

Implementasi *AI* Multifungsi Untuk Menelusuri Berita Dan Mengidentifikasi Potensi Hoaks Berbasis *Multi AI* Dan *Tools* Eksternal

1st Daryl Loudi Wardhana

Fakultas Ilmu Terapan

Universitas Telkom

Bandung, Indonesia

darrylloudi@student.telkomuniversity.ac.id

2nd Muhammad Rizqy Alfarisi

Fakultas Ilmu Terapan

Universitas Telkom

Bandung, Indonesia

mrizkyalfarisi@telkomuniversity.ac.id

3rd Periyadi

Fakultas Ilmu Terapan

Universitas Telkom

Bandung, Indonesia

periyadi@tass.telkomuniversity.ac.id

Abstrak - Penelitian ini mengembangkan aplikasi berbasis AI untuk mengatasi tantangan disinformasi yang berkembang pesat, didorong oleh kesadaran akan potensi AI dan keterbatasan model bahasa besar (LLM) seperti *knowledge cut-off* dan halusinasi. Aplikasi ini dirancang untuk menyediakan informasi real-time dan mendeteksi potensi misinformasi, memanfaatkan workflow *n8n* untuk mengintegrasikan berbagai tools AI dan sumber data. Workflow utama melibatkan pemrosesan input pengguna melalui Webhook, analisis gambar menggunakan OpenAI, pencarian informasi web melalui Firecrawl, dan sintesis respons yang relevan dan terverifikasi oleh LLM (Gemini, Grok, dan model lainnya). Data di-chunk dan disimpan di Qdrant untuk pengambilan informasi yang efisien. Pengembangan ini bertujuan untuk memberdayakan pengguna dalam membedakan informasi faktual dari misinformasi, yang relevan mengingat kerentanan masyarakat Indonesia terhadap disinformasi akibat literasi rendah dan pengaruh emosional. Hasil yang diharapkan adalah alat bantu deteksi berita potensial, yang dapat dikembangkan lebih lanjut untuk menghadapi tantangan penyebaran informasi cepat oleh AI. Aplikasi ini menyoroti bagaimana *n8n* dapat digunakan untuk membangun sistem AI kompleks yang memproses dan memverifikasi informasi secara efisien. Kata kunci - LLM, Gemini, Grok, AI, knowledge.

I. PENDAHULUAN

Kita berada di era digital yang ditandai oleh arus informasi tak terbatas dan kemudahan akses melalui platform online. Perkembangan teknologi informasi ini mengalami akselerasi luar biasa dengan kemajuan pesat Kecerdasan Buatan (*Artificial Intelligence*)[1]. Pengamatan langsung menunjukkan bahwa AI berkembang dengan kecepatan yang bahkan melampaui kapasitas belajar manusia, menandakan potensi perkembangan di masa depan. Model bahasa besar (*Large Language Models - LLM*) seperti *GPT*, *Gemini*, dan *Grok* adalah manifestasi nyata dari kemajuan ini[2][3], menunjukkan kemampuan mengolah dan menghasilkan bahasa alami secara canggih, serta semakin terintegrasi dalam berbagai aplikasi. Kondisi ini menciptakan peluang sekaligus tantangan baru yang fundamental. Di balik kemudahan akses informasi dan kecanggihan AI, terdapat tantangan krusial berupa penyebaran masif disinformasi dan

