

DAFTAR PUSTAKA

- Aisah, A. (2019). *STUDENTS' CHALLENGES IN UNDERGRADUATE THESIS DEFENSE*.
- Alberts, I. L., Mercolli, L., Pyka, T., Prenosil, G., Shi, K., Rominger, A., & Afshar-Oromieh, A. (2023). Large language models (LLM) and ChatGPT: what will the impact on nuclear medicine be? In *European Journal of Nuclear Medicine and Molecular Imaging* (Vol. 50, Issue 6, pp. 1549–1552). Springer Science and Business Media Deutschland GmbH. <https://doi.org/10.1007/s00259-023-06172-w>
- Alfath Khairudin, & Raharjo Fauzi Fajar. (2019). TEKNIK PENGOLAHAN HASIL ASESMEN: TEKNIK PENGOLAHAN DENGAN MENGGUNAKAN PENDEKATAN ACUAN NORMA (PAN) DAN PENDEKATAN ACUAN PATOKAN (PAP). *Jurnal Komunikasi Dan Pendidikan Islam*, 8(1).
- Alhindi, T., Chakrabarty, T., Musi, E., & Muresan, S. (2023). *Multitask Instruction-based Prompting for Fallacy Recognition*. <http://arxiv.org/abs/2301.09992>
- Anusha, M., Pavan Kumar, K., Vemuri, S., Reddy, V. M., & Vaishnavi, T. (2012). IJFANS INTERNATIONAL JOURNAL OF FOOD AND NUTRITIONAL SCIENCES Speech-to-Text and Text-to-Speech Recognition. *Rights Reserved*, 11, 2022. <https://doi.org/10.48047/IJFANS/V11/Splis5/45>
- Balaguer, A., Benara, V., Cunha, R. L. de F., Filho, R. de M. E., Hendry, T., Holstein, D., Marsman, J., Mecklenburg, N., Malvar, S., Nunes, L. O., Padilha, R., Sharp, M., Silva, B., Sharma, S., Aski, V., & Chandra, R. (2024). RAG vs Fine-tuning: Pipelines, Tradeoffs, and a Case Study on Agriculture. *ArXiv*, 2401.08406. <http://arxiv.org/abs/2401.08406>
- Chen, L., Deng, Y., Bian, Y., Qin, Z., Wu, B., Chua, T.-S., & Wong, K.-F. (2023). Beyond Factuality: A Comprehensive Evaluation of Large Language Models as Knowledge Generators. *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, 6325–6341. <http://arxiv.org/abs/2310.07289>

- DocsBot AI. (2024). *RAG-Fusion: Multi-Query Retrieval & Rank Fusion Techniques*. <https://docsbot.ai/article/advanced-rag-techniques-multiquery-and-rank-fusion>
- Es, S., James, J., Espinosa-Anke, L., & Schockaert, S. (2023). RAGAS: Automated Evaluation of Retrieval Augmented Generation. *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics: System Demonstrations*, 150–158. <http://arxiv.org/abs/2309.15217>
- Fatmawati, J., & Laksmiwati, H. (2022). Hubungan antara Efikasi Diri dengan Kecemasan Menghadapi Ujian Skripsi pada Mahasiswa Hermien Laksmiwati. *Jurnal Penelitian Psikologi*, 9(8).
- Fitria, T. N. (2022). ANALYSIS OF EFL STUDENTS' DIFFICULTIES IN WRITING AND COMPLETING ENGLISH THESIS. *LLT Journal: Journal on Language and Language Teaching*, 25(1), 295–309. <https://doi.org/10.24071/llt.v25i1.3607>
- Gao, Y., Xiong, Y., Gao, X., Jia, K., Pan, J., Bi, Y., Dai, Y., Sun, J., Wang, M., & Wang, H. (2023). *Retrieval-Augmented Generation for Large Language Models: A Survey*. <http://arxiv.org/abs/2312.10997>
- Han, Y., Liu, C., & Wang, P. (2023). A Comprehensive Survey on Vector Database: Storage and Retrieval Technique, Challenge. *ArXiv*, 2310.11703. <http://arxiv.org/abs/2310.11703>
- Huang, H., Qu, Y., Bu, X., Zhou, H., Liu, J., Yang, M., Xu, B., & Zhao, T. (2024a). An Empirical Study of LLM-as-a-Judge for LLM Evaluation: Fine-tuned Judge Model is not a General Substitute for GPT-4. *ArXiv*, 2403.02839. <http://arxiv.org/abs/2403.02839>
- Huang, H., Qu, Y., Bu, X., Zhou, H., Liu, J., Yang, M., Xu, B., & Zhao, T. (2024b). An Empirical Study of LLM-as-a-Judge for LLM Evaluation: Fine-tuned Judge Model is not a General Substitute for GPT-4. *ArXiv*, 2403.02839. <http://arxiv.org/abs/2403.02839>

