

Jurnal JTIK (Jurnal Teknologi Informasi dan Komunikasi)

DOI: <https://doi.org/10.35870/jtik.v10i1.5247>

Optimalisasi Peringkasan Artikel Teks Bahasa Indonesia dengan Kombinasi *TextRank* dan *Graph Neural Network* Sederhana

Muhammad Rifqi Syatria¹, Dadang Iskandar Mulyana^{2*}

^{1,2*} Program Studi Teknik Informatika, Sekolah Tinggi Ilmu Komputer Cipta Karya Informatika, Kota Jakarta Timur, Daerah Khusus Ibukota Jakarta, Indonesia.

article info

Article history:

Received 4 August 2025

Received in revised form

20 August 2025

Accepted 1 November 2025

Available online January 2026.

Keywords:

Automatic Text

Summarization; TextRank;

Graph Neural Network;

SimpleGNN, ROUGE;

Indonesian; Liputan6; Natural

Language Processing.

Kata Kunci:

Peringkasan Teks Otomatis;

TextRank; Graph Neural

Network; SimpleGNN,

ROUGE; Bahasa Indonesia;

Liputan6; Pemrosesan Bahasa

Alami.

abstract

The delivery of digital information in Indonesian-language news presents challenges in efficiently capturing the core information. This study proposes a combination of the TextRank algorithm and a simple Graph Neural Network (GNN) to improve the quality of automatic text summarization. TextRank is used to construct a sentence graph based on TF-IDF similarity and cosine similarity, followed by training a SimpleGNN model to optimize sentence scores. Evaluations were conducted on 1,000 articles from the Liputan6 dataset using the ROUGE metric (ROUGE-1, ROUGE-2, and ROUGE-L). The results show that this combined method improves performance compared to pure TextRank, especially in capturing semantic relationships between sentences. This study demonstrates that the integration of a simple GNN can enrich representations in graphs and provide more informative and contextual summaries.

abstrak

Ledakan informasi digital dalam berita berbahasa Indonesia menimbulkan tantangan dalam menyerap inti informasi secara efisien. Penelitian ini mengusulkan metode kombinasi antara algoritma TextRank dan Graph Neural Network (GNN) sederhana untuk meningkatkan kualitas peringkasan teks otomatis. TextRank digunakan untuk membentuk graf kalimat berdasarkan kemiripan TF-IDF dan cosine similarity, kemudian dilanjutkan dengan pelatihan model SimpleGNN untuk mengoptimalkan skor kalimat. Evaluasi dilakukan terhadap 1000 artikel dari dataset Liputan6 menggunakan metrik ROUGE (ROUGE-1, ROUGE-2, dan ROUGE-L). Hasil menunjukkan bahwa metode gabungan ini meningkatkan performa peringkasan dibandingkan TextRank murni, khususnya dalam menangkap relasi semantik antar kalimat. Penelitian ini membuktikan bahwa integrasi GNN sederhana dapat memperkaya representasi dalam graf dan memberikan ringkasan yang lebih informatif serta kontekstual.

Corresponding Author. Email: mahvin2012@gmail.com^{2}.

1. Pendahuluan

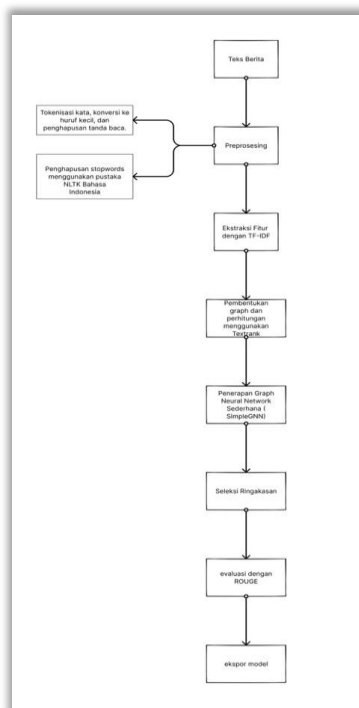
Di era digital saat ini, manusia dihadapkan pada ledakan informasi yang terus bertambah setiap detik. Jumlah dokumen digital seperti berita daring, artikel ilmiah, laporan bisnis, dan konten media sosial meningkat secara signifikan, sehingga memunculkan fenomena *information overload* (Gunawan *et al.*, 2016). Dalam situasi tersebut, pengguna sering kesulitan untuk menemukan dan memahami informasi yang benar-benar relevan secara efisien. Salah satu pendekatan yang banyak dikembangkan untuk mengatasi tantangan tersebut adalah peringkasan teks (*text summarization*) (Gunawan *et al.*, 2016). Peringkasan merupakan proses mereduksi sejumlah besar informasi menjadi ringkasan singkat dan koheren, namun tetap mencakup poin-poin penting serta gagasan utama (Hossain *et al.*, 2023). Peringkasan teks otomatis bertujuan menghasilkan ringkasan yang memuat kalimat-kalimat esensial dan mencakup informasi relevan dari dokumen asal (Widyassari *et al.*, 2022; Khor *et al.*, 2021; Christanti & Pragantha, 2017; Ogundele *et al.*, 2023). Beragam metode telah dikembangkan untuk mengekstraksi kata kunci atau kalimat penting secara otomatis, tetapi *TextRank* masih menjadi salah satu metode dasar yang banyak diadopsi (Zhang *et al.*, 2020). Meskipun kerap digunakan karena tak berlabel (*unsupervised*) (Kazemi *et al.*, 2020), *TextRank* memiliki beberapa keterbatasan: ketergantungan pada parameter yang ditetapkan secara manual misalnya ukuran jendela koeksistensi kata dan jumlah iterasi yang dapat memengaruhi akurasi ringkasan, serta kelemahan dalam menangkap makna semantik antarkalimat (Mallick *et al.*, 2018). Sejalan dengan itu, Zhang *et al.* (2020) menegaskan bahwa performa *TextRank* sangat bergantung pada penyetelan parameter agar tetap optimal dan stabil terhadap panjang teks. Chen *et al.* (2015) menunjukkan bahwa peringkasan *broadcast news* memerlukan model yang mampu memahami struktur informasi jangka panjang. Pendekatan berbasis *bag-of-words* dinilai terlalu sederhana dan tidak memadai untuk memodelkan hubungan semantik intrakalimat. *Graph Neural Network* (GNN) dimanfaatkan karena memiliki daya ungkap tinggi terhadap relasi kompleks pada struktur teks (Gao & Huang, 2025).

Oleh sebab itu, pendekatan jaringan saraf seperti RNNLM atau GNN menawarkan keunggulan dengan mengintegrasikan penggunaan kata dan keterhubungan semantik antarkalimat. Pada saat yang sama, Koto, Lau, dan Baldwin (2020) mencatat bahwa penelitian peringkasan otomatis berbahasa Indonesia masih relatif sedikit, baik dari segi jumlah studi maupun pemanfaatan korpus berskala besar. Selain itu, Pan, Li, dan Dai (2019) menunjukkan bahwa varian *TextRank* yang disempurnakan bersifat lebih umum dan dapat meningkatkan akurasi ekstraksi kata kunci dibanding versi standar. Hal tersebut mengukuhkan *TextRank* sebagai algoritma dasar yang tetap relevan dalam pengembangan sistem peringkasan berbasis graf. Berdasarkan latar belakang tersebut, penelitian ini berupaya mengembangkan dan mengevaluasi sistem peringkasan berita berbahasa Indonesia secara otomatis dengan membandingkan kinerja *TextRank* klasik dan *TextRank* yang diperkuat *Graph Neural Network* (GNN) sederhana. Evaluasi dilakukan menggunakan metrik ROUGE (ROUGE-1, ROUGE-2, dan ROUGE-L) untuk menilai kualitas ringkasan secara objektif.

