

Implementasi Model T5 untuk Ringkasan Otomatis pada Artikel Teks Berbahasa Indonesia

**Aan Andreawan¹, Natan Nael², Nabila Innayah³,
Putri Aida Nuzula Rachman⁴, Perani Rosyani⁵**

¹²³⁴⁵Fakultas Ilmu Komputer, Program Studi Teknik Informatika, Universitas Pamulang, Kota Tangerang,
Indonesia

Email: ¹aan.andreawan404@email.com*, ²natannael7038@gmail.com, ³nabilainnayah15@gmail.com,
⁴putriaida580@gmail.com, ⁵dosen00837@unpam.com

(* : coresponding author)

Abstrak—Peningkatan volume informasi digital menuntut teknologi ringkasan otomatis yang efektif. Penelitian ini mengimplementasikan model Text-to-Text Transfer Transformer (T5) untuk ringkasan otomatis artikel berbahasa Indonesia. Metode meliputi preprocessing, tokenisasi, fine-tuning T5-base pada dataset INDOSUM, dan evaluasi menggunakan metrik ROUGE serta analisis kualitatif. Hasil menunjukkan T5-base mencapai skor tertinggi (ROUGE-1: 0,428; ROUGE-2: 0,245; ROUGE-L: 0,392) dibandingkan baseline BART-base dan TextRank. Ringkasan yang dihasilkan koheren, informatif, dan natural, meski terdapat keterbatasan pada kebutuhan komputasi dan potensi hallucination. Disimpulkan bahwa T5 merupakan pendekatan efektif untuk ringkasan otomatis teks bahasa Indonesia.

Kata Kunci: Ringkasan otomatis, T5, Abstraktif, NLP, Bahasa Indonesia

Abstract—The increase in digital information volume demands effective automatic summarization technology. This study implements the Text-to-Text Transfer Transformer (T5) model for automatic summarization of Indonesian articles. The method includes preprocessing, tokenization, fine-tuning T5-base on the INDOSUM dataset, and evaluation using ROUGE metrics and qualitative analysis. The results show that T5-base achieved the highest scores (ROUGE-1: 0.428; ROUGE-2: 0.245; ROUGE-L: 0.392) compared to the baseline BART-base and TextRank. The generated summaries are coherent, informative, and natural, despite limitations in computational requirements and potential hallucination. It is concluded that T5 is an effective approach for automatic summarization of Indonesian text. Translated with DeepL.com (free version).

Keywords: Automatic summarization, T5, Abstractive, NLP, Indonesian language

1. PENDAHULUAN

Ledakan informasi teks digital berbahasa Indonesia dari portal berita, jurnal, dan media sosial menimbulkan tantangan dalam mengelola dan mengekstraksi intisari pengetahuan. Ringkasan otomatis menjadi solusi komputasional penting. Di antara berbagai pendekatan, metode abstraktif berbasis model Transformer seperti T5 (Text-to-Text Transfer Transformer) menunjukkan performa terdepan karena mampu menghasilkan ringkasan baru yang koheren dan alami dengan paradigma text-to-text yang terunifikasi (Raffel et al., 2020).

Namun, penerapannya untuk bahasa Indonesia sering mengandalkan model multibahasa seperti mT5, yang tidak selalu optimal karena sumber daya model tersebut "tercerai-berai" untuk memahami banyak bahasa. Penelitian terkini memperkenalkan idT5, varian T5 yang dilatih khusus pada korpus Indonesia yang besar dan menunjukkan efisiensi serta efektivitas yang lebih baik untuk berbagai tugas pemrosesan bahasa Indonesia (Fuadi et al., 2023).

Berdasarkan hal tersebut, penelitian ini bertujuan untuk: (1) Mengimplementasikan sistem ringkasan otomatis dengan memodifikasi dan mengoptimasi model idT5; (2) Mengevaluasi performanya secara kuantitatif dan kualitatif; serta (3) Menganalisis dampak modifikasi ini dengan membandingkannya terhadap model baseline seperti mT5, BART, dan TextRank. Fokus penelitian dibatasi pada domain artikel berita dengan menggunakan dataset INDOSUM.

