

BAB V PENUTUP

5.1 Kesimpulan

Implementasi model IndoBERT untuk klasifikasi sentimen dan deteksi ujaran kebencian pada komentar YouTube berbahasa Indonesia telah berhasil dilakukan melalui proses *fine-tuning* dengan konfigurasi yang meliputi strategi pembobotan kelas serta arsitektur WeightedTrainer guna mengatasi ketidakseimbangan data. Kinerja model menunjukkan keunggulan dibandingkan *baseline*, dengan mencapai akurasi sebesar 80,0 persen dan *F1-score* makro 0,750 pada analisis sentimen, serta mampu mendeteksi 52 persen ujaran kebencian pada tugas deteksi meskipun akurasi keseluruhan lebih rendah. Analisis pola opini mengungkap dominasi sentimen negatif sebesar 58,6 persen dalam diskusi, namun hanya 14,3 persen di antaranya yang tergolong ujaran kebencian, menunjukkan bahwa sebagian besar respons negatif bersifat kritik argumentatif dan bukan serangan terhadap gender.

5.2 Saran

Berdasarkan hasil penelitian, beberapa pengembangan berikut dapat dipertimbangkan pada penelitian selanjutnya:

1. Perluasan dan Penyeimbangan Dataset

Ketidakseimbangan kelas pada ujaran kebencian masih memengaruhi performa model. Penelitian lanjutan disarankan menggunakan dataset yang lebih besar dan seimbang, baik melalui penambahan sumber data maupun teknik penyeimbangan seperti *oversampling*.

2. Peningkatan Kualitas Pelabelan Manual

Pelabelan komentar berpotensi bersifat subjektif, terutama pada batas antara keras dan ujaran kebencian implisit. Penggunaan lebih dari satu anotator serta pengukuran *inter-annotator agreement* dapat meningkatkan reliabilitas label

3. Eksplorasi Strategi Penanganan Imbalance

Penelitian selanjutnya dapat mengeksplorasi metode penanganan ketidakseimbangan kelas lainnya seperti oversampling, undersampling, focal loss, atau augmentasi data teks. Perbandingan beberapa pendekatan tersebut dapat membantu menentukan strategi yang paling efektif dalam meningkatkan performa model pada kelas minoritas.

4. Mengeksplorasi model transformel lain untuk Bahasa Indonesia

Penelitian selanjutnya dapat mempertimbangkan penggunaan model lain yang dirancang untuk teks media sosial, seperti IndoBERTweet, XLM-RoBERTa, atau arsitektur *hibrida transformer-LSTM*, sebagai pembanding tambahan untuk melihat variasi performa model pada teks informal dan komentar pendek.

5. Perluasan Topik dan Domain Data

Dataset penelitian ini berfokus pada satu topik, yaitu kesetaraan gender pada satu video. Untuk menguji kemampuan generalisasi model, penelitian berikutnya disarankan menggunakan data dari berbagai topik sosial dan berbagai sumber platform.