2. Metodologi Penelitian

Penelitian ini menggunakan pendekatan komparatif untuk menilai perbedaan kinerja antara *TextRank* klasik dan *TextRank* yang dikombinasikan dengan *Graph Neural Network* (GNN) sederhana dalam menghasilkan ringkasan teks otomatis pada korpus berita *Liputan6*. Proses penelitian dimulai dengan tahap praproses teks, yaitu membersihkan setiap artikel dari karakter yang tidak diperlukan, melakukan pemisahan teks menjadi kalimat, tokenisasi kata, serta penghapusan *stopwords* Bahasa Indonesia. Tahap ini bertujuan untuk memperoleh representasi kalimat yang bersih dan siap diolah. Setelah proses praproses, setiap kalimat diubah menjadi vektor numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), yang memungkinkan penilaian bobot kata berdasarkan frekuensi kemunculannya dalam artikel dan keseluruhan korpus. Representasi TF-IDF tersebut kemudian digunakan untuk menghitung kemiripan antar kalimat melalui *cosine similarity*, yang hasilnya dimanfaatkan untuk membentuk graf tak berarah berbobot, di mana simpul graf merepresentasikan kalimat dan bobot sisi

menunjukkan tingkat kemiripan antar kalimat. Graf yang terbentuk selanjutnya diproses menggunakan algoritma TextRank untuk menghasilkan skor kepentingan bagi setiap kalimat. Skor ini kemudian digunakan sebagai nilai awal dalam tahap pelatihan model SimpleGNN. Model SimpleGNN menerima masukan berupa matriks fitur TF-IDF dan matriks kedekatan antarkalimat, lalu mempelajari pola hubungan antar simpul untuk memperbaiki pemeringkatan kalimat secara lebih adaptif. Pelatihan model dilakukan menggunakan fungsi kerugian *Mean Squared Error* (MSE) dan algoritma optimisasi Adam selama sepuluh *epoch*. Setelah model selesai dilatih, proses pemilihan ringkasan dilakukan dengan mengambil lima kalimat dengan skor tertinggi dari masing-masing pendekatan, yaitu TextRank klasik dan TextRank + SimpleGNN. Untuk menilai kualitas ringkasan yang dihasilkan, penelitian ini menggunakan metrik ROUGE (ROUGE-1, ROUGE-2, dan ROUGE-L) yang membandingkan hasil peringkasan model dengan ringkasan acuan (*gold summary*). Evaluasi ini memungkinkan penilaian objektif terhadap kesesuaian dan ketepatan ringkasan yang dihasilkan oleh kedua metode, sehingga perbedaan kinerja dapat dianalisis secara terukur.



Gambar 1. Flowchart

Praproses Teks

Pada tahap praproses teks, setiap artikel berita terlebih dahulu melalui proses pembersihan untuk menghilangkan karakter atau tanda baca yang tidak diperlukan. Prosedur ini mengikuti pendekatan yang dikemukakan oleh Algoritma *et al.* (2025), di mana proses pembersihan dilakukan untuk mengurangi *noise* pada data teks sehingga struktur kalimat menjadi lebih teratur. Setelah proses pembersihan selesai, teks dipecah menjadi satuan kalimat menggunakan metode *sentence tokenization* dari pustaka *Natural Language Toolkit* (NLTK) sebagaimana dijelaskan oleh Novie Tri Lestari *et al.* (2022). Setiap artikel direpresentasikan sebagai kumpulan kalimat dalam bentuk:

$$D = \{S_1, S_2, \dots, S_n\}$$

Keterangan:

D = Artikel berita yang terdiri atas Kumpulan kalimat

S_i = Artikel Ke-i dalam berita

Ekstraksi Fitur Kalimat dengan TF-IDF

Tahap berikutnya setelah praproses teks adalah ekstraksi fitur kalimat menggunakan metode *Term Frequency–Inverse Document Frequency* (TF-IDF). Setiap kalimat dalam artikel diubah menjadi vektor numerik yang menggambarkan bobot atau tingkat kepentingan setiap kata dalam konteks dokumen. Proses ini dilakukan dengan memanfaatkan pustaka *scikit-learn* sebagai alat bantu dalam perhitungan bobot. Representasi TF-IDF berfungsi sebagai masukan utama untuk pembentukan graf serta sebagai fitur input ke dalam model Graph Neural Network (GNN). Sebelum digunakan dalam proses pembentukan graf, nilai hasil perhitungan TF-IDF diukur tingkat kemiripannya menggunakan *Cosine Similarity*. Pendekatan ini berperan untuk mengidentifikasi seberapa dekat dua kalimat berdasarkan arah vektor dalam ruang multidimensi. Menurut Ogundele *et al.* (2023), *cosine similarity* merupakan ukuran kemiripan antara dua vektor non-nol yang dihitung berdasarkan nilai kosinus dari sudut di antara keduanya. Semakin kecil sudut antar vektor, semakin tinggi nilai kemiripannya, yang menandakan bahwa kedua kalimat memiliki kesamaan konteks semantik. Secara umum, algoritma TF-IDF terdiri atas dua komponen utama, yaitu *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF).

Term Frequency digunakan untuk mengukur seberapa sering suatu kata muncul dalam satu dokumen dibandingkan dengan total jumlah kata dalam dokumen tersebut. Nilai TF dihitung menggunakan rumus (Zhang *et al.*, 2020):

$$TF(t) = \frac{\text{Jumlah kemunculan kata } t \text{ dalam dokumen}}{\text{Total jumlah kata dalam dokumen}}$$

Sementara itu, *Inverse Document Frequency* berfungsi menilai tingkat keunikan atau kelangkaan suatu kata di seluruh kumpulan dokumen. Konsep ini menekankan bahwa kata yang jarang muncul dalam korpus memiliki nilai IDF lebih tinggi, karena dianggap lebih spesifik terhadap dokumen tertentu. Rumus matematis IDF dinyatakan sebagai berikut (Zhang *et al.*, 2020):

$$IDF(t) = \log \frac{\text{Jumlah seluruh dokumen}}{\text{Jumlah dokumen yang mengandung kata } t}$$

Bobot akhir TF-IDF diperoleh dengan mengalikan kedua nilai tersebut. Secara matematis dapat ditulis sebagai:

$$TFIDF(t) = TF_t \times IDF_t$$

Menurut Widyassari *et al.* (2022), nilai TF-IDF yang tinggi menunjukkan bahwa suatu kata sering muncul dalam dokumen tertentu, tetapi jarang ditemukan pada dokumen lain, sehingga dianggap penting dan representatif bagi isi dokumen tersebut. Bobot inilah yang kemudian digunakan untuk menggambarkan karakteristik semantik setiap kalimat sebelum dimasukkan ke tahap pembentukan graf dan analisis lanjutan menggunakan GNN.

Pembentukan Graf dan Perhitungan TextRank

Tahap selanjutnya adalah pembentukan graf dan perhitungan skor kalimat menggunakan algoritma *TextRank*. Setelah diperoleh representasi vektor dari setiap kalimat melalui metode TF-IDF, dilakukan pengukuran cosine similarity untuk menentukan tingkat kedekatan antar kalimat. Menurut Ogundele *et al.* (2023), *cosine similarity* digunakan untuk menghitung kesamaan dua vektor non-nol dengan mengukur nilai kosinus dari sudut di antara keduanya. Semakin kecil sudut yang terbentuk, semakin besar

nilai kemiripan antara dua vektor tersebut. Hubungan matematis ini dapat dinyatakan sebagai:

$$\text{similarity}(A, B) = \frac{\sum_{i=1}^n A_i B_i}{\sqrt{\sum_{i=1}^n A_i^2} \sqrt{\sum_{i=1}^n B_i^2}}$$

