

ABSTRAK

Di era digital, pesatnya penyebaran informasi daring menghadirkan tantangan besar dalam mengidentifikasi konten yang relevan, khususnya dalam ranah berita *online*. Rendahnya budaya membaca di Indonesia serta preferensi terhadap konten instan semakin memperburuk permasalahan ini. Peringkasan teks otomatis menawarkan solusi yang menjanjikan dengan cara merangkum teks berita panjang menjadi ringkasan yang padat dan informatif. Studi ini membandingkan efektivitas dua algoritma peringkasan: BART, model abstraktif berbasis *deep learning*, dan *TextRank*, metode ekstraktif berbasis graf. Evaluasi dilakukan secara kuantitatif dan kualitatif. Untuk penilaian kuantitatif digunakan metrik ROUGE. BART secara konsisten mengungguli *TextRank* pada semua varian ROUGE—dengan rata-rata skor F1 yang lebih tinggi (ROUGE-1: 0.42 vs. 0.31, ROUGE-2: 0.23 vs. 0.12, ROUGE-L: 0.31 vs. 0.22). Survei pengguna yang melibatkan 21 partisipan juga mendukung keunggulan BART, dengan penilaian rata-rata yang lebih tinggi dalam hal kebermanaknaan dan kejelasan informasi. Namun, analisis waktu eksekusi menunjukkan adanya *trade-off* antara kualitas dan efisiensi. *TextRank* menyelesaikan peringkasan hanya dalam waktu 0,01–0,02 detik, sedangkan BART membutuhkan 18 sampai 20 detik, mencerminkan kebutuhan komputasi yang tinggi.

Kata kunci: Peringkasan Teks, Berita Daring, BART, *TextRank*, ROUGE

ABSTRACT

In the digital era, the rapid spread of online information presents significant challenges in identifying relevant content, particularly in the realm of online news. Indonesia's low reading culture and preference for instant content further exacerbate this issue. Automatic text summarization offers a promising solution by condensing lengthy articles into concise and informative summaries. This study compares the effectiveness of two summarization algorithms: BART, an abstractive model based on deep learning, and TextRank, a graph-based extractive method. Evaluations were conducted both quantitatively and qualitatively. The ROUGE metric was used for quantitative assessment. BART consistently outperformed TextRank across all ROUGE variants—with higher average F1-scores (ROUGE-1: 0.42 vs. 0.31, ROUGE-2: 0.23 vs. 0.12, ROUGE-L: 0.31 vs. 0.22). A user survey involving 21 participants also supported BART's superiority, with higher average ratings in terms of informativeness and clarity. However, execution time analysis revealed a trade-off between quality and efficiency. TextRank completed summarization in just 0.01–0.02 seconds, while BART required 18 to 20 seconds, reflecting its high computational demands.

Keywords: *Text Summarization, Online News, BART, TextRank, ROUGE*