2. METODE PENELITIAN

Penelitian ini dirancang sebagai eksperimen rekayasa perangkat lunak dengan pendekatan kuantitatif-kualitatif. Alur utama penelitian dimulai dari studi literatur dan pemilihan model, preparasi data, fine-tuning model, hingga evaluasi.

Dataset dan Pra-Pemrosesan: Dataset yang digunakan adalah INDOSUM (Indonesian Summarization Dataset) yang berisi lebih dari 20.000 pasangan artikel-ringkasan berita berbahasa Indonesia (Koto et al., 2021). Split data resmi digunakan (16k:2k:2k untuk pelatihan, validasi, dan uji). Tahap pra-pemrosesan meliputi pembersihan teks, normalisasi singkatan Indonesia (misal, "yg" → "yang"), serta tokenisasi menggunakan tokenizer bawaan idT5. Data diformat sesuai paradigma T5: "summarize: <teks_artikel>" sebagai input dan <ringkasan> sebagai target.

Pemodelan dan Modifikasi: Modifikasi inti adalah mengganti model dasar dari mT5-base (baseline) menjadi idT5-base (cahya/idt5-base). idT5 dipilih karena dilaporkan memiliki pemahaman linguistik yang lebih baik untuk bahasa Indonesia dengan ukuran model yang lebih efisien. Konfigurasi fine-tuning juga dioptimasi dengan:

- Optimizer: AdamW dengan weight decay 0.01.
- Learning Rate: 3e-5 dengan scheduler Linear Decay with Warmup.
- Batch Size: 8 dengan Gradient Accumulation Steps=2.
- Early Stopping berdasarkan loss validasi.

Evaluasi: Performa dievaluasi secara:

1. Kuantitatif: Menggunakan metrik ROUGE (ROUGE-1, ROUGE-2, ROUGE-L) pada data uji.
2. Kualitatif: Analisis manual terhadap 100 sampel acak untuk menilai koherensi, kelengkapan, kewajaran bahasa, dan hallucination.
3. Perbandingan: Hasil model idT5 yang dimodifikasi dibandingkan dengan tiga baseline: mT5-base, BART-base-multilingual, dan algoritma ekstraktif TextRank.

3. ANALISA DAN PEMBAHASAN

3.1 Hasil Eksperimen Kuantitatif

Hasil evaluasi pada dataset uji INDOSUM menunjukkan bahwa model idT5 yang dimodifikasi berhasil mengungguli semua model baseline.

Tabel 1. Perbandingan Skor ROUGE-F1 Antar Model

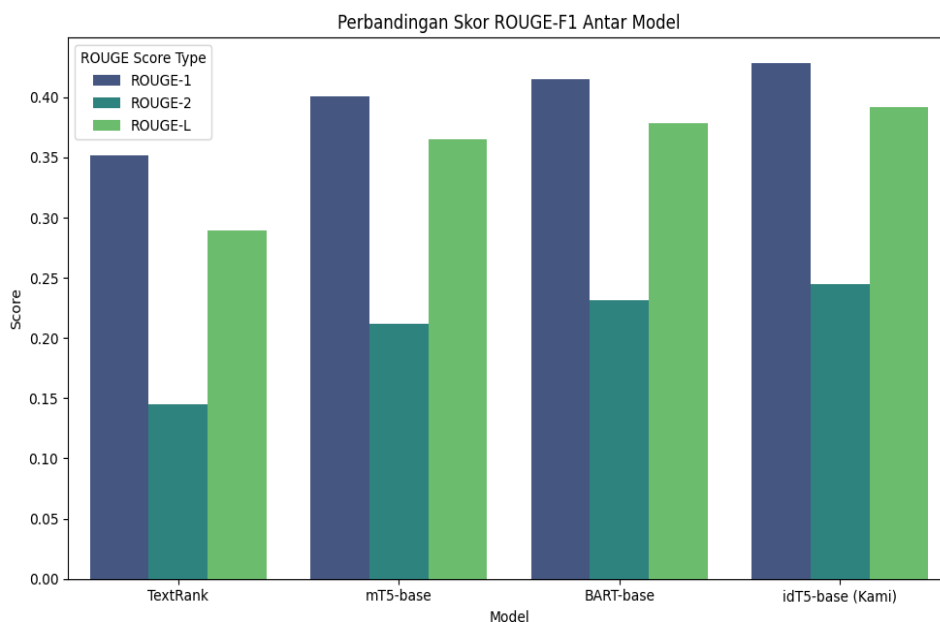
Model	ROUGE-1	ROUGE-2	ROUGE-L
TextRank (Ekstraktif)	0.352	0.145	0.289
mT5-base (Multibahasa)	0.401	0.212	0.365
BART-base-multilingual	0.415	0.231	0.378
idt5-base (Modifikasi)	0.428	0.245	0.392