misinformasi[4]. Konten salah atau menyesatkan ini menyebar dengan cepat di ruang digital, mengancam kepercayaan publik, dan proses demokrasi. Penanganan masalah ini semakin terasa di Indonesia, mengingat data menunjukkan tingkat literasi digital yang perlu ditingkatkan[5][6], serta adanya kerentanan terhadap narasi emosional yang memicu penyebaran[10]. Kegagalan mengatasi ini berdampak nyata pada berbagai aspek kehidupan berbangsa. Teknologi LLM yang potensial menjadi bagian dari solusi, justru memiliki kelemahan yang perlu dicermati saat digunakan untuk verifikasi fakta. Salah satu kelemahan utama adalah *knowledge cut-off*. Sebagai contoh, model seperti *GPT* mungkin memiliki basis pengetahuan yang terbatas hingga periode tertentu (misalnya, 2022 hingga awal 2024)[7], membuatnya tidak dapat diandalkan untuk informasi mengenai peristiwa terkini. Keterbatasan pengetahuan ini seringkali memicu kelemahan kedua, yaitu halusinasi. LLM dapat menghasilkan informasi yang terdengar sangat meyakinkan namun sebenarnya keliru, tidak akurat, atau bahkan mengarang sumber. Pengamatan menunjukkan kasus di mana LLM menciptakan tautan web yang fiktif dan tidak valid ketika mencoba menjawab pertanyaan di luar batas pengetahuannya. Model bahasa besar standar memiliki dua keterbatasan utama, beberapa larangan mengakses informasi terkini dan potensi menghasilkan informasi yang tidak benar[8]. Hal ini menyebabkan LLM tersebut kurang cocok dan berisiko jika digunakan sendiri untuk mengatasi disinformasi yang terus berubah. Menyadari masalah disinformasi serta kelemahan spesifik yang teramati pada LLM (*knowledge cut-off* dan halusinasi seperti pembuatan link fiktif), penelitian ini mengusulkan sebuah solusi melalui pengembangan aplikasi AI. Berbeda dengan mengandalkan satu LLM saja, aplikasi ini dirancang secara khusus untuk mengatasi keterbatasan tersebut dengan memungkinkan penelusuran informasi secara real-time. Solusi ini memanfaatkan platform workflow *n8n* untuk secara sistematis mengintegrasikan beberapa tools: kemampuan analisis LLM (*Gemini*, *Grok*)[2][3], digabungkan dengan pencarian web aktual via *Firecrawl*, kemampuan analisis gambar *OpenAI*, dan pengelolaan data efisien via *Qdrant*. Tujuannya adalah agar respons yang diberikan kepada pengguna tidak hanya didasarkan pada pengetahuan internal

LLM yang mungkin usang atau salah, tetapi telah diverifikasi silang dengan informasi terbaru dari web. Meskipun dari segi cakupan pengetahuan murni mungkin tidak seluas model dasar *LLM*[9], aplikasi ini dirancang untuk unggul dalam menyediakan informasi yang lebih akurat, relevan secara waktu, dan terverifikasi. Dengan demikian, penelitian ini bertujuan mengembangkan alat bantu praktis yang memberdayakan pengguna untuk memvalidasi informasi secara kritis dan lebih efektif dalam menghadapi tantangan disinformasi di era *AI*.

A. Rumusan Masalah

1. Bagaimana cara mengatasi penyebaran disinformasi dan misinformasi yang masif di era digital, terutama dengan mempertimbangkan keterbatasan model bahasa besar seperti *knowledge cut-off* dan halusinasi?
2. Bagaimana mengembangkan aplikasi berbasis *AI* yang mampu menelusuri informasi secara *real-time* dan memverifikasi keakuratannya dari berbagai sumber web untuk mendeteksi potensi hoaks?
3. Bagaimana mengintegrasikan berbagai *tools AI* dan sumber data eksternal dalam sebuah sistem yang efisien menggunakan platform *workflow n8n* untuk meningkatkan akurasi deteksi informasi hoaks?

B. Tujuan

Dalam proyek ini penulis bertujuan untuk mengembangkan aplikasi yang dapat:

1. Mengidentifikasi suatu berita atau informasi menggunakan *multi model* yang terus berkembang dan lebih bagus dari waktu ke waktu.
2. Aplikasi dapat mengidentifikasi potensi berita hoaks pada suatu informasi yang telah tersebar.
3. Mengintegrasikan *multi AI*, *tools* dan sumber data eksternal menggunakan platform *n8n* untuk meningkatkan akurasi dalam pendeteksian hoaks

C. Batasan Masalah

Sistem masih dalam pengembangan karena ada beberapa batasan masalah *Scraping/Crawling* terkadang bermasalah dan memakan waktu sangat lama karena keamanan *website* itu contoh:

