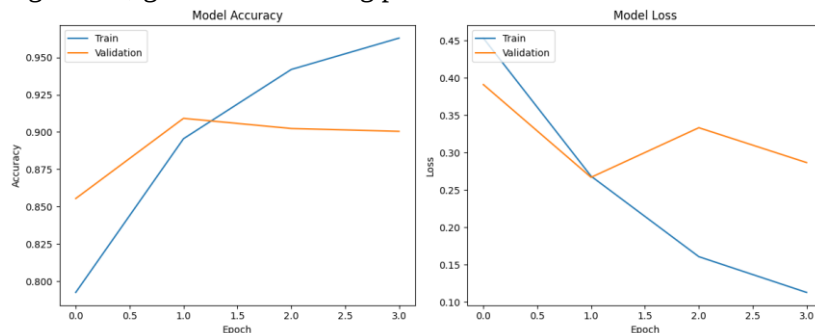
Gambar 3 menunjukkan hasil pelabelan dari 17.613 *tweet* didapatkan 73.04% data memiliki sentimen negatif dan 26.96% memiliki sentimen positif pada Gambar 3. Proporsi yang tidak seimbang antara 2 kelas ini akan diseimbangkan dengan *SMOTE* dan *Random Oversampling* pada tahap berikutnya. Dari hasil pelabelan dilakukan visualisasi dengan *word cloud* pada Gambar 4 untuk mengetahui berbagai kata yang muncul di setiap sentimen pada data yang sudah melalui *stemming*.



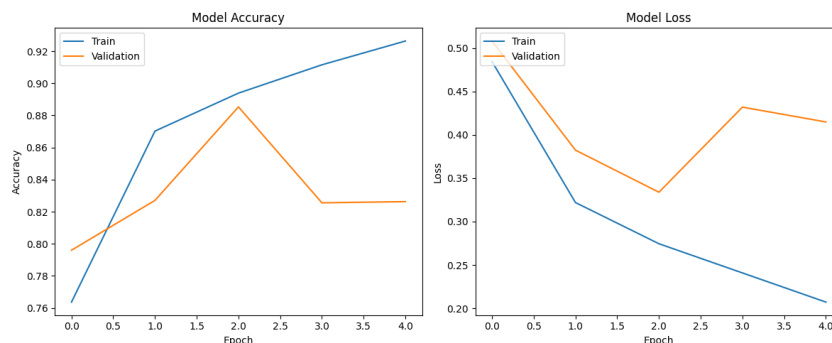
Gambar 4. *Word Cloud* Sentimen Positif dan Negatif

Dari Gambar 4, sentimen positif menunjukkan kata-kata seperti “rohingya”, “arti”, dan “bantu”, yang mencerminkan simpati terhadap pengungsi Rohingya. Sedangkan pada sentimen negatif, kata-kata seperti “rohingya”, “problem”, dan “issue” mencerminkan pandangan pengungsi sebagai masalah, dengan aspek agama terlihat pada kata “muslim” dan “hindu”.

- d. Pengembangan LSTM dimulai dengan membagi data hasil preprocessing menjadi 15% testing data (2.643), 15% validation data (2.642), serta 70% training data (12.329). Data training yang terdiri dari 9.005 sentimen negatif dan 3.324 sentimen positif diseimbangkan menggunakan Random Oversampling dan SMOTE sesuai jumlah sentimen negative, yang selanjutnya digunakan dalam auto-tuning agar memperoleh parameter terbaik yang selanjutnya digunakan pada training LSTM, grafik hasil training pada Gambar 5 dan 6.

