

BAB V KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan serangkaian proses penelitian yang telah dilakukan, mulai dari pengumpulan data, pra-pemrosesan, hingga implementasi dan evaluasi model IndoBERT untuk analisis sentimen terhadap berita mega korupsi Pertamina, dapat ditarik beberapa kesimpulan sebagai berikut:

a. Efektivitas Model IndoBERT:

Penerapan metode Deep Learning menggunakan model pra-latih (*pre-trained*) IndoBERT terbukti efektif dalam melakukan analisis sentimen pada teks berbahasa Indonesia yang tidak baku (media sosial). Model mampu memahami konteks kalimat secara semantik tanpa memerlukan proses *stemming* atau penghapusan stopwords, sehingga informasi penting terkait negasi dan struktur kalimat tetap terjaga.

b. Peningkatan Performa melalui Penanganan Imbalanced Data:

Masalah ketidakseimbangan kelas (*class imbalance*) yang dominan pada sentimen Negatif dan Positif berhasil ditangani menggunakan strategi *Custom Trainer* dengan mekanisme *Class Weighting*.

c. Wawasan Sentimen Publik:

Berdasarkan hasil klasifikasi, sentimen masyarakat terhadap kasus mega korupsi Pertamina didominasi oleh sentimen Negatif. Hal ini mencerminkan tingginya tingkat kekecewaan dan ketidakpercayaan publik terhadap integritas pengelolaan perusahaan negara tersebut.

d. Kinerja Arsitektur:

Penggunaan parameter fine-tuning dengan learning rate $3e-5$, batch size 16, dan 4 epoch memberikan hasil yang optimal dalam menjaga stabilitas proses

pembelajaran model (*convergence*), sehingga model mampu menghasilkan metrik evaluasi (F1-Score) yang seimbang antar kelas.

5.2 Saran

Penelitian ini masih memiliki keterbatasan yang dapat diperbaiki untuk pengembangan selanjutnya. Berikut adalah beberapa saran yang dapat dipertimbangkan bagi peneliti masa depan:

a. Perluasan Sumber Data:

Penelitian ini hanya berfokus pada data yang bersumber dari media sosial X (Twitter). Penelitian selanjutnya disarankan untuk menggabungkan dataset dari berbagai platform lain seperti Instagram, TikTok, atau kolom komentar portal berita daring untuk mendapatkan gambaran sentimen publik yang lebih komprehensif dan beragam.

b. Peningkatan Kualitas Pelabelan:

Pelabelan data pada penelitian ini menggunakan metode *Lexicon Based* yang kemudian divalidasi oleh model. Untuk meningkatkan validitas ground truth, disarankan menggunakan metode pelabelan manual oleh ahli bahasa atau pakar domain (*human annotation*) guna meminimalisir kesalahan pelabelan pada kalimat yang mengandung sarkasme atau makna implisit.

c. Eksperimen Model Lain:

Peneliti selanjutnya dapat membandingkan performa IndoBERT dengan model bahasa varian lain yang lebih baru atau lebih ringan, seperti IndoRoBERTa, XLM-R, atau DistilBERT, untuk melihat efisiensi komputasi berbanding dengan akurasi yang dihasilkan.

d. Penanganan Sarkasme Khusus:

Mengingat banyaknya komentar politik yang mengandung sarkasme, penelitian di masa depan dapat menambahkan modul khusus deteksi sarkasme (*sarcasm detection*) sebelum tahap analisis sentimen utama agar akurasi prediksi pada sentimen negatif terselubung dapat lebih ditingkatkan.