- Izacard, G., Lewis, P., Lomeli, M., Hosseini, L., Petroni, F., Schick, T., Dwivedi-Yu, J., Joulin, A., Riedel, S., & Grave, E. (2022). Atlas: Few-shot Learning with Retrieval Augmented Language Models. *Journal of Machine Learning Research*, 24(251), 1–43. <http://arxiv.org/abs/2208.03299>
- Jasiah, Kusumawati Rahmania Ita, Sutiharni, Febrina Wetri, & S Eflina Yetti. (2023). Pelatihan Sistematika Penulisan Skripsi bagi Mahasiswa. *Masyarakat Berdaya Dan Inovasi*, 4(1), 58–64.
- Jeffrey Ip. (2025a). *Answer Relevancy | DeepEval - The Open-Source LLM Evaluation Framework*. <https://docs.confident-ai.com/docs/metrics-answer-relevancy>
- Jeffrey Ip. (2025b). *Contextual Relevancy | DeepEval - The Open-Source LLM Evaluation Framework*. <https://docs.confident-ai.com/docs/metrics-contextual-relevancy>
- Jeffrey Ip. (2025c). *Faithfulness | DeepEval - The Open-Source LLM Evaluation Framework*. <https://docs.confident-ai.com/docs/metrics-faithfulness>
- Jiang, Z., Xu, F. F., Gao, L., Sun, Z., Liu, Q., Dwivedi-Yu, J., Yang, Y., Callan, J., & Neubig, G. (2023). Active Retrieval Augmented Generation. *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, 7969–7992. <http://arxiv.org/abs/2305.06983>
- Kong, A., Zhao, S., Chen, H., Li, Q., Qin, Y., Sun, R., Zhou, X., Wang, E., & Dong, X. (2023). *Better Zero-Shot Reasoning with Role-Play Prompting*. <http://arxiv.org/abs/2308.07702>
- Kothadiya, D., Pise, N., & Bedekar, M. (2020). Different Methods Review for Speech to Text and Text to Speech Conversion. *International Journal of Computer Applications*, 175(20), 9–12. <https://doi.org/10.5120/ijca2020920727>
- Lakshmi, K., Chandra, M. T., & Rao, S. (n.d.). Design And Implementation Of Text To Speech Conversion Using Raspberry PI. In *IJITR) INTERNATIONAL*

JOURNAL OF INNOVATIVE TECHNOLOGY AND RESEARCH (Vol. 4, Issue 6). <http://www.ijitr.com>

Lantsoght, E. O. L. (2021). *Preparation for the Doctoral Defense: Methods and Relation to Defense Outcome and Perception*. <https://doi.org/10.20944/preprints202109.0481.v1>

Lantsoght, E. O. L. (2022). Effectiveness of Doctoral Defense Preparation Methods. *Education Sciences*, 12(7). <https://doi.org/10.3390/educsci12070473>

Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. *In Proceedings of the 37th International Conference on Neural Information Processing Systems*, 1–29. <http://arxiv.org/abs/2005.11401>

Li, H., Su, Y., Cai, D., Wang, Y., & Liu, L. (2022). A Survey on Retrieval-Augmented Text Generation. *ArXiv*, 2202.01110. <http://arxiv.org/abs/2202.01110>

Ma, X., Gong, Y., He, P., Zhao, H., & Duan, N. (2023). *Query Rewriting for Retrieval-Augmented Large Language Models*. <http://arxiv.org/abs/2305.14283>