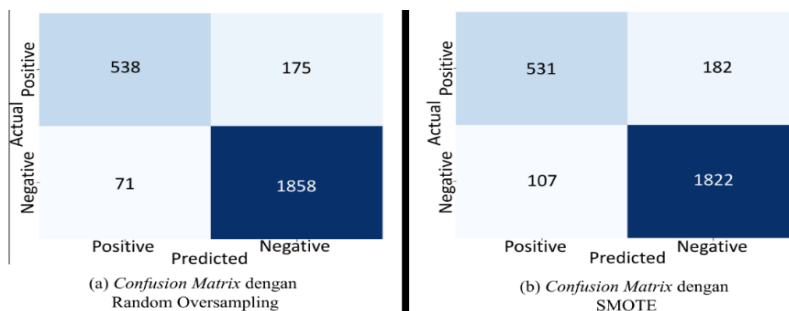


Gambar 5. Training LSTM dengan data Random Oversampling



Gambar 6. Training LSTM dengan data SMOTE

Pada tahap *training* menggunakan *callback early_stopping* pada *validation loss* untuk memberhentikan *training* ketika performa model tidak membaik pada *epochs* terakhir, pada Gambar 5 terlihat Random Oversampling berakhir pada *epochs* ketiga dan pada Gambar 6 terlihat SMOTE berakhir pada *epochs* keempat. Proses selanjutnya adalah evaluasi menggunakan visualisasi *confusion matrix* di Gambar 7 serta *classification report* di Tabel 6. Tahap evaluasi juga membandingkan hasil dari penyeimbang Random Oversampling dan SMOTE.



Gambar 7. Confusion Matrix Random Oversampling dan SMOTE

Tabel 6. Hasil Evaluasi *Classification Report* LSTM

Penyeimbang Data		<i>recall</i>	<i>precision</i>	skor f1	<i>support</i>
<i>Random Oversampling</i>	positif	0.7546	0.8834	0.8139	713
	negatif	0.9632	0.9139	0.9379	1929
	<i>accuracy</i> (semua kelas)			0.9069	2642
	rata-rata makro	0.8589	0.8987	0.8759	2642
	rata-rata tertimbang	0.9069	0.9057	0.9044	2642
	semua kelas	0.9632	0.9139	0.9379	-
<i>SMOTE</i>	positif	0.7447	0.8323	0.7861	713
	negatif	0.9445	0.9092	0.9265	1929
	<i>accuracy</i> (semua kelas)			0.8906	2642
	rata-rata makro	0.8446	0.8707	0.8563	2642
	rata-rata tertimbang	0.8906	0.8884	0.8886	2642
	semua kelas	0.9445	0.9092	0.9265	-

Hasil *classification report* pada Tabel 6 memperlihatkan perbandingan antara metode *Random Oversampling* dan *SMOTE* pada kelas "positif", *Random Oversampling* mencapai *precision* 0.8834, skor f1 0.8139, serta *recall* 0.7546, dibandingkan *SMOTE* yang menghasilkan *precision* 0.8323, skor f1 0.7861, serta *recall* 0.7447. Sementara itu, untuk kelas "negatif", keduanya memiliki hasil yang cukup dekat, namun *Random Oversampling* tetap unggul dengan *precision* 0.9139, skor f1 0.9379, serta *recall* 0.9632 dibandingkan *SMOTE* dengan *precision* 0.9092, skor f1 0.9265, serta *recall* 0.9445. Secara keseluruhan, *Random Oversampling* memberikan *recall* 0.9632, *precision* 0.9139, skor f1 0.9379, serta akurasi 0.9069 lebih tinggi dibandingkan *SMOTE* yang memiliki *recall* 0.9445, skor f1 0.9265, *precision* 0.9092, serta akurasi 0.8906 membuat *Random Oversampling* lebih efektif dalam menjaga keseimbangan performa antar kelas. Hasil ini juga terlihat pada Gambar 7 *confusion matrix* dimana data dengan *Random Oversampling* menghasilkan *true label* sedikit lebih banyak daripada hasil dengan *SMOTE*.