1. Situs seperti *youtube* atau yang memiliki konten berupa video, karena saat ini belum ada model *LLM* atau *Tools* yang mampu melakukan itu dan itu perlu penerapan *tools* lain seperti *Speech Recognition* dan *Computer Vision* yang lebih kompleks.
2. Situs lain yang memiliki *anti scraping*, *anti bot*, dan keamanan pendeteksi *traffic* mencurigakan, ataupun situs seperti (*scribd.com*) yang terkadang sulit untuk diambil datanya karena di kunci dengan keamanan yang tidak bisa ditembus oleh *tools scraper* ini.
3. *Timeout* dari keamanan *Cloudflare*, *Cloudflare* merupakan layanan perlindungan *website* yang sering kali memicu *error timeout 524* ketika *tools scraper* mencoba mengakses situs yang dilindungi. *Error 524* ini terjadi ketika *Cloudflare* mendeteksi bahwa *server* asal (*origin server*) membutuhkan waktu yang terlalu lama untuk merespons

permintaan, melebihi batas waktu yang ditetapkan oleh *Cloudflare* biasanya 100 detik.

II. KAJIAN TEORI

A. Kecerdasan Buatan dan Model Bahasa Besar (*LLM*)

Kecerdasan Buatan (*AI*) memungkinkan sistem melakukan tugas cerdas seperti pemrosesan bahasa alami [1]. *Large Language Models (LLM)* seperti *Gemini*, *Grok*, dan *GPT* mampu menghasilkan teks mirip manusia, namun memiliki keterbatasan berupa *knowledge cut-off* (batasan pengetahuan hingga periode tertentu) dan halusinasi (informasi tidak akurat) [7][8]. Keterbatasan ini menghambat verifikasi informasi terkini, terutama untuk deteksi hoaks.

B. Pendekatan *Multi AI*

Menggunakan beberapa model *AI* (*Gemini*, *Grok*) meningkatkan akurasi dengan meminimalkan halusinasi melalui verifikasi silang [2][3]. Pendekatan ini efektif untuk mendeteksi hoaks yang memerlukan analisis konteks dari berbagai sumber.

C. *Workflow n8n*

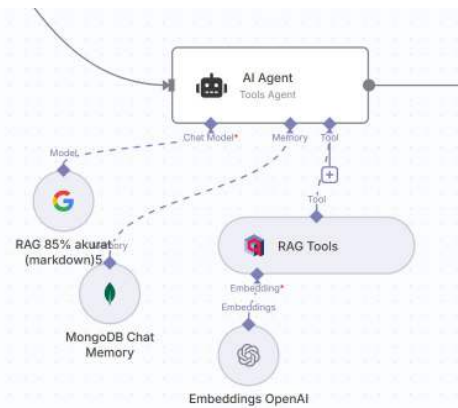
n8n adalah platform *open-source* untuk otomatisasi *workflow* yang memungkinkan integrasi berbagai aplikasi, *API*, dan layanan eksternal melalui *nodes* yang dapat dikonfigurasi [9]. Dalam konteks penelitian ini, *n8n* digunakan untuk mengintegrasikan model *AI* seperti *Gemini*, *Grok*, dan *OpenAI*, serta alat seperti *Firecrawl* untuk *web scraping* dan *Qdrant* untuk pengelolaan data. *Workflow n8n* memungkinkan pemrosesan input pengguna secara sistematis, mulai dari penerimaan melalui *Webhook*, analisis data, hingga pengiriman respons yang telah diverifikasi. Pendekatan ini memastikan bahwa sistem dapat menghasilkan informasi yang relevan, terverifikasi, dan bebas dari halusinasi dengan memanfaatkan data *real-time* dari sumber eksternal.

III. METODE

Proses *workflow* saat ini melibatkan 3 *workflow*:

1. *General Workflow*:

- *Workflow* bertugas untuk menangani percakapan biasa berdasarkan *knowledge* dari model nya sendiri dan juga referensi data dari *vector store* nya.
- Setelah pengguna menyelesaikan *workflow* kedua, yaitu fitur *Pencarian Web*, data *markdown* (hasil *scraping*) disimpan ke dalam dua *database* terpisah. Satu *database* khusus untuk data *markdown* mentah, dan *database* lainnya menyimpan hasil akhir yang telah diproses dan dibersihkan dari elemen *HTML*, *CSS*, dan lain-lain, sehingga berisi jawaban murni dari alur kerja kedua.