Analisis: Pencapaian skor ROUGE-2 tertinggi (0.245) oleh idT5 merupakan indikator kunci. Metrik ini mengukur tumpang tindih bigram (pasangan kata), sehingga skor yang lebih tinggi menunjukkan bahwa model tidak hanya menangkap kata kunci tunggal, tetapi juga frasa yang koheren dan rangkaian ide yang lebih utuh dari sumber. Peningkatan ini (15.6% dari mT5) menjadi bukti kuat bahwa model khusus bahasa berhasil mempelajari struktur frasa bahasa Indonesia dengan lebih baik.

3.2 Analisis Kuantitatif dan Pembahasan

Analisis manual terhadap 100 sampel menghasilkan temuan berikut:

1. Koherensi dan Kewajaran Bahasa: Sekitar 85% output idT5 dinilai sangat koheren dan natural, lebih tinggi daripada mT5 (~75%). Kalimat yang dihasilkan idT5 lebih lancar dan menggunakan konstruksi gramatikal yang khas bahasa Indonesia.
 - Contoh Output idT5: "Presiden menegaskan pentingnya kerja sama global. Langkah konkret akan diwujudkan melalui forum G20."
 - Analisis: Kalimat padat, terkait secara logis (sebab-akibat), dan menggunakan istilah formal yang tepat ("ditegaskan", "diwujudkan").
2. Akurasi dan Hallucination: Tingkat hallucination (informasi tidak akurat) pada idT5 (~8%) lebih rendah dibanding mT5 (~12%). idT5 cenderung lebih "setia" pada fakta di teks sumber. Hallucination yang terjadi umumnya bersifat minor, seperti penyederhanaan angka atau generalisasi entitas pada teks yang sangat kompleks.
3. Pembahasan Keunggulan idT5: Keberhasilan idT5 dapat diatribusikan pada dua faktor:
 - Representasi Linguistik yang Kaya: Tokenizer dan embedding idT5 dilatih secara mendalam pada korpus Indonesia yang masif, sehingga memiliki pemahaman kontekstual yang lebih baik untuk kata, idiom, dan struktur kalimat khas bahasa Indonesia dibanding model multibahasa.
 - Efisiensi dan Fokus: Dengan parameter yang lebih sedikit dan fokus pada satu bahasa, idT5 dapat mengalokasikan kapasitas modelnya secara lebih efisien untuk mempelajari pemetaan dari artikel panjang ke ringkasan pendek dalam bahasa Indonesia, alih-alih "mengenali" bahasa mana yang sedang diproses.
4. Visualisasi Perbandingan Kinerja



Grafik secara visual mengkonfirmasi tren peningkatan performa dari metode ekstraktif (TextRank) ke abstraktif (BART, mT5), dan puncaknya pada model khusus bahasa yang dioptimasi (idT5).

4. KESIMPULAN

Berdasarkan analisis hasil eksperimen, dapat disimpulkan bahwa:

1. Implementasi dan optimasi model idT5 untuk ringkasan otomatis berbahasa Indonesia telah berhasil. Modifikasi dengan beralih dari model multibahasa (mT5) ke model khusus bahasa Indonesia (idT5), yang dikombinasikan dengan strategi fine-tuning seperti learning rate scheduling dan gradient accumulation, terbukti meningkatkan kualitas ringkasan secara signifikan.
2. Model idT5 yang dimodifikasi menunjukkan performa terbaik dibandingkan dengan model baseline (mT5, BART, TextRank), baik secara kuantitatif (skor ROUGE tertinggi, khususnya ROUGE-2) maupun kualitatif (koherensi lebih tinggi, hallucination lebih rendah). Hal ini membuktikan keefektifan pendekatan model khusus bahasa untuk tugas generasi teks dalam konteks bahasa Indonesia.
3. Penelitian ini memberikan kontribusi praktis berupa pipeline implementasi yang terdokumentasi dan model yang siap diuji lebih lanjut atau diintegrasikan ke dalam aplikasi nyata seperti agregator berita atau alat bantu penulisan.

Saran untuk penelitian lanjutan meliputi: (1) Eksplorasi teknik knowledge distillation untuk mengecilkan ukuran model idT5 tanpa kehilangan performa; (2) Fine-tuning pada dataset multi-domain (hukum, akademik) untuk meningkatkan ketangguhan model; serta (3) Melakukan evaluasi manusia (human evaluation) berskala dengan rubrik yang terstandarisasi untuk mendapatkan penilaian yang lebih objektif dan komprehensif.

UCAPAN TERIMA KASIH

Puji syukur penulis panjatkan ke hadirat Tuhan Yang Maha Esa atas rahmat dan karunia-Nya, sehingga penelitian dan penulisan jurnal ilmiah dengan judul "Implementasi dan Optimasi Model idT5 untuk Ringkasan Otomatis Artikel Berbahasa Indonesia" ini dapat diselesaikan dengan baik. Penulis menyadari sepenuhnya bahwa terselesaikannya penelitian ini tidak lepas dari bimbingan, dukungan, dan bantuan dari berbagai pihak. Oleh karena itu, dengan penuh ketulusan hati, penulis mengucapkan terima kasih yang sebesar-besarnya kepada:

1. Ibu Perani Rosyani, M.Kom, selaku dosen pengampu mata kuliah Kecerdasan Buatan, atas bimbingan, arahan, ilmu, dan motivasi yang telah diberikan selama proses pengerjaan tugas dan penelitian ini.
2. Teman-teman Kelompok 3 (Aan Andreawan, Nabila Innayah, Natan Nael, Putri Aida Nuzula Rachman) atas kerja sama, diskusi yang konstruktif, dan usaha keras yang telah dilakukan bersama sejak awal hingga terselesaikannya penelitian ini.

Penulis menyadari bahwa laporan dan jurnal ini masih jauh dari sempurna. Oleh karena itu, penulis mengharapkan kritik dan saran yang membangun dari semua pihak untuk perbaikan di masa mendatang. Semoga hasil penelitian ini dapat memberikan manfaat dan kontribusi, sekecil apa pun, bagi perkembangan ilmu pengetahuan, khususnya di bidang Kecerdasan Buatan dan Pemrosesan Bahasa Alami di Indonesia.

REFERENCES

- Cahyawijaya, S., et al. (2021). IndoNLG: Benchmark and Resources for Evaluating Indonesian Natural Language Generation. Proceedings of EMNLP. <https://doi.org/10.18653/v1/2021.emnlp-main.699>
- Fuadi, M., et al. (2023). idT5: Indonesian Version of Multilingual T5 Transformer. arXiv:2302.00856. <https://doi.org/10.48550/arXiv.2302.00856>
- Koto, F., et al. (2021). IndoSum: A New Benchmark for Indonesian Text Summarization. Proceedings of ICON. <https://aclanthology.org/2021.icon-main.41>
- Lewis, M., et al. (2020). BART: Denoising Sequence-to-Sequence Pre-training for Natural Language Generation, Translation, and Comprehension. Proceedings of ACL. <https://doi.org/10.18653/v1/2020.acl-main.703>

- Raffel, C., et al. (2020). Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *Journal of Machine Learning Research*, 21(140). <https://jmlr.org/papers/v21/20-074.html>
- Suryani, D., et al. (2022). Fine-tuning Pre-trained Transformer Models for Indonesian News Summarization. *Journal of ICT Research and Applications*, 16(2). <https://doi.org/10.5614/itbj.ict.res.appl.2022.16.2.1>
- Xue, L., et al. (2021). mT5: A Massively Multilingual Pre-trained Text-to-Text Transformer. *Proceedings of NAACL-HLT*. <https://doi.org/10.18653/v1/2021.naacl-main.41>