4. KESIMPULAN DAN SARAN

Model IndoBERTweet yang telah dilatih pada dataset IndoNLU mampu memperlihatkan performa yang unggul dalam klasifikasi sentimen dengan akurasi keseluruhan 94.53% pada fase 1. Model ini kemudian digunakan untuk melabeli data tentang pengungsi Rohingya menghasilkan proporsi sentimen negatif yang dominan. Pada fase 2, penerapan LSTM dengan penyeimbangan data menggunakan metode *Random Oversampling* dan *SMOTE* menunjukkan bahwa *Random Oversampling* memberikan hasil yang lebih baik dengan akurasi 90.69% dibandingkan 89.06% dari *SMOTE*. Performa ini juga sama metrik *precision*, *recall*, serta *f1 score* yang konsisten lebih tinggi pada metode *Random Oversampling*, terutama pada kelas sentimen positif. Hasil ini menunjukkan bahwa penyeimbangan data sangat berpengaruh terhadap performa klasifikasi pada data yang tidak seimbang. Hasil yang baik pada model LSTM menunjukkan bahwa model IndoBERTweet mampu menangkap pola sentimen serta memberikan label yang relevan.

Sentimen mengenai pengungsi Rohingya menunjukkan kecenderungan negatif dengan presentase 73.04% bersentimen negatif dan 26.96% bersentimen positif dari total 17.613 data *tweet*. Sentimen negatif ini mengindikasikan kekhawatiran masyarakat terkait dampak kedatangan pengungsi. Sebagai saran, proses *preprocessing* dapat ditambahkan normalisasi kata dan penambahan *stopwords*, termasuk *slang*, untuk mengurangi *noise*. Pengujian lebih lanjut dengan variasi parameter pada model LSTM dan metode penyeimbangan data juga dianjurkan untuk hasil analisis sentimen yang lebih optimal.