Vektor A dan B merepresentasikan dua kalimat yang dibandingkan, sementara A_i dan B_i adalah nilai elemen ke- i dari masing-masing vektor. Hasil perhitungan kemiripan ini digunakan untuk membentuk graf tak berarah berbobot, di mana setiap simpul (node) merepresentasikan sebuah kalimat, dan setiap sisi (edge) menggambarkan tingkat kemiripan antar kalimat berdasarkan nilai cosine similarity. Semakin besar nilai kemiripan, semakin kuat bobot sisi yang menghubungkan kedua simpul tersebut. Graf yang telah dibentuk kemudian diproses menggunakan algoritma *TextRank*, yaitu adaptasi dari algoritma *PageRank* yang dikembangkan oleh Brin dan Page untuk sistem peringkat laman web. Pada penelitian ini, algoritma tersebut digunakan untuk menghitung skor kepentingan setiap kalimat dalam graf. Rumus perhitungan *TextRank* untuk simpul V_i dapat dituliskan sebagai berikut (Zhang *et al.*, 2020):

$$S(V_i) = (1 - d) + d \sum_{V_j \in \text{In}(V_i)} \frac{w_{ji}}{\sum_{V_k \in \text{Out}(V_j)} w_{jk}} S(V_j)$$

Keterangan dari persamaan di atas adalah sebagai berikut: $S(V_i)$ merupakan skor *TextRank* untuk simpul kalimat V_i , d adalah *damping factor* yang umumnya bernilai 0,85 untuk mengontrol penyebaran skor antar simpul, $\text{In}(V_i)$ menunjukkan himpunan simpul yang mengarah ke V_i , sedangkan w_{ji} merupakan bobot sisi antara simpul V_j dan V_i . Semakin tinggi bobot keterhubungan dan skor kalimat pengarah, semakin besar nilai skor *TextRank* yang diperoleh oleh suatu simpul. Proses ini memungkinkan sistem peringkat otomatis menyeleksi kalimat-kalimat yang memiliki hubungan kuat dengan kalimat lain di dalam teks, sekaligus mempertimbangkan keterkaitan semantik yang terbangun dalam struktur graf. Dengan demikian, hasil akhir yang diperoleh berupa serangkaian kalimat dengan nilai kepentingan tertinggi yang paling mewakili isi utama artikel. Pendekatan berbasis graf seperti ini terbukti efektif dalam mengidentifikasi kalimat inti tanpa memerlukan data pelatihan berlabel (Mallick *et al.*, 2018).

Penerapan *Graph Neural Network* (SimpleGNN)

Untuk memperkuat pemeringkatan kalimat, penelitian menerapkan Simple *Graph Neural Network* (SimpleGNN) dua lapis pada PyTorch. Masukan model terdiri atas matriks fitur TF-IDF setiap kalimat serta matriks kedekatan (*adjacency*) yang dibangun dari nilai *cosine similarity*. Strategi ini memungkinkan penyebaran informasi antarkalimat melalui struktur graf sehingga skor pentingnya kalimat tidak hanya ditentukan oleh bobot leksikal, tetapi juga oleh kedekatan relasional pada graf (Zhou *et al.*, 2020; Khaliq *et al.*, 2024). Pelatihan dilakukan dengan skema *weakly supervised*: skor TextRank digunakan sebagai target awal sehingga GNN bertindak sebagai pengalibrasi yang menyaring ulang peringkat berdasarkan pola keterhubungan yang lebih kaya (Yang *et al.*, 2023). Fungsi kerugian yang digunakan adalah *Mean Squared Error* (MSE) dengan Adam sebagai pengoptimal, dijalankan selama sepuluh *epoch*. Arsitektur maju (*forward pass*) mengikuti persamaan:

$$H^{(1)} = \text{ReLU}(A X W^{(1)} + b^{(1)}), \quad Y = H^{(1)} W^{(2)} + b^{(2)},$$

Dengan X sebagai matriks fitur TF-IDF, A sebagai matriks kedekatan, $w^{(1)}, w^{(2)}$ bobot lapisan, serta $b^{(1)}, b^{(2)}$ vektor *bias*. Keluaran Y diperlakukan sebagai skor akhir pemeringkatan kalimat. Pendekatan berbasis graf semacam ini terbukti efektif untuk tugas ekstraktif karena mampu menggabungkan informasi isi kalimat dan struktur hubungan antarkalimat tanpa memerlukan anotasi manual berskala besar (Zhou *et al.*, 2020).

Seleksi Ringkasan dan Evaluasi ROUGE

Setelah pelatihan, seluruh kalimat pada setiap artikel diranking menggunakan dua pendekatan TextRank klasik dan TextRank + SimpleGNN—lalu lima kalimat dengan skor tertinggi dipilih sebagai ringkasan. Mutu ringkasan dinilai menggunakan metrik ROUGE-1, ROUGE-2, dan ROUGE-L melalui pustaka *rouge_scorer*, yang masing-masing mengukur tumpang-tindih unigram, bigram, dan subsekuensi terpanjang antara ringkasan model dan *gold summary* (Widyassari *et al.*, 2022). Pemakaian gabungan ROUGE-1/2/L memberikan gambaran kuantitatif yang seimbang antara cakupan unit kata, kestabilan frasa pendek, dan keselarasan urutan

kalimat. Evaluasi diterapkan pada data uji yang sama untuk kedua pendekatan agar perbandingan bersifat adil dan dapat direplikasi.

Ekspor Model

Model GNN terlatih disimpan dalam format .pth dan diekspor ke ONNX (*Open Neural Network Exchange*) untuk memudahkan pemanfaatan lintas *framework*—misalnya integrasi ke TensorFlow atau CoreML—tanpa perubahan arsitektur. Praktik ini mempermudah penyebaran model pada berbagai lingkungan produksi dan memperkecil ketergantungan terhadap satu ekosistem perangkat lunak (Zhou *et al.*, 2020). Seluruh eksperimen menggunakan set kalimat yang identik ketika membandingkan TextRank dengan TextRank + SimpleGNN, sehingga hasil perbandingan kinerja tetap setara dari sisi input dan hanya berbeda pada mekanisme pemeringkatan.

3. Hasil dan Pembahasan

Hasil

Penerapan Model TextRank–SimpleGNN

TextRank merupakan algoritma pemeringkatan berbasis graf yang bersifat *unsupervised* dan telah banyak digunakan dalam sistem peringkasan teks otomatis karena kesederhanaannya serta kemampuannya menangkap hubungan antarkalimat melalui keterkaitan semantik atau *co-occurrence* (Zhang *et al.*, 2020). Meskipun demikian, algoritma ini memiliki keterbatasan utama, yaitu sifatnya yang statis dan hanya bergantung pada struktur graf tanpa memperhitungkan representasi fitur konten kalimat secara mendalam (Mallick *et al.*, 2018). Untuk mengatasi kelemahan tersebut, penelitian ini mengintegrasikan *Graph Neural Network* (GNN) sebagai pendekatan lanjutan guna memperkaya representasi informasi dalam simpul (kalimat) pada graf TextRank. Pendekatan ini memungkinkan penyebaran (*propagation*) informasi antarkalimat secara terstruktur melalui hubungan bobot yang tercermin dalam *adjacency matrix*. Penilaian terhadap suatu kalimat tidak hanya bergantung pada jumlah koneksi seperti dalam TextRank tradisional, tetapi juga pada kualitas hubungan semantik dan kekuatan konten antar kalimat (Yang *et al.*, 2023). Model GNN yang digunakan adalah SimpleGNN, sebuah arsitektur dua lapisan linier (*fully connected layers*) yang sederhana

namun efektif. Lapisan pertama mengubah vektor fitur masukan hasil TF-IDF menjadi representasi laten, sedangkan lapisan kedua menghasilkan skor akhir prediksi. Proses pelatihan dilakukan menggunakan pendekatan *weakly supervised learning*, di mana skor hasil TextRank digunakan sebagai target supervisi tanpa memerlukan label manual. Dengan demikian, model GNN belajar menyesuaikan bobot parameter berdasarkan hasil algoritma ranking tradisional yang sudah ada (Zhou *et al.*, 2020).