Modran, H., Bogdan, I. C., Ursuțiu, D., Samoila, C., & Modran, P. L. (2024). *LLM Intelligent Agent Tutoring in Higher Education Courses using a RAG Approach*. <https://doi.org/10.20944/preprints202407.0519.v1>

Neumann, A. T., Yin, Y., Sowe, S., Decker, S., & Jarke, M. (2024). An LLM-Driven Chatbot in Higher Education for Databases and Information Systems. *IEEE Transactions on Education*, 1–14. <https://doi.org/10.1109/TE.2024.3467912>

Nursidiq Cahyana. (2016). HUBUNGAN REGULASI DIRI DENGAN KECEMASAN MENGHADAPI UJIAN SKRIPSI PADA MAHASISWA PROGRAM STUDI PENDIDIKAN EKONOMI UNIVERSITAS

- MUHAMMADIYAH PURWOREJO. *Jurnal Ilmiah Ekonomi Dan Pembelajarannya*, 4(6).
- Patel, B., & Chand, B. B. (2024). RAG-based Chat Application using LLMs: A Case Study of Vikram Sarabhai Library IIM Ahmedabad. *12th Convention PLANNER-2024 Rajiv Gandhi University*.
- Pezeshkpour, P. (2023). Measuring and Modifying Factual Knowledge in Large Language Models. *International Conference on Machine Learning and Applications (ICMLA)*, 831–838. <http://arxiv.org/abs/2306.06264>
- pgpaul. (2023, August 24). *Dokumen Pendukung Skripsi*. <https://pgpaul.uad.ac.id/menu-skripsi.pgpaul-uad>
- Pinecone. (2024). *Rerankers and Two-Stage Retrieval*. <https://www.pinecone.io/learn/series/rag/rerankers/>
- Putri Anugrah. (2023, February 23). *Bagaimana Tahapan Kuliah S1 Sampai Bisa Wisuda? | Danacita*. <https://danacita.co.id/blog/bagaimana-tahapan-kuliah-s1-sampai-bisa-wisuda/>
- Reimers, N., & Gurevych, I. (2019). *Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks*. <http://arxiv.org/abs/1908.10084>
- Salton, G., Wong, A., & Yang, C. S. (1975). *A Vector Space Model for Automatic Indexing*.
- Saputra, D., Haryani, H., Surniadari, A., Martias, M., & Akbar, F. (2022). Sistem Informasi Bimbingan Tugas Akhir Mahasiswa Berbasis Website Menggunakan Metode Waterfall. *MATRIK : Jurnal Manajemen, Teknik Informatika Dan Rekayasa Komputer*, 21(2), 403–416. <https://doi.org/10.30812/matrik.v21i2.1591>
- Sari, D. R., Windarto, A. P., Hartama, D., & Solikhun, S. (2018). Decision Support System for Thesis Graduation Recommendation Using AHP-TOPSIS Method. *Jurnal Teknologi Dan Sistem Komputer*, 6(1), 1–6. <https://doi.org/10.14710/jtsiskom.6.1.2018.1-6>

