

ABSTRAK

Dokumen regulasi kelapa sawit Indonesia memiliki struktur teks yang panjang dan sarat terminologi hukum terminologi hukum, sehingga menyulitkan pemangku kepentingan memahami substansinya secara efisien. Penelitian ini mengeksplorasi efektivitas strategi *prompt engineering* pada peringkasan abstraktif regulasi kelapa sawit menggunakan dua *Large Language Model* (LLM), yaitu GPT-5.4-mini dan Llama 3.1:8B, terhadap 157 dokumen regulasi. Eksperimen mengombinasikan tiga strategi *prompting* (*zero-shot*, *one-shot*, *few-shot*) dengan empat varian *context configuration* yang dibentuk dari kombinasi dua komponen, yaitu *persona* dan *domain knowledge* (*general* sebagai baseline, *persona*, *domain knowledge*, dan *domain knowledge + persona*), menghasilkan 24 skenario dan 3.768 ringkasan yang dievaluasi menggunakan BERTScore, BLEU, ROUGE-L, dan METEOR. Hasil menunjukkan bahwa *zero-shot* mendominasi tiga dari empat metrik (BLEU 0,1469, ROUGE-L 0,3267, BERTScore 0,7571), karena penambahan contoh pada *one-shot* dan *few-shot* berpotensi memicu *overprompting* terhadap gaya bahasa dokumen contoh. Konfigurasi *persona* terbukti terbaik pada semua metrik, sementara *domain knowledge* justru menekan performa akibat beban konteks yang berlebih. GPT-5.4-mini mengungguli Llama 3.1:8B pada seluruh metrik, dengan kesenjangan terbesar pada METEOR (0,4279 vs. 0,2835). Secara keseluruhan, kombinasi *zero-shot persona* GPT-5.4-mini menjadi konfigurasi terbaik dengan BERTScore 0,7847, BLEU 0,1786, dan ROUGE-L 0,3675.

Kata Kunci: *Abstractive Summarization, Prompt Engineering, Large Language Model.*