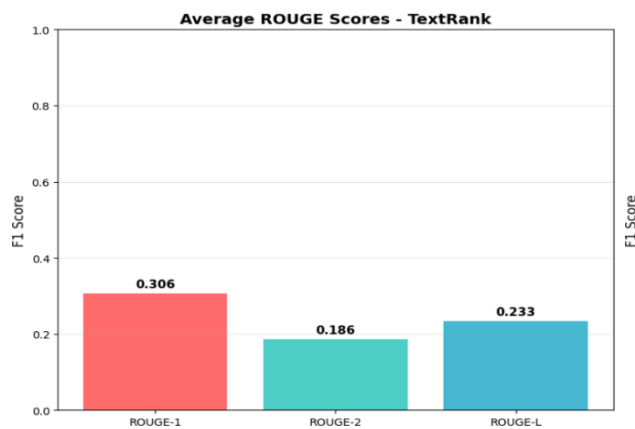
Dataset

Penelitian ini menggunakan dataset Liputan6, yakni kumpulan artikel berita berbahasa Indonesia yang pertama kali diperkenalkan oleh Koto, Lau, dan Baldwin (2020) dalam konferensi *AACL-IJCNLP 2020*. Dataset ini tersedia secara publik melalui repositori GitHub resmi penulis dan mencakup lebih dari 215.827 pasangan artikel dan ringkasan referensi. Dataset disusun dalam format JSON, yang berisi elemen-elemen seperti *judul artikel (title)*, *isi berita dalam bentuk daftar kalimat (clean_article)*, serta *ringkasan referensi (clean_summary)* yang berfungsi sebagai *gold summary* dalam tahap evaluasi. Dari keseluruhan dataset, penelitian ini menggunakan 1.000 artikel sebagai sampel penelitian. Pemilihan jumlah tersebut didasarkan pada pertimbangan efisiensi komputasi, manajemen memori, serta kecepatan pengembangan model. Pertama, pembatasan jumlah data membantu menghemat waktu pelatihan serta kebutuhan sumber daya (*computational efficiency*). Kedua, ukuran dataset yang lebih kecil mencegah *memory overflow* selama proses perhitungan matriks TF-IDF dan propagasi graf (*memory management*). Ketiga, jumlah data yang moderat mempermudah proses *debugging* dan iterasi cepat dalam tahap pengembangan model (*development speed*). Dataset Liputan6 dipilih karena representatif terhadap gaya bahasa jurnalistik Indonesia dan telah digunakan secara luas dalam berbagai penelitian peringkasan teks otomatis, baik ekstraktif maupun abstraktif (Koto *et al.*, 2020).

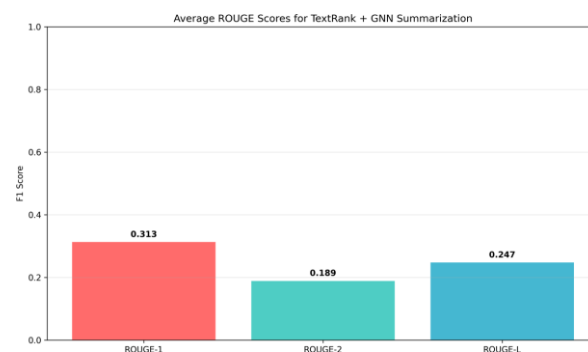
Penerapan Algoritma TextRank dengan TextRank-SimpleGNN

Evaluasi performa sistem peringkasan dilakukan dengan membandingkan hasil ringkasan yang dihasilkan oleh TextRank klasik dan TextRank yang diperkuat dengan SimpleGNN. Kedua model diuji

menggunakan 1.000 artikel berita dari dataset Liputan6, dengan metrik evaluasi utama ROUGE (Recall-Oriented Understudy for Gisting Evaluation). Metrik ini terdiri atas ROUGE-1, ROUGE-2, dan ROUGE-L, yang masing-masing mengukur tingkat kesamaan antara hasil ringkasan model dan ringkasan referensi berdasarkan *unigram overlap*, *bigram overlap*, serta *longest common subsequence* (Widyassari *et al.*, 2022). Hasil evaluasi menunjukkan bahwa integrasi SimpleGNN pada TextRank mampu meningkatkan kualitas hasil ringkasan, yang tercermin pada kenaikan nilai rata-rata ROUGE dibandingkan dengan metode TextRank klasik. Perbandingan skor rata-rata dari kedua metode divisualisasikan melalui diagram batang dan gambar berikut, yang menampilkan peningkatan kinerja model secara kuantitatif. Hasil ini mengindikasikan bahwa pendekatan berbasis GNN mampu memperbaiki keterbatasan 356emantic356l TextRank dengan menambahkan lapisan pemahaman 356emantic yang lebih dalam terhadap hubungan antarkalimat.



Gambar 2. Diagram Batang Performa TextRank



Gambar 3. Diagram Batang Performa TextRank + SimpleGNN

Tabel 1. Hasil Uji Perbandingan skor ROUGE

Jumlah Artikel	Model	ROUGE-1	ROUGE-2	ROUGE-L
1000	TextRank	0.3067	0.1857	0.2335
	TextRank + SimpleGNN	0.3127	0.1888	0.2474

Perbandingan Skor ROUGE dan Analisis Kinerja Model

Berdasarkan hasil pada diagram batang dan tabel perbandingan nilai rata-rata ROUGE, dapat diamati bahwa metode TextRank + SimpleGNN menunjukkan peningkatan performa pada seluruh metrik evaluasi dibandingkan dengan TextRank klasik. Secara kuantitatif, skor ROUGE-1 meningkat sebesar +0,006 poin, yang mengindikasikan bahwa model mampu menangkap lebih banyak *unigram* penting dari teks sumber. Peningkatan ini menandakan bahwa fitur semantik yang dipelajari melalui GNN berkontribusi terhadap penentuan kata-kata utama yang lebih representatif. Selanjutnya, nilai ROUGE-2 mengalami peningkatan sebesar +0,0031, meskipun relatif kecil, menunjukkan bahwa kombinasi metode mampu mempertahankan struktur frasa dua kata (*bigram*) yang mencerminkan keterpaduan kalimat lebih baik dibandingkan dengan TextRank murni. Sementara itu, skor ROUGE-L meningkat secara signifikan sebesar +0,0139, yang menandakan bahwa hasil ringkasan dari model GNN memiliki urutan kalimat yang lebih mendekati struktur ringkasan referensi (*gold summary*) dibandingkan hasil dari *TextRank* klasik.

Meskipun peningkatan yang diperoleh tidak terlalu besar secara absolut, hasil ini menunjukkan konsistensi arah peningkatan di semua metrik, yang berarti integrasi *Graph Neural Network* berhasil memperkaya kemampuan sistem dalam mengagregasi informasi antarkalimat secara lebih kontekstual. Hal ini selaras dengan tujuan utama penelitian, yakni mengembangkan model peringkasan teks yang tidak hanya mempertimbangkan hubungan struktural antar kalimat seperti pada pendekatan berbasis graf tradisional, tetapi juga memperhitungkan makna semantik dari isi kalimat melalui propagasi informasi dalam graf. GNN berfungsi sebagai lapisan representasi tambahan yang memperhalus hasil peringkat kalimat berdasarkan kesamaan konten dan

hubungan semantik yang lebih kompleks. Selain itu, penggunaan metrik ROUGE terbukti efektif sebagai ukuran kuantitatif dalam mengidentifikasi perbedaan kualitas antara dua model. Hasil pengukuran menunjukkan bahwa integrasi GNN berdampak positif terhadap peningkatan ketepatan dan kelengkapan ringkasan. Secara konseptual, hal ini memperlihatkan bahwa metode berbasis neural graph mampu memberikan pemahaman kontekstual yang lebih baik terhadap isi teks, dibandingkan algoritma graf konvensional yang hanya mengandalkan bobot koneksi antar simpul. Hasil ini juga menegaskan potensi scalability dan fleksibilitas pendekatan TextRank + SimpleGNN untuk diterapkan pada jenis teks lain di luar berita daring, seperti artikel ilmiah, teks opini, atau laporan kebijakan publik, termasuk pada dokumen multibahasa. Dengan kemampuan adaptif GNN terhadap variasi fitur linguistik dan hubungan semantik, pendekatan ini dapat menjadi fondasi bagi sistem peringkasan otomatis yang lebih cerdas dan kontekstual di masa mendatang.