DAFTAR PUSTAKA

- [1] W. O. Zalmatin, Ruslan, and S. Syukur, "Konflik Rohingya dan Pengakuan Kewarganegaraannya," *Edu Sociata J. Pendidik. Sociol.*, vol. 6, no. 2, pp. 558–568, 2023.
- [2] C. A. Sopamena, "Pengungsi Rohingya Dan Potensi Konflik & Kemajemukan Horizontal Di Aceh," *J. Caraka Prabhu*, vol. 7, no. 2, pp. 85–115, 2023.
- [3] R. Maria, R. U. Umayah, S. Mahardinny, D. N. Kalana, and D. D. Saputra, "Analisis Sentimen Persepsi Masyarakat Terhadap Penggunaan Aplikasi My Pertamina Pada Media Sosial Twitter Menggunakan Metode Naïve Bayes Classifier," *J. Komput. Antart.*, vol. 1, no. 1, pp. 1–10, 2023.
- [4] A. Firdaus, W. I. Firdaus, P. Studi, T. Informatika, M. Digital, and P. N. Sriwijaya, "Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi (Sebuah Ulasan)," *J. JUPITER (Jurnal Penelit. dan Ilmu Komputer)*, vol. 13, no. 1, pp. 66–78, 2021.
- [5] S. M. Fani, R. Santoso, and S. Suparti, "Penerapan Text Mining untuk Melakukan Clustering Data Tweet Akun Blibli pada Media Sosial Twitter Menggunakan K-means Clustering," *J. Gaussian*, vol. 10, no. 4, pp. 583–593, 2021.
- [6] S. Raschka and V. Mirjalili, *Python Machine Learning*, 3rd ed. Packt Publishing Ltd., 2019.
- [7] S. Mutmatimah, Khairunnas, and Khairunnisa, "Metode Deep Learning LSTM dalam Analisis Sentimen Aplikasi PeduliLindungi," *J. Comput. Sci. Informatics*, vol. 1, no. 1, pp. 9–19, 2024.
- [8] M. Z. Rahman, Y. A. Sari, and N. Yudistira, "Analisis Sentimen Tweet COVID-19 menggunakan Word Embedding dan Metode Long Short-Term Memory (LSTM)," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 5, no. 11, pp. 5120–5127, 2021.
- [9] P. L. Parameswari and Prihandoko, "Penggunaan Convolutional Neural Network Untuk Analisis Sentimen Opini Lingkungan Hidup Kota Depok Di Twitter," *J. Ilm. Teknol. dan Rekayasa*, vol. 27, no. 1, pp. 29–42, 2022.
- [10] A. R. Maulana, S. H. Wijoyo, and Y. T. Mursityo, "Analisis Sentimen Kebijakan Penerapan Kurikulum Merdeka Sekolah Dasar dan Sekolah Menengah pada Media Sosial Twitter dengan Menggunakan Metode Word Embedding Dan Long Short-term Memory Networks (LSTM)," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 3, pp. 523–530, 2023.
- [11] Y. A. Pradana, I. Cholissodin, and D. Kurnianingtyas, "Analisis Sentimen Pemindahan Ibu Kota Indonesia pada Media Sosial Twitter menggunakan Metode LSTM dan Word2Vec," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 7, no. 5, pp. 2389–2397, 2023.
- [12] D. Ananda and R. R. Suryono, "Analisis Sentimen Publik Terhadap Pengungsi Rohingya di Indonesia dengan Metode Support Vector Machine dan Naïve Bayes," *J. Media Inform. Budidarma*, vol. 8, no. 2, p. 748, 2024.
- [13] A. F. Pratama, T. B. Kurniawan, Misinem, and D. A. Dewi, "Implementasi Analisis Sentimen dan Model Deep Learning Untuk Prediksi Harga Bitcoin," *JUPITER J. Penelit. Ilmu Dan Teknol. Komput.*, vol. 15, no. 1b, pp. 403–412, 2023.
- [14] R. Merdiansah, S. Siska, and A. Ali Ridha, "Analisis Sentimen Pengguna X Indonesia Terkait Kendaraan Listrik Menggunakan IndoBERT," *J. Ilmu Komput. dan Sist. Inf.*, vol. 7, no. 1, pp. 221–228, Mar. 2024.
- [15] J. F. Kusuma and A. Chowanda, "Indonesian Hate Speech Detection Using IndoBERTweet and BiLSTM on Twitter," *JOIV Int. J. Informatics Vis.*, vol. 7, no. 3, pp. 773–780, 2023.
- [16] F. Koto, J. H. Lau, and T. Baldwin, "INDOBERTWEET: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization," *EMNLP 2021-2021 Conference on Empirical Methods in Natural Language Processing, Proceedings*, 2021, pp. 10660–10668.
- [17] Z. Tan et al., "Large Language Models for Data Annotation: A Survey," in *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, 2024, pp. 930–957.
- [18] U. Hasanah, A. Mohamad Soleh, and K. Sadik, "Effect of Random Under sampling , Oversampling, and SMOTE on the Performance of Cardiovascular Disease Prediction

- Models,” *J. Mat. Stat. dan Komputasi*, vol. 21, no. 1, pp. 88–102, 2024.
- [19] M. V. Shyahrin, Y. Sibaroni, and D. Puspandari, “Penerapan Metode Long Short-Term Memory dan Word2Vec dalam Analisis Sentimen Ulasan pada Aplikasi Ferizy,” *Techno.Com*, vol. 22, no. 4, pp. 833–842, 2023.
- [20] D. I. Af'idah, D. Dairoh, S. F. Handayani, and R. W. Pratiwi, “Pengaruh Parameter Word2Vec terhadap Performa Deep Learning pada Klasifikasi Sentimen,” *J. Inform. J. Pengemb. IT*, vol. 6, no. 3, pp. 156–161, 2021.
- [21] G. Y. Christiawan, R. A. Putra, A. Sulaiman, E. Poerbaningtyas, and S. W. P. Listio, “Penerapan Metode Convolutional Neural Network Dalam Mengklasifikasikan Penyakit Daun Tanaman Padi (CNN),” *J-INTECH (Journal Inf. Technol.*, vol. 11, no. 2, pp. 294–306, 2023.