

DAFTAR PUSTAKA

- Ainslie, J., Lee-Thorp, J., de Jong, M., Zemlyanskiy, Y., Lebrón, F., & Sanghai, S. (2023). GQA: Training generalized multi-query transformer models from multi-head checkpoints. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing* (pp. 4895–4901). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2023.emnlp-main.298>
- Banerjee, S., & Lavie, A. (2005). METEOR: An automatic metric for MT evaluation with improved correlation with human judgments. In *Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization* (pp. 65–72). Association for Computational Linguistics. <https://aclanthology.org/W05-0909>
- Barry, E. S., Merkebu, J., & Varpio, L. (2022). State-of-the-art literature review methodology: A six-step approach for knowledge synthesis. *Perspectives on Medical Education*, 11(5), 281–288. <https://doi.org/10.1007/s40037-022-00725-9>
- Borhan, I., & Bajaj, A. (2024). The effect of prompt types on text summarization performance with large language models. *Journal of Database Management*, 35(1), 1–23. <https://doi.org/10.4018/JDM.358475>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. In *Advances in Neural Information Processing Systems* (Vol. 33, pp. 1877–1901). Curran Associates.
- Chalkidis, I., Fergadiotis, M., Malakasiotis, P., & Aletras, N. (2020). LEGAL-BERT: The muppets straight out of law school. In *Findings of the Association for Computational Linguistics: EMNLP 2020* (pp. 2898–2904). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.findings-emnlp.261>
- Deroy, A., Ghosh, K., & Ghosh, S. (2023). How ready are pre-trained abstractive models and LLMs for legal case judgement summarization? [Preprint]. arXiv. <https://arxiv.org/abs/2306.01248>
- Dougherty, J., Kohavi, R., & Sahami, M. (1995). Supervised and unsupervised discretization of continuous features. In *Machine Learning Proceedings 1995* (pp. 194–202). Morgan Kaufmann. <https://doi.org/10.1016/B978-1-55860-377-6.50032-3>

- García, S., Luengo, J., Sáez, J. A., López, V., & Herrera, F. (2013). A survey of discretization techniques: Taxonomy and empirical analysis in supervised learning. *IEEE Transactions on Knowledge and Data Engineering*, 25(4), 734–750. <https://doi.org/10.1109/TKDE.2012.35>
- Gaveau, D. L. A., Locatelli, B., Salim, M. A., Husnayaen, Manurung, T., Descals, A., Angelsen, A., Meijaard, E., & Sheil, D. (2022). Slowing deforestation in Indonesia follows declining oil palm expansion and lower oil prices. *PLOS ONE*, 17(3), Article e0266178. <https://doi.org/10.1371/journal.pone.0266178>
- Grattafiori, A., Dubey, A., Jauhri, A., Pandey, A., Kadian, A., Al-Dahle, A., Letman, A., Mathur, A., Schelten, A., Vaughan, A., Yang, A., Fan, A., Goyal, A., Hartshorn, A., Yang, A., Mitra, A., Sravankumar, A., Korenev, A., Hinsvark, A., ... Ma, Z. (2024). The Llama 3 herd of models [Preprint]. arXiv. <https://arxiv.org/abs/2407.21783>
- Hemanth Kumar, M., Jayanth, P., & Anand Kumar, M. (2024). Large language models for Indian legal text summarisation. In 2024 IEEE International Conference on Electronics, Computing and Communication Technologies (CONECCT) (pp. 1–5). IEEE. <https://doi.org/10.1109/CONECCT62155.2024.10677065>
- Hwang, W., Eom, S., Lee, H., Park, H. J., & Seo, M. (2022). Data-efficient end-to-end information extraction for statistical legal analysis [Preprint]. arXiv. <https://arxiv.org/abs/2211.01692>
- Itmam, M. S. (2021). Pengantar ilmu hukum. Nusa Litera Inspirasi.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1–38. <https://doi.org/10.1145/3571730>
- Kanapala, A., Pal, S., & Pamula, R. (2019). Text summarization from legal documents: A survey. *Artificial Intelligence Review*, 51(3), 371–402. <https://doi.org/10.1007/s10462-017-9566-2>
- Lin, C.-Y. (2004). ROUGE: A package for automatic evaluation of summaries. In *Text Summarization Branches Out* (pp. 74–81). Association for Computational Linguistics. <https://aclanthology.org/W04-1013>
- Liu, N. F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2023). Lost in the middle: How language models use long contexts [Preprint]. arXiv. <https://arxiv.org/abs/2307.03172>
- Liu, P., Yuan, W., Fu, J., Jiang, Z., Hayashi, H., & Neubig, G. (2021). Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing [Preprint]. arXiv. <https://arxiv.org/abs/2107.13586>

- Moro, G., Magnani, L. D. M., & Ragazzi, L. (2026). Legal lay summarization: Exploring methods and data generation with large language models. *Artificial Intelligence Review*, 59(1), Article 21. <https://doi.org/10.1007/s10462-025-11392-7>
- Nenkova, A., & McKeown, K. (2011). Automatic summarization. *Foundations and Trends in Information Retrieval*, 5(2–3), 103–233. <https://doi.org/10.1561/15000000015>
- Ng, J.-P., & Abrecht, V. (2015). Better summarization evaluation with word embeddings for ROUGE. In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing* (pp. 1925–1930). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D15-1222>
- OpenAI. (2026, March 17). Introducing GPT-5.4 mini and nano. <https://openai.com/index/introducing-gpt-5-4-mini-and-nano/>
- Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C. L., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., Schulman, J., Hