

BAB V PENUTUP

5.2 Kesimpulan

Berdasarkan hasil penelitian mengenai analisis sentimen pada dataset ulasan pengguna aplikasi *Grab* Indonesia, dapat disimpulkan bahwa model *Bi-Directional LSTM* dengan penambahan *Multi-Head Attention* menunjukkan performa yang lebih baik dibandingkan dengan model *Stacked LSTM*. *Bi-Directional LSTM*, yang memproses data sekuensial dengan mempertimbangkan konteks dari kedua arah, berhasil meningkatkan akurasi validasi dan mengurangi kesalahan prediksi, terutama pada kelas negatif. Model ini memiliki *precision* dan *recall* yang lebih tinggi untuk kelas negatif, yang menunjukkan kemampuannya yang lebih baik dalam mengidentifikasi sentimen negatif pengguna. Selain itu, penerapan teknik *SMOTE* untuk menangani ketidakseimbangan kelas juga berkontribusi pada hasil yang lebih seimbang antara kedua kelas.

Pada model *Bi-Directional LSTM*, akurasi validasi tertinggi tercatat sebesar 87%, dengan *F1-score* untuk kelas negatif mencapai 0,90 dan untuk kelas positif 0,82, menunjukkan keseimbangan yang lebih baik antara *precision* dan *recall* pada kedua kelas. Sementara itu, model *Stacked LSTM* hanya mencapai akurasi 84%, dengan *F1-score* untuk kelas negatif 0,87 dan untuk kelas positif 0,78. Hasil ini menegaskan bahwa model *Bi-Directional LSTM* memiliki kemampuan yang lebih baik dalam memahami hubungan kontekstual dalam data ulasan yang bersifat kompleks dan sekuensial.

Secara keseluruhan, penelitian ini menunjukkan bahwa penggunaan *Bi-Directional LSTM* dengan *Multi-Head Attention* dapat menghasilkan model yang lebih efektif dalam menganalisis sentimen ulasan pengguna pada aplikasi *Grab*, dengan akurasi yang lebih tinggi dan kinerja yang lebih baik dalam mengklasifikasikan sentimen positif dan negatif.

5.1 Saran

Dari hasil yang didapatkan, terdapat beberapa saran untuk pengembangan penelitian di masa mendatang. Penelitian selanjutnya dapat mengeksplorasi model yang lebih unggul, seperti *Transformer* atau *BERT*, untuk meningkatkan performa

analisis sentimen. Selain itu, penggunaan dataset yang lebih besar dan beragam, dapat membantu meningkatkan generalisasi model dan juga dapat mencegah *overfitting* dikarenakan kurangnya data yang digunakan untuk *testing*. Teknik pra-pemrosesan lanjutan, seperti augmentasi data atau penggunaan leksikon sentimen, juga dapat diterapkan untuk meningkatkan akurasi. Optimisasi seperti *Word2vec* serta *hyperparameter* dengan metode seperti *Grid Search* yang dapat membantu menemukan konfigurasi terbaik untuk model. Dengan saran-saran tersebut, diharapkan penelitian di masa depan dapat memberikan kontribusi yang lebih besar dalam pengembangan analisis sentimen dan teknologi pemrosesan bahasa alami.