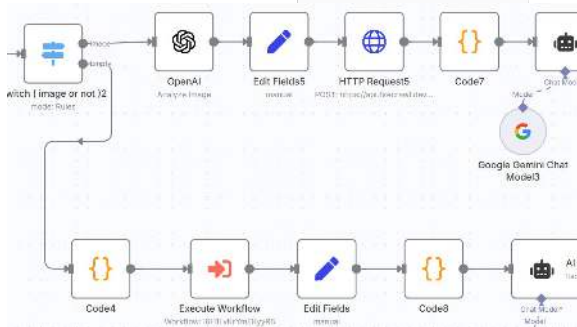
GAMBAR 1
(WORKFLOW GENERAL)

Sistem cara kerjanya

Pengguna memasukkan input, memicu *webhook* untuk mengklasifikasikan kelengkapan data yang diterima. Jika tidak ada gambar yang dikirim, alur berlanjut untuk memproses *input*. Agen pertama menerima *input*, memecah *query* menjadi beberapa *query* berbeda agar pencarian di *vector store* lebih akurat dan relevan, lalu mengirimkan jawaban ke agen yang bertugas memproses *query* pengguna berdasarkan data dari agen pertama. Agen kedua memeriksa riwayat percakapan (*recall memory*), meninjau topik yang telah dibahas, dan mencocokkannya dengan *query* saat ini. Hasil akhir diproduksi dan dikirim kembali melalui *webhook* untuk ditampilkan di antarmuka pengguna.

2. Pencarian *Web*:

Workflow ini berperan penting dalam aplikasi ini karena ini adalah fungsi utama nya, pada saat user melakukan *input data query* tersebut terkirim ke *workflow* ini melalui *webhook*, dan ini sekaligus mengaktifkan fungsi pencarian *SerpAPI* yang mana berfungsi untuk melakukan pencarian informasi tambahan.



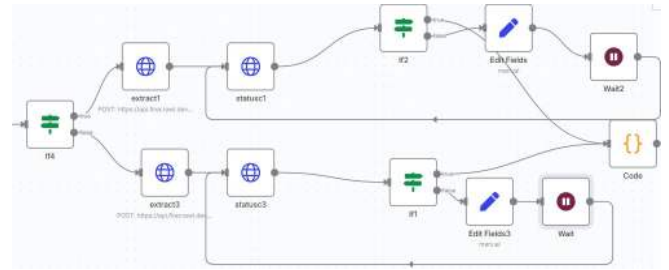
GAMBAR 2
(WORKFLOW WEB SEARCH)

Cara kerja nya kurang lebih sama seperti *Workflow* pertama, Pada *workflow* ini ada pemanggilan *workflow* lain yang bertugas untuk *scraping* data di *website*:

3. *Scraper*:

Workflow ini memanggil fungsi dari *firecrawl* untuk melakukan *scraping* dan penelusuran data berdasarkan *query* yang diterima sebelumnya oleh *workflow* kedua untuk mencari informasi dari 8 situs memakan waktu kurang lebih 30 detik hingga 1 menit tergantung dari keamanan *website* nya itu sendiri, setelah itu hasil mentah dari *scraping data* nya

kembali ke *workflow* kedua untuk diproses kembali dimasukkan ke dua *database* berbeda untuk filterisasi data bersih dan data mentah nya.



GAMBAR 3
(WORKFLOW SCRAPER)

TABEL 1
(DATA PENGUJIAN)

No	Nama Tool/Model	Type	Keunggulan	Pertimbangan
1	<i>Gpt 4.1</i>	<i>LLM</i>	Unggul dalam hal pertimbangan dan penalaran lebih lanjut tetapi kekurangannya waktu pemrosesan lama dan sangat mahal	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal
2	<i>o3</i>	<i>LLM</i>	Unggul dalam hal pertimbangan dan penalaran lebih lanjut tetapi kekurangannya waktu pemrosesan lama dan sangat mahal	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal
3	<i>o3-mini</i>	<i>LLM</i>	Unggul dalam hal thinking mode dan lebih cepat dibandingkan o3 serta lebih murah	Model cocok digunakan untuk menalar hasil akhir dari output LLM sebelum diberikan ke user, tetapi waktu pemrosesan lebih lama dan biaya cukup mahal
4	<i>o4</i>	<i>LLM</i>	Model khusus untuk thinking atau reasoning yang sangat unggul di kelasnya,	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal

			tetapi biaya nya sangat mahal						
5	<i>Gpt 4.1 mini</i>	<i>LLM</i>	Unggul untuk menalar hasil output dari model LLM, dan waktu pemrosesan nya lumayan cepat	Model cocok digunakan untuk melakukan pemrosesan input dan output saja, tetapi dalam hal menalar potensi hoax model ini kurang cocok					
6	<i>Gpt 4.1 nano</i>	<i>LLM</i>	Sangat unggul dalam memproses Input dan output, model tercepat di kelasnya dan biaya yang sangat murah	Model cocok digunakan untuk melakukan pemrosesan input dan output saja, tetapi dalam hal menalar potensi hoax model ini kurang cocok					
7	<i>Gpt 4o mini</i>	<i>LLM</i>	Unggul untuk menalar hasil output dari model LLM	Model cocok digunakan untuk melakukan pemrosesan input dan output saja, tetapi dalam hal menalar potensi hoax model ini kurang cocok					
8	<i>Gpt 4o</i>	<i>LLM</i>	Unggul dalam hal pertimbangan dan penalaran lebih lanjut tetapi kekurangan nya waktu pemrosesan lama dan sangat mahal	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal					
9	<i>Gemini 2.0 Flash</i>	<i>LLM</i>	Unggul untuk menalar hasil output dari model LLM, dan waktu pemrosesan nya lumayan cepat	Model cocok digunakan untuk melakukan pemrosesan input dan output saja, tetapi dalam hal menalar potensi hoax model ini kurang cocok					
10	<i>Gemini 2.0 Flash Thinking</i>	<i>LLM</i>	Unggul dalam hal reasoning model dan memiliki kecepatan pemrosesan lumayan cepat, dan juga murah	Model cocok digunakan untuk melakukan pemrosesan input dan output serta waktu pemrosesan yang lebih cepat					
11	<i>Gemini 2.0</i>	<i>LLM</i>	Unggul	Model cocok					
	<i>Flash lite</i>						dalam kecepatan pemrosesan lumayan cepat, dan juga murah	digunakan untuk melakukan pemrosesan input dan output saja, tetapi dalam hal menalar potensi hoax model ini kurang cocok	
12	<i>Gemini 2.5 Flash Thinking</i>	<i>LLM</i>					Lebih unggul dibandingkan versi 2.0 nya dan pemrosesan sedikit lebih lambat	Model cocok digunakan untuk melakukan penelitian potensi hoax yang memerlukan pemrosesan lebih lama tetapi output lebih detail.	
13	<i>Gemini 2.5 Pro</i>	<i>LLM</i>					Unggul dalam hal reasoning dan pemrosesan input/output, saat ini adalah model terbaik menurut benchmark dibandingkan model LLM lainnya	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal	
14	<i>Grok 3 mini</i>	<i>LLM</i>					Unggul dalam pemrosesan input dan output, model tercepat saat ini lebih cepat dibandingkan GPT 4.1 nano, biaya hampir sama	Model cocok digunakan untuk melakukan penelitian potensi hoax yang memerlukan pemrosesan lebih lama tetapi output lebih detail.	
15	<i>Grok 3</i>	<i>LLM</i>					Unggul dalam hal reasoning dan pemrosesan input output, biaya cukup mahal	Setelah uji coba sejauh ini model tidak digunakan karena waktu yang diperlukan untuk memproses seluruh workflow terlalu lama dan biaya sangat mahal	
16	<i>Claude 3 Sonnet + Claude 3.5 Sonnet</i>	<i>LLM</i>					Unggul dalam hal reasoning dan pemrosesan input output, biaya mahal	Setelah uji coba sejauh ini model tidak support dan memerlukan penelitian lebih lanjut untuk mengatur format output yang tidak terstruktur	
17	<i>Claude 3 Haiku + Claude 3.5 Haiku</i>	<i>LLM</i>					Unggul dalam hal reasoning dan pemrosesan	Setelah uji coba sejauh ini model tidak support dan memerlukan penelitian lebih	