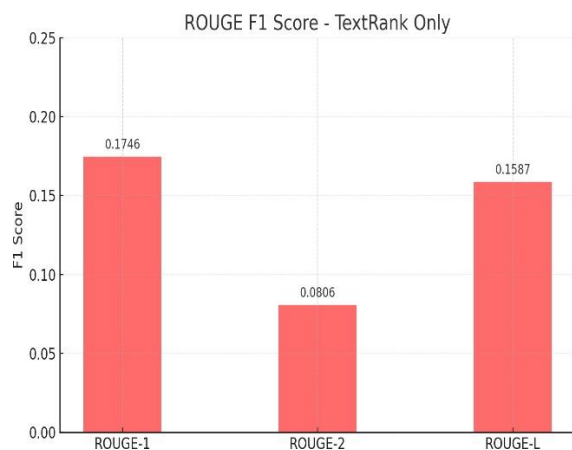
Sebagai bagian dari validasi lebih lanjut terhadap efektivitas metode, penelitian ini juga melakukan pengujian studi kasus individual terhadap satu artikel uji yang tidak termasuk dalam data pelatihan. Artikel tersebut diambil dari portal berita *Liputan6.com* dengan topik mengenai kebijakan pemerintah terhadap para konglomerat penandatangan *Master Settlement and Acquisition Agreement* (MSAA). Pengujian ini bertujuan untuk menilai kemampuan model dalam melakukan peringkasan pada data yang benar-benar baru, dengan meninjau hasil ringkasan otomatis terhadap ringkasan manual (*gold summary*) baik secara kuantitatif melalui skor ROUGE maupun secara kualitatif melalui analisis isi. Artikel uji tersebut memiliki ringkasan manual yang terdiri atas dua kalimat inti, sebagaimana disajikan pada tabel berikut, yang kemudian digunakan sebagai pembanding terhadap hasil ringkasan otomatis dari kedua metode yang diuji.

Tabel 2. Artikel Studi Kasus

Artikel Asli
Liputan6 com, Jakarta: Pemerintah masih memberikan waktu dua minggu lagi kepada seluruh konglomerat yang telah menandatangani perjanjian pengembalian bantuan likuiditas Bank Indonesia dengan jaminan aset (MSAA), untuk secepatnya menyerahkan jaminan pribadi serta aset. Jika lewat dari tenggat tersebut, pemerintah akan menerapkan tindakan hukum. Hal tersebut dikemukakan Menteri Koordinator Bidang Perekonomian Rizal Ramli di Jakarta, baru- baru ini. Rizal mengakui bahwa permintaan untuk meminta jaminan pribadi atau personal guarantee pada awalnya ditentang sejumlah konglomerat. Sebab para debitor menganggap tindakan tersebut memungkinkan pemerintah untuk menyita seluruh aset mereka baik yang berada di dalam maupun luar negeri. Sejauh ini, penilaian jaminan MSAA baru dilakukan atas aset milik Grup Salim. Tetapi, nilai aset yang dijaminan Kelompok Salim atas utang BLBI Bank Central Asia diperkirakan tak lebih dari Rp 20 triliun. Padahal, kewajiban mereka mencapai Rp 52 triliun. Sementara itu, pemerintah dengan DPR sepakat, hingga akhir Oktober mendatang, para konglomerat penandatangan MSAA harus sudah menutupi kekurangan mereka dengan menyerahkan aset baru. Selain itu, para pengutang tersebut diwajibkan memberikan jaminan pribadinya. (TNA/Merdi Sofansyah dan Anto Susanto).
Golden Summary
Pemerintah memberikan tenggat 14 hari kepada para konglomerat penandatangan MSAA untuk menyerahkan aset. Jika mangkir, mereka bakal dihukum.

Tabel 3. Ringkasan oleh TextRank

“com, Jakarta: Pemerintah masih memberikan waktu dua minggu lagi kepada seluruh konglomerat yang telah menandatangani perjanjian pengembalian bantuan likuiditas Bank Indonesia dengan jaminan aset (MSAA), untuk secepatnya menyerahkan jaminan pribadi serta aset. Sejauh ini, penilaian jaminan MSAA baru dilakukan atas aset milik Grup Salim. Sementara itu, pemerintah dengan DPR sepakat, hingga akhir Oktober mendatang, para konglomerat penandatangan MSAA harus sudah menutupi kekurangan mereka dengan menyerahkan aset baru. Sebab para debitor menganggap tindakan tersebut memungkinkan pemerintah untuk menyita seluruh aset mereka baik yang berada di dalam maupun luar negeri. Tetapi, nilai aset yang dijaminan Kelompok Salim atas utang BLBI Bank Central Asia diperkirakan tak lebih dari Rp 20 triliun.”

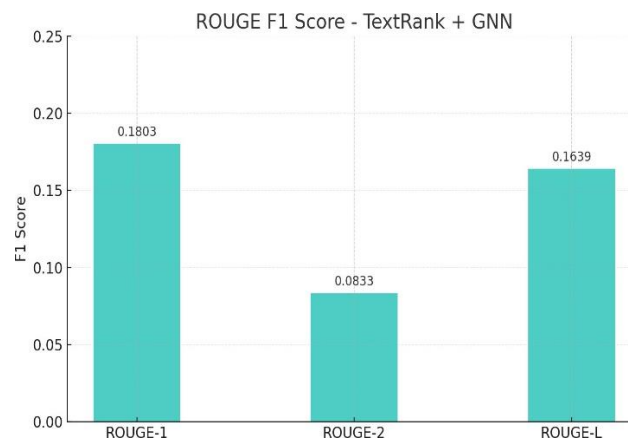


Gambar 4. Nilai ROUGE untuk TexRank pada Studi Kasus

Sementara itu, TextRank+GNN menambahkan dimensi relasional dan temporal ke dalam ringkasan, termasuk informasi kolaborasi antara pemerintah dan DPR, serta pendapat tokoh utama (Rizal). Kalimat-kalimat ini mencerminkan nuansa kebijakan dan respon publik yang lebih luas. Sebagaimana dalam tabel berikut:

Tabel 4. Ringkasan oleh Textrank-SimpleGNN

“com, Jakarta: Pemerintah masih memberikan waktu dua minggu lagi kepada seluruh konglomerat yang telah menandatangani perjanjian pengembalian bantuan likuiditas Bank Indonesia dengan jaminan aset (MSAA). untuk secepatnya menyerahkan jaminan pribadi serta aset. Sejauh ini, penilaian jaminan MSAA baru dilakukan atas aset milik Grup Salim. Sementara itu, pemerintah dengan DPR sepakat, hingga akhir Oktober mendatang, para konglomerat penandatangan MSAA harus sudah menutupi kekurangan mereka dengan menyerahkan aset baru. Sebab para debitor menganggap tindakan tersebut memungkinkan pemerintah untuk menyita seluruh aset mereka baik yang berada di dalam maupun luar negeri. Rizal mengakui bahwa permintaan untuk meminta jaminan pribadi atau personal guarantee pada awalnya ditentang sejumlah konglomerat.”