- Schiele, Dr.-I. M., Gittmann, Y., Ilchmann, S., Gojsalićgojsalić, A., Jurinčićjurinčić, D., Klempt, P., Gojsalić, A., Jurinčić, D., Zagreb, S., & Ai, C. A. (2024). Voting Advice Applications: Implementation of RAG-supported LLMs. *Institute of Electrical and Electronics Engineers (IEEE)*. <https://doi.org/10.36227/techrxiv.172115156.64500701/v1>
- Syahputra, A., Putri Kinanti, I., & Ilmu Kesehatan, F. (2023). PENYULUHAN KESEHATAN TENTANG KEPERCAYAAN DIRI DENGAN KECEMASAN PADA MAHASISWA KESEHATAN MASYARAKAT SEMESTER VII UNIVERSITAS UBUDYAH INDONESIA YANG AKAN MENGHADAPI SKRIPSI. In *Jurnal Pengabdian Kepada Masyarakat (Kesehatan)* (Vol. 5).
- Tan, H., Luo, Q., Jiang, L., Zhan, Z., Li, J., Zhang, H., & Zhang, Y. (2024). *Prompt-based Code Completion via Multi-Retrieval Augmented Generation*. <http://arxiv.org/abs/2405.07530>
- Trivedi, A., Pant, N., Shah, P., Sonik, S., & Agrawal, S. (2018). Speech to text and text to speech recognition systems-Areview. *IOSR Journal of Computer Engineering*, 20(2), 36–43. <https://doi.org/10.9790/0661-2002013643>
- Tuzzahrah Fatimah. (2024, November 24). *7 Tips Mengatasi Rasa Tegang Jelang Sidang Skripsi, Relaks*. <https://www.idntimes.com/life/education/fatimah-tuzzahrah/mengatasi-tegang-jelang-skripsi-c1c2>
- Universitas Multimedia Nusantara. (2022, October 8). *Kiat Sukses Sidang Skripsi, Auto Dipuji Dosen Penguji | Universitas Multimedia Nusantara*. <https://www.umn.ac.id/kiat-sukses-sidang-skripsi-auto-dipuji-dosen-penguji/>
- USLU, İ. B., Tora, H., & Karamehmet, T. (2017). IMPLEMENTATION OF TURKISH TEXT-TO-SPEECH SYNTHESIS ON A VOICE SYNTHESIZER CARD WITH PROSODIC FEATURES. *ANADOLU UNIVERSITY JOURNAL OF SCIENCE AND TECHNOLOGY A - Applied Sciences and Engineering*, 1–1. <https://doi.org/10.18038/aubtda.283172>
- Vinnarasu, A., & Jose, D. V. (2019). Speech to text conversion and summarization for effective understanding and documentation. *International Journal of*

- Electrical and Computer Engineering*, 9(5), 3642–3648.
<https://doi.org/10.11591/ijece.v9i5.pp3642-3648>
- Wang, J., & Dong, Y. (2020). Measurement of text similarity: A survey. In *Information (Switzerland)* (Vol. 11, Issue 9, pp. 1–17). MDPI AG.
<https://doi.org/10.3390/info11090421>
- Wang, J., Wang, C., Tan, C., Huang, J., & Gao, M. (2024). Knowledgeable In-Context Tuning: Exploring and Exploiting Factual Knowledge for In-Context Learning. *Findings of the Association for Computational Linguistics: NAACL 2024*, 3261–3280. <http://arxiv.org/abs/2309.14771>
- Wang, L., Ma, C., Feng, X., Zhang, Z., Yang, H., Zhang, J., Chen, Z., Tang, J., Chen, X., Lin, Y., Zhao, W. X., Wei, Z., & Wen, J. (2024). A survey on large language model based autonomous agents. In *Frontiers of Computer Science* (Vol. 18, Issue 6). Higher Education Press Limited Company.
<https://doi.org/10.1007/s11704-024-40231-1>
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Ichter, B., Xia, F., Chi, E., Le, Q., & Zhou, D. (2022). *Chain-of-Thought Prompting Elicits Reasoning in Large Language Models*. <http://arxiv.org/abs/2201.11903>
- Yan, L., Sha, L., Zhao, L., Li, Y., Martinez-Maldonado, R., Chen, G., Li, X., Jin, Y., & Gašević, D. (2024). Practical and ethical challenges of large language models in education: A systematic scoping review. In *British Journal of Educational Technology* (Vol. 55, Issue 1, pp. 90–112). John Wiley and Sons Inc. <https://doi.org/10.1111/bjet.13370>
- Yao, S., Yu, D., Zhao, J., Shafran, I., Griffiths, T. L., Cao, Y., & Narasimhan, K. (2023). *Tree of Thoughts: Deliberate Problem Solving with Large Language Models*. <http://arxiv.org/abs/2305.10601>
- Zheng, L., Chiang, W.-L., Sheng, Y., Zhuang, S., Wu, Z., Zhuang, Y., Lin, Z., Li, Z., Li, D., Xing, E. P., Zhang, H., Gonzalez, J. E., & Stoica, I. (2023). Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena. In *Proceedings of the 37th International Conference on Neural Information Processing Systems*, 1–29. <http://arxiv.org/abs/2306.05685>