			input output, biaya cukup mahal	lanjut untuk mengatur format output yang tidak terstruktur
18	<i>Claude 3 Opus</i>	<i>LLM</i>	Unggul dalam Pemrosesan input output, biaya cukup mahal	Setelah uji coba sejauh ini model tidak support dan memerlukan penelitian lebih lanjut untuk mengatur format output yang tidak terstruktur
19	<i>Claude 3.7 Sonnet</i>	<i>LLM</i>	Unggul dalam reasoning thinking model peringkat ketiga setelah o4 dari OpenAI menurut data benchmark	Setelah uji coba sejauh ini model tidak support dan memerlukan penelitian lebih lanjut untuk mengatur format output yang tidak terstruktur
20	<i>Apify Instagram, Tiktok, X Scraper</i>	<i>Tools Scraping</i>	Unggul dalam scraping data dari media sosial serta google, tetapi biaya sangat mahal	Sangat cocok digunakan untuk produksi kedepannya tetapi memiliki biaya yang sangat mahal
21	<i>Firecrawl Extract + Scraper + Mapping</i>	<i>Tools Scraping</i>	Unggul dalam scraping data dari google, tidak support media sosial, biaya murah	Cocok digunakan untuk keperluan produksi kecil seperti bahan pembelajaran, dan tidak disarankan untuk digunakan produksi besar
22	<i>SerpAPI</i>	<i>Tools Snippet</i>	Unggul dalam melakukan Snippet layaknya user melakukan pencarian ke google langsung	Tools ini bekerja seperti user melakukan pencarian langsung di Google, namun kekurangannya adalah tidak dapat menelusuri isi web secara langsung kecuali dengan berlangganan dengan biaya yang cukup mahal
23	<i>text embedding 3 large</i>	<i>embedding</i>	Unggul dalam pemrosesan pemecahan data menjadi chunk	Sangat cocok digunakan untuk embedding vector khususnya untuk kasus ini karena model ini berperan sangat penting untuk memecah data hasil scraping menjadi beberapa chunk kedalam database.

IV. HASIL DAN PEMBAHASAN

Hasil akhir dari proyek tugas akhir ini adalah sebuah sistem aplikasi berbasis *AI* multifungsi yang dirancang untuk menelusuri berita dan mengidentifikasi potensi hoaks. Sistem ini memanfaatkan integrasi beberapa model bahasa besar (*LLM*) dan berbagai *tools* eksternal melalui platform *workflow n8n*. Sistem ini bertujuan untuk memberikan informasi yang akurat, relevan, dan terverifikasi kepada pengguna dengan mengatasi keterbatasan *LLM*, seperti *knowledge cut-off* dan halusinasi.

Secara rinci, luaran dari proyek ini meliputi:

1. **Sistem Aplikasi Terintegrasi**
Sebuah sistem aplikasi yang mampu menerima input pengguna berupa teks atau gambar melalui *Webhook*, memprosesnya menggunakan berbagai *tools AI* dan sumber data, dan memberikan jawaban yang relevan.
2. **Platform Workflow n8n**
Sebagai pusat untuk mengintegrasikan berbagai *tools AI* (*Gemini*, *Grok*, *OpenAI*) dan sumber data eksternal (*Firecrawl*, *Qdrant*, *SerpAPI*). *Workflow* ini mencakup langkah-langkah seperti:
 - Penerimaan input pengguna melalui *Webhook*.
 - Analisis gambar menggunakan *OpenAI*
 - Pencarian informasi web secara *real-time* menggunakan *Firecrawl* dan *SerpAPI*
 - Penyimpanan dan pengolahan data hasil *scraping* di *Qdrant* menggunakan *text embedding 3 large*.
 - Strategi jawaban menggunakan berbagai *LLM* (*Gemini 2.0 Flash Thinking*, *Gemini 2.5 Flash Thinking*, *Grok 3 mini*) dengan pemilihan model dinamis sesuai kebutuhan pemrosesan dan biaya.
 - Pengiriman jawaban kembali ke pengguna melalui *Webhook*.
3. **Database Vektor Qdrant**
Database vektor yang berfungsi untuk menyimpan dan mengelola data hasil *scraping* yang *di-chunk* dan *di-embed* menggunakan model *text embedding 3 large*. *Database* ini memungkinkan pengambilan informasi yang efisien dan relevan berdasarkan *query* pengguna.
4. **Pemilihan dan Evaluasi Model LLM**
Hasil evaluasi dan perbandingan berbagai model *LLM* yang menghasilkan pemilihan model *Gemini 2.0 Flash Thinking*, *Gemini 2.5 Flash Thinking*, dan *Grok 3 mini* sebagai model yang paling efisien dan efektif untuk kebutuhan sistem, dengan pertimbangan kecepatan, akurasi, dan biaya.
5. **Script dan Konfigurasi**
Berbagai *script* dan konfigurasi yang diperlukan untuk mengintegrasikan *tools* dan *API* yang berbeda, serta mengkonfigurasi *workflow n8n*. Ini termasuk konfigurasi *API keys*, *endpoints*, dan *request parameters*.