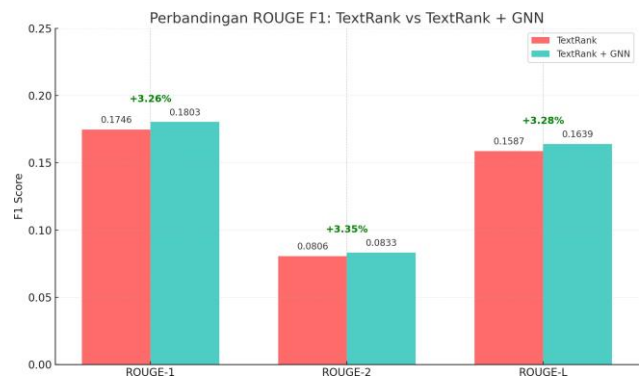


Gambar 5. Nilai ROUGE Textrank + GNN pada Studi Kasus

Tabel 5. Hasil Evaluasi ROUGE dari Studi Kasus

Metrik	F1 (GNN)	F1 (TextRank)	Precision (GNN)	Precision (TextRank)	Recall (GNN)	Recall (TextRank)
ROUGE-1	0.1803	0.1746	0.1058	0.1019	0.6111	0.6111
ROUGE-2	0.0833	0.0806	0.0485	0.0467	0.2941	0.2941
ROUGE-L	0.1639	0.1587	0.0962	0.0926	0.5556	0.5556

Dari tabel di atas, terlihat bahwa metode TextRank + GNN mengalami peningkatan performa di seluruh metrik ROUGE dibandingkan TextRank murni, dengan peningkatan tertinggi pada ROUGE-1 sebesar +3.26%, sebesar +3.35% pada ROUGE-2, dan juga sebesar +3.28% pada ROUGE-L. Meskipun secara absolut nilainya tergolong sedang, hasil ini memperkuat temuan sebelumnya bahwa GNN mampu meningkatkan kualitas rangkuman dengan mempertimbangkan fitur semantik dan struktur graf secara simultan.



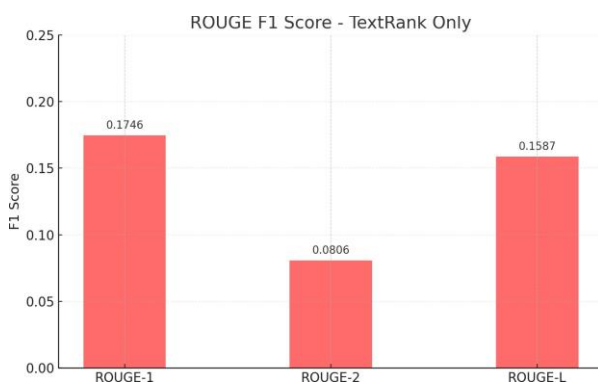
Gambar 6. Hasil Evaluasi ROUGE dari Studi Kasus

Secara kualitatif, ringkasan yang dihasilkan oleh GNN mencakup kalimat yang lebih informatif dan mencerminkan konteks narasi secara lebih utuh, misalnya dengan memasukkan unsur aktor (pemerintah, DPR), objek (aset), dan waktu (akhir

Oktober), dibandingkan TextRank yang lebih fokus pada detail angka nominal dan entitas tunggal. Hasil uji studi kasus ini memperkuat argumen bahwa kombinasi TextRank dan GNN memberikan nilai tambah dalam Menyusun ringkasan teks berita yang relevan dan kontekstual.

Tabel 6. Ringkasan oleh TextRank

“com, Jakarta: Pemerintah masih memberikan waktu dua minggu lagi kepada seluruh konglomerat yang telah menandatangani perjanjian pengembalian bantuan likuiditas Bank Indonesia dengan jaminan aset (MSAA), untuk secepatnya menyerahkan jaminan pribadi serta aset. Sejauh ini, penilaian jaminan MSAA baru dilakukan atas aset milik Grup Salim. Sementara itu, pemerintah dengan DPR sepakat, hingga akhir Oktober mendatang, para konglomerat penandatanganan MSAA harus sudah menutupi kekurangan mereka dengan menyerahkan aset baru. Sebab para debitor menganggap tindakan tersebut memungkinkan pemerintah untuk menyita seluruh aset mereka baik yang berada di dalam maupun luar negeri. Tetapi, nilai aset yang dijaminan Kelompok Salim atas utang BLBI Bank Central Asia diperkirakan tak lebih dari Rp 20 triliun.”

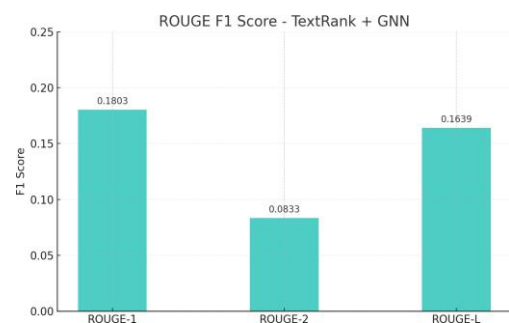


Gambar 7. Nilai ROUGE untuk TextRank pada Studi Kasus

Sementara itu, TextRank+GNN menambahkan dimensi relasional dan temporal ke dalam ringkasan, termasuk informasi kolaborasi antara pemerintah dan DPR, serta pendapat tokoh utama (Rizal). Kalimat-kalimat ini mencerminkan nuansa kebijakan dan respon publik yang lebih luas. Sebagaimana dalam tabel berikut:

Tabel 7. Ringkasan oleh TextRank-SimpleGNN

“com, Jakarta : Pemerintah masih memberikan waktu dua minggu lagi kepada seluruh konglomerat yang telah menandatangani perjanjian pengembalian bantuan likuiditas Bank Indonesia dengan jaminan aset (MSAA), untuk secepatnya menyerahkan jaminan pribadi serta aset. Sejauh ini, penilaian jaminan MSAA baru dilakukan atas aset milik Grup Salim. Sementara itu, pemerintah dengan DPR sepakat, hingga akhir Oktober mendatang, para konglomerat penandatanganan MSAA harus sudah menutupi kekurangan mereka dengan menyerahkan aset baru. Sebab para debitor menganggap tindakan tersebut memungkinkan pemerintah untuk menyita seluruh aset mereka baik yang berada di dalam maupun luar negeri. Rizal mengakui bahwa permintaan untuk meminta jaminan pribadi atau personal guarantee pada awalnya ditentang sejumlah konglomerat.”

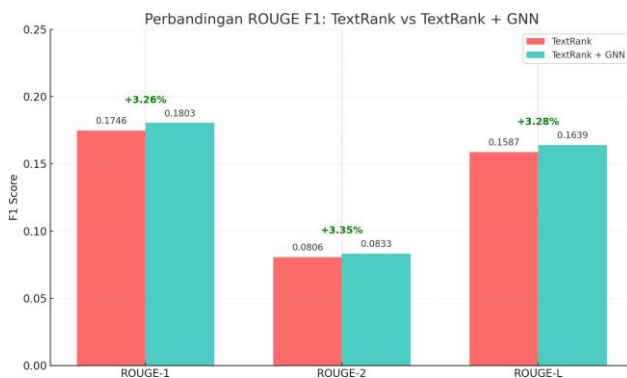


Gambar 8. Nilai ROUGE TextRank + GNN pada Studi Kasus

Tabel 8. Hasil Evaluasi ROUGE dari Studi Kasus

Metrik	F1 (GNN)	F1 (TextRank)	Precision (GNN)	Precision (TextRank)	Recall (GNN)	Recall (TextRank)
ROUGE- 1	0.1803	0.1746	0.1058	0.1019	0.6111	0.6111
ROUGE- 2	0.0833	0.0806	0.0485	0.0467	0.2941	0.2941
ROUGE- L	0.1639	0.1587	0.0962	0.0926	0.5556	0.5556

Dari tabel di atas, terlihat bahwa metode TextRank + GNN mengalami peningkatan performa di seluruh metrik ROUGE dibandingkan TextRank murni, dengan peningkatan tertinggi pada ROUGE-1 sebesar +3.26%, sebesar +3.35% pada ROUGE-2, dan juga sebesar +3.28% pada ROUGE-L. Meskipun secara absolut nilainya tergolong sedang, hasil ini memperkuat temuan sebelumnya bahwa GNN mampu meningkatkan kualitas rangkuman dengan mempertimbangkan fitur semantik dan struktur graf secara simultan.



Gambar 9. Hasil Evaluasi ROUGE dari Studi Kasus

Secara kualitatif, ringkasan yang dihasilkan oleh GNN mencakup kalimat yang lebih informatif dan mencerminkan konteks narasi secara lebih utuh, misalnya dengan memasukkan unsur aktor (pemerintah, DPR), objek (aset), dan waktu (akhir Oktober), dibandingkan TextRank yang lebih fokus pada detail angka nominal dan entitas tunggal. Hasil uji studi kasus ini memperkuat argumen bahwa kombinasi TextRank dan GNN memberikan nilai tambah dalam menyusun ringkasan teks berita yang relevan dan kontekstual.

