

ABSTRAK

Penelitian ini bertujuan menerapkan metode Bidirectional Long Short-Term Memory (BiLSTM) untuk mendeteksi indikasi depresi berdasarkan teks berbahasa Indonesia. Deteksi dini terhadap depresi penting dilakukan mengingat tingginya jumlah penderita dan keterbatasan akses terhadap layanan kesehatan mental. Penelitian menggunakan dataset publik sebanyak 6.976 data berlabel, terdiri atas 733 data depresi dan 6.243 data tidak depresi. Sebelum pelatihan model, data diproses melalui tahapan *preprocessing* yang meliputi *cleaning*, *case folding*, normalisasi kata, *stopword removal*, *stemming*, tokenisasi, dan *padding sequence*. Untuk mengatasi ketidakseimbangan data diterapkan teknik *class weight*. Model BiLSTM dibangun menggunakan TensorFlow dan Keras, kemudian diimplementasikan menggunakan Streamlit sebagai antarmuka pengguna. Proses klasifikasi menggunakan fungsi aktivasi sigmoid dengan nilai *threshold* sebesar 0,5. Evaluasi model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, dan *Area Under Curve* (AUC). Hasil pengujian menunjukkan nilai *accuracy* 99,81%, *precision* 98,21%, *recall* 100%, *F1-score* 99,10%, dan AUC 1,0000. Hasil penelitian menunjukkan bahwa metode BiLSTM mampu menghasilkan performa klasifikasi yang sangat baik sehingga berpotensi menjadi alat bantu skrining awal indikasi depresi berdasarkan teks, tanpa menggantikan diagnosis oleh psikolog atau psikiater.

Kata kunci: *Bidirectional Long Short-Term Memory* (BiLSTM), deteksi depresi, klasifikasi teks, pemrosesan bahasa alami, kesehatan mental.

ABSTRACT

This study aims to apply the Bidirectional Long Short-Term Memory (BiLSTM) method to detect depression indications from Indonesian-language text. Early detection of depression is essential due to the increasing number of affected individuals and the limited access to mental health services. The study employed a publicly available dataset consisting of 6,976 labeled text samples, including 733 depression samples and 6,243 non-depression samples. Before model training, the data underwent several preprocessing stages, including cleaning, case folding, word normalization, stopword removal, stemming, tokenization, and padding sequence. To address the class imbalance, the class weight technique was applied. The BiLSTM model was developed using TensorFlow and Keras and implemented through Streamlit as a user interface. The classification process utilized a sigmoid activation function with a 0.5 threshold to distinguish between depression and non-depression classes. Model performance was evaluated using accuracy, precision, recall, F1-score, and Area Under the Curve (AUC). The experimental results achieved an accuracy of 99.81%, precision of 98.21%, recall of 100%, F1-score of 99.10%, and an AUC of 1.0000. These findings indicate that the proposed BiLSTM model provides excellent classification performance and has the potential to support the early screening of depression based on textual data. However, the system is intended only as a screening tool and does not replace professional diagnosis by psychologists or psychiatrists.

Keywords: *Bidirectional Long Short-Term Memory (BiLSTM), depression detection, text classification, natural language processing, mental health.*