Hasil ini secara bersama-sama membentuk sebuah sistem yang berfungsi untuk menelusuri berita dari berbagai sumber *web*, memverifikasi keakuratannya, dan mendeteksi potensi

hoaks dengan memanfaatkan kekuatan berbagai model *AI* dan *tools* eksternal. Sistem ini diharapkan dapat memberikan solusi praktis untuk mengatasi penyebaran disinformasi di era digital.

V. KESIMPULAN

Penelitian ini berhasil mengembangkan aplikasi *AI* multifungsi yang terbukti mampu mengidentifikasi kebenaran suatu berita dengan efektif melalui validasi terhadap sumber-sumber informasi terpercaya. Hasil pengujian menunjukkan bahwa sistem ini mencapai tingkat akurasi yang signifikan, berkisar antara 80% hingga 95%. Tingkat akurasi ini bergantung pada kuantitas dan sejauh mana informasi terkait telah tersebar serta dipublikasikan oleh media-media besar yang kredibel, yang menjadi acuan bagi sistem dalam proses identifikasi.

REFERENSI

- [1] “Empowered Minds: AI and the New Era of Digital Mindfulness,” Nov. 2024, doi: 10.59646/dm/285.
- [2] J. Chen and Y. Bao, “A Multi-Agent Large Language Model (Llm) Framework for Code-Complying Design Automation of Concrete Structures,” Elsevier BV, 2025. Accessed: May 13, 2025. [Online]. Available: <https://doi.org/10.2139/ssrn.5193679>
- [3] B. Nadimi and H. Zheng, “A Multi-Expert Large Language Model Architecture for Verilog Code Generation,” in 2024 IEEE LLM Aided Design Workshop (LAD), IEEE, Jun. 2024, pp. 1–5. Accessed: May 13, 2025. [Online]. Available: <https://doi.org/10.1109/lad62341.2024.10691683>
- [4] A. Aminudin, “Menghadapi Disinformasi Konten Berita Digital di Era Post Truth,” JURNAL LENSEA MUTIARA KOMUNIKASI, vol. 6, no. 2, pp. 283–292, Dec. 2022, doi: 10.51544/jlmk.v6i2.3137.
- [5] A. N. Ma’rifah, “Tingkat Literasi Aksesibilitas Wisatawan Domestik di Indonesia,” Ekodestinas, vol. 1, no. 1, pp. 20–26, Mar. 2023, doi: 10.59996/ekodestinas.v1i1.35.
- [6] S. Rahma, R. Ramly, and N. Nurhusna, “TINGKAT LITERASI DIGITAL GURU BAHASA INDONESIA DALAM MENYAJIKAN PEMBELAJARAN TINGKAT SMA DI KABUPATEN GOWA,” Titik Dua: Jurnal Pembelajaran Bahasa dan Sastra Indonesia, vol. 3, no. 3, Oct. 2023, doi: 10.59562/titikdua.v3i3.47380.
- [7] S. Kim, D. Kim, S. Hwang, K.-H. Lee, and K. Lee, “LLM-Assisted Ontology Restriction Verification With Clustering-Based Description Generation,” IEEE Access, vol. 13, pp. 73603–73618, 2025, doi: 10.1109/access.2025.3562560.
- [8] “Review for ‘AI for AI: Using AI methods for classifying AI science documents,’” Feb. 2022, doi: 10.1162/qss_a_00223/v1/review1.
- [9] T. Händler, “A Taxonomy for Autonomous LLM-Powered Multi-Agent Architectures,” in Proceedings of the

15th International Joint Conference on Knowledge Discovery, Knowledge Engineering and Knowledge Management, SCITEPRESS - Science and Technology Publications, 2023, pp. 85–98. Accessed: May 13, 2025. [Online]. Available: <https://doi.org/10.5220/0012239100003598>

[10] T. Konflik, K. Kriminal, B. Di, and P. Dengan Pendekatan, “Analisis Parsing Data Sosial Media”.