Pembahasan

Penelitian ini mengevaluasi kinerja dua pendekatan peringkasan teks otomatis: TextRank klasik dan TextRank yang diperkaya dengan Graph Neural Network (GNN) sederhana. Evaluasi dilakukan dengan menggunakan metrik ROUGE (ROUGE-1, ROUGE-2, dan ROUGE-L) pada 1000 artikel berita berbahasa Indonesia dari dataset Liputan6. Hasil menunjukkan bahwa kombinasi TextRank dan GNN memberikan hasil yang lebih baik dibandingkan dengan TextRank murni pada semua metrik yang diuji. Peningkatan skor ROUGE pada ROUGE-1, ROUGE-2, dan ROUGE-L mengindikasikan bahwa

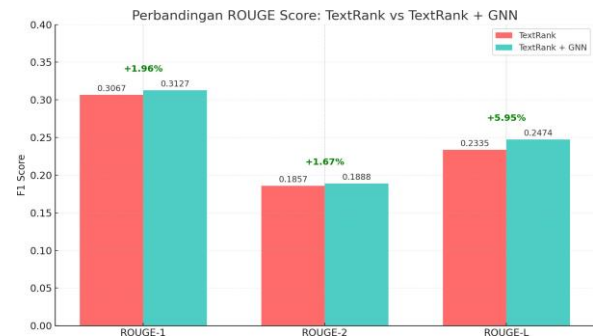
metode gabungan ini lebih efektif dalam menangkap kata-kata kunci, struktur frasa dua kata, dan hubungan semantik antar kalimat. Khususnya, ROUGE-L yang mengukur kesamaan struktur kalimat mengalami peningkatan signifikan, yang menunjukkan bahwa ringkasan yang dihasilkan lebih koheren dan mencerminkan struktur teks yang penting. Hasil ini sejalan dengan penelitian sebelumnya, seperti yang ditemukan oleh Zhang *et al.* (2020), yang menyatakan bahwa integrasi GNN dalam peringkasan teks dapat memperkaya representasi graf dan membantu model memahami hubungan semantik antar kalimat secara lebih mendalam. Gao dan Huang (2025) juga menunjukkan bahwa GNN memiliki kemampuan ekspresif yang lebih tinggi dibandingkan dengan model berbasis *bag-of-words*, yang memungkinkan sistem untuk menangkap hubungan kompleks dalam teks, sebuah keunggulan yang terlihat jelas dalam penelitian ini. Selain itu, analisis kualitatif terhadap ringkasan yang dihasilkan menunjukkan bahwa TextRank + GNN mampu menghasilkan ringkasan yang lebih komprehensif dan kontekstual. Ringkasan dari TextRank klasik cenderung terfokus pada detail angka dan entitas tertentu, seperti nominal utang, namun kurang menekankan pada aspek penting lainnya, seperti urgensi waktu atau peran otoritas negara dalam kebijakan yang dijelaskan.

Sebaliknya, TextRank + GNN berhasil mencakup informasi yang lebih luas, termasuk hubungan antara pemerintah dan DPR serta tenggat waktu yang diberikan kepada konglomerat, yang mencerminkan aspek yang lebih holistik dari narasi tersebut. Penerapan GNN pada TextRank memberikan peningkatan kualitas yang signifikan dengan memperhitungkan keterkaitan semantik antar kalimat. Fajriyah (2022) juga mencatat bahwa GNN mampu mengatasi keterbatasan yang dimiliki oleh model berbasis graf statis, seperti TextRank, dengan mengintegrasikan fitur-fitur semantik dan hubungan antar kalimat secara lebih efektif. Dengan GNN, setiap kalimat dalam graf tidak hanya dinilai berdasarkan kedekatannya secara statistik, tetapi juga berdasarkan kualitas hubungan semantiknya dengan kalimat lainnya, yang menghasilkan pemerinkatan yang lebih informatif dan akurat. Dalam pengujian kasus dunia nyata, yang melibatkan artikel yang tidak digunakan dalam pelatihan, TextRank + GNN berhasil menghasilkan ringkasan yang lebih kaya akan

informasi, mencakup elemen-elemen penting seperti aktor (pemerintah, DPR), objek (aset), dan konteks waktu (akhir Oktober). Ini menunjukkan bahwa metode ini dapat lebih baik memahami dan menangkap informasi kontekstual yang relevan dalam teks, yang sangat penting untuk menghasilkan ringkasan yang berguna dan akurat. Secara keseluruhan, hasil penelitian ini memperkuat argumen bahwa kombinasi TextRank dengan GNN dapat meningkatkan kualitas peringkasan teks otomatis. Dengan memperkaya representasi graf dan memperhitungkan hubungan semantik secara lebih mendalam, metode ini berhasil menghasilkan ringkasan yang lebih relevan, koheren, dan informatif dibandingkan dengan TextRank klasik. Temuan ini sejalan dengan penelitian oleh Yang *et al.* (2023) yang menekankan pentingnya menggali potensi Graph Neural Network dalam meningkatkan kualitas peringkasan teks. Sebagai langkah lanjutan, penelitian ini menyarankan untuk mengembangkan lebih lanjut arsitektur GNN yang lebih canggih, seperti Graph Attention Networks (GAT) atau GraphSAGE, untuk meningkatkan performa pada graf yang lebih besar dan kompleks. Selain itu, penambahan fitur linguistik atau semantik tambahan, seperti *Part-of-Speech Tagging* atau *Named Entity Recognition*, juga dapat meningkatkan kualitas representasi dan kemampuan model untuk memahami teks dengan lebih baik.

4. Kesimpulan dan Saran

Penelitian ini bertujuan untuk meningkatkan kualitas peringkasan teks berita secara otomatis dengan membandingkan metode TextRank klasik dan TextRank berbasis Graph Neural Network (GNN). Dalam implementasinya, metode klasik menggunakan algoritma TextRank berbasis cosine similarity dan bobot TF-IDF antar kalimat, sementara pendekatan berbasis GNN mengintegrasikan fitur node dan struktur graf kalimat untuk memprediksi skor pentingnya kalimat secara lebih kontekstual. Berdasarkan hasil evaluasi terhadap 1000 berita dari dataset Liputan6, diperoleh beberapa temuan penting.



Gambar 10. Perbandingan skor

Sebagai tindak lanjut dari penelitian ini, disarankan agar penelitian di bidang peringkasan teks otomatis dikembangkan lebih lanjut dengan mengeksplorasi arsitektur Graph Neural Network (GNN) yang lebih kompleks. Salah satu arsitektur yang patut dipertimbangkan adalah Graph Attention Network (GAT), yang memiliki kemampuan untuk memberikan bobot perhatian yang berbeda pada setiap hubungan antar simpul dalam graf, sehingga informasi semantik antar kalimat dapat diproses dengan lebih selektif dan sesuai prioritas. Alternatif lainnya adalah GraphSAGE, yang dapat menggeneralisasi representasi node dengan cara mengambil sampel tetangga secara efisien, serta cocok digunakan pada graf berukuran besar dan dinamis. Penggunaan model-model GNN yang lebih canggih ini diharapkan dapat memperkaya representasi hubungan semantik antar kalimat dan secara signifikan meningkatkan akurasi sistem peringkasan teks otomatis.

Selain itu, untuk meningkatkan kemampuan model dalam menangani kompleksitas bahasa alami, disarankan agar penelitian selanjutnya menambahkan fitur linguistik atau semantik tambahan pada setiap node dalam graf. Fitur-fitur ini bisa meliputi hasil dari *Part-of-Speech Tagging* (POS-tagging), *Named Entity Recognition* (NER), serta embeddings dari model bahasa berbasis deep learning seperti BERT. Dengan penambahan fitur ini, setiap kalimat dalam proses peringkasan tidak hanya akan dinilai berdasarkan distribusi kata-kata semata, tetapi juga akan mempertimbangkan aspek gramatikal dan semantik yang lebih mendalam. Hal ini akan memungkinkan model untuk lebih adaptif terhadap kompleksitas bahasa alami dan meningkatkan kemampuan model dalam menyaring informasi yang paling relevan. Dalam hal evaluasi, disarankan agar sistem

peringkasan teks tidak hanya bergantung pada metrik otomatis seperti ROUGE, yang meskipun bermanfaat, tidak sepenuhnya mencakup semua aspek kualitas ringkasan. Oleh karena itu, evaluasi manual melalui penilaian human judgment sangat penting untuk memberikan gambaran yang lebih komprehensif, terutama terkait dengan keterbacaan, koherensi, dan relevansi informasi yang dihasilkan. Kolaborasi dengan evaluator atau pakar bahasa yang memiliki pengetahuan domain yang relevan akan memberikan penilaian yang lebih objektif dan dapat dipercaya, meningkatkan kredibilitas hasil yang diperoleh. Selanjutnya, untuk meningkatkan generalisasi dan adaptasi dari sistem peringkasan teks ini, akan sangat bermanfaat jika pengujian sistem diperluas ke domain-domain lain, seperti artikel ilmiah, dokumen hukum, atau teks fiksi. Pendekatan ini bertujuan untuk menguji kemampuan model dalam menghadapi berbagai jenis teks yang memiliki karakteristik yang berbeda, serta untuk memastikan efektivitas pendekatan GNN dalam beragam konteks dunia nyata, yang sering kali memiliki kompleksitas dan variasi gaya bahasa yang lebih besar. Akhirnya, untuk mendukung efisiensi pelatihan model GNN yang lebih besar dan kompleks di masa depan, peneliti disarankan untuk menerapkan teknik optimasi seperti pruning (pengurangan parameter yang tidak penting), batching (pengelompokan data untuk pemrosesan paralel), serta memanfaatkan akselerasi perangkat keras seperti GPU dan TPU. Langkah-langkah ini sangat penting agar proses pengembangan dan eksperimen model dapat dilakukan lebih efisien, baik dari segi waktu maupun penggunaan sumber daya, tanpa mengurangi kualitas hasil yang diperoleh.

5. Daftar Pustaka

- Aulia, V. (2017). Utilizing mind mapping to summarize English text with the theme "American culture". *Journal of Educational Science and Technology*, 3(3), 218-225.
- Chen, K. Y., Liu, S. H., Chen, B., Wang, H. M., Jan, E. E., Hsu, W. L., & Chen, H. H. (2015). Extractive broadcast news summarization leveraging recurrent neural network language modeling techniques. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 23(8), 1322-1334.
<https://doi.org/10.1109/TASLP.2015.2432578>.
- Eris, E., & Pragantha, J. (2017). Penerapan algoritma textrank untuk automatic summarization pada dokumen berbahasa indonesia. *Jurnal Ilmu Teknik dan Komputer*, 1(1), 71-78.
- Gao, Q., & Huang, J. (2025). Design and Implementation of Classical Literature Sentiment Analysis System Based on Ensemble Learning and Graph Neural Network. *International Journal of Cognitive Computing in Engineering*.
- Gulati, V., Kumar, D., Popescu, D. E., & Hemanth, J. D. (2023). Extractive article summarization using integrated TextRank and BM25+ algorithm. *Electronics*, 12(2), 372.
- Hernawan, Y. F., Adikara, P. P., & Wihandika, R. C. (2022). Peringkasan Artikel Berbahasa Indonesia Menggunakan TextRank dengan Pembobotan BM25. *Jurnal Teknologi Informasi Dan Ilmu Komputer (JTIIK)*, 9(1), 61-68.
- Hossain, M. M., Anselma, L., & Mazzei, A. (2023, November). Exploring sentiments in summarization: SentiTextRank, an Emotional Variant of TextRank. In *Proceedings of the 9th Italian Conference on Computational Linguistics (CLiC-it 2023)* (pp. 535-539).
- Jain, R., Singh, P., & Puri, S. (2023, September). Summarization of Daily News Using TextRank and TF-IDF Algorithm. In *Congress on Intelligent Systems* (pp. 313-324). Singapore: Springer Nature Singapore.
- Kazemi, A., Pérez-Rosas, V., & Mihalcea, R. (2020). Biased TextRank: Unsupervised graph-based content extraction. *arXiv preprint arXiv:2011.01026*.
<https://doi.org/10.48550/arXiv.2011.01026>.
- Khaliq, A., Awan, S. A., Ahmad, F., Zia, M. A., & Iqbal, M. Z. (2024). Enhanced topic-aware

- summarization using statistical graph neural networks. *Computers, Materials and Continua*, 80(2), 3221-3242.
- Khor, Y. K., Tan, C. W., & Lim, T. M. (2021, October). Text summarization on amazon food reviews using textrank. In *International Conference on Digital Transformation and Applications (ICDXA)* (Vol. 25, pp. 113-120).
- Koto, F., Lau, J. H., & Baldwin, T. (2020). Liputan6: A large-scale Indonesian dataset for text summarization. *arXiv preprint arXiv:2011.00679*.
<https://doi.org/10.48550/arXiv.2011.00679>.
- Lestari, I. N. T., & Mulyana, D. I. (2022). Implementation of OCR (Optical Character Recognition) using Tesseract in detecting character in quotes text images. *Journal of Applied Engineering and Technological Science (JAETS)*, 4(1), 58-63.
- Mallick, C., Das, A. K., Dutta, M., Das, A. K., & Sarkar, A. (2018, August). Graph-based text summarization using modified TextRank. In *Soft Computing in Data Analytics: Proceedings of International Conference on SCD A 2018* (pp. 137-146). Singapore: Springer Singapore.
- Mulyana, D. I., Yel, M. B., & Rahmanto, M. D. (2023). Optimasi Penerapan Metode Text Recognition Dalam Fitur Catatan Otomatis Berbasis Mobile. *Jurasik (Jurnal Riset Sistem Informasi dan Teknik Informatika)*, 8(2), 489-498.
- Pan, S., Li, Z., & Dai, J. (2019, May). An improved TextRank keywords extraction algorithm. In *Proceedings of the ACM Turing Celebration Conference-China* (pp. 1-7).
- Widyassari, A. P., Rustad, S., Shidik, G. F., Noersasongko, E., Syukur, A., Affandy, A., & Setiadi, D. R. I. M. (2022). Review of automatic text summarization techniques & methods. *Journal of King Saud University-Computer and Information Sciences*, 34(4), 1029-1046.
- Wu, L., Cui, P., Pei, J., Zhao, L., & Guo, X. (2022, August). Graph neural networks: foundation, frontiers and applications. In *Proceedings of the 28th ACM SIGKDD conference on knowledge discovery and data mining* (pp. 4840-4841).
<https://doi.org/10.1145/3534678.3542609>.
- Yang, F., Zhang, H., Tao, S., & Fan, X. (2024). Simple hierarchical PageRank graph neural networks. *Journal of Supercomputing*, 80(4).
- Yang, H., & Gonçalves, T. (2025). MultiLTR: Text Ranking with a Multi-Stage Learning-to-Rank Approach. *Information*, 16(4), 308.
- Zhang, M., Li, X., Yue, S., & Yang, L. (2020). An empirical study of TextRank for keyword extraction. *IEEE access*, 8, 178849-178858.
- Zhou, J., Cui, G., Hu, S., Zhang, Z., Yang, C., Liu, Z., ... & Sun, M. (2020). Graph neural networks: A review of methods and applications. *AI open*, 1, 57-81.
- Zidan, M. R., & Setiawan, K. (2025). Implementasi Algoritma Rabin-Karp dalam Pendeteksian Plagiarisme pada Dokumen Makalah Mahasiswa. *Jurnal Indonesia: Manajemen Informatika dan Komunikasi*, 6(1), 273-284.
<https://doi.org/10.35870/jimik.v6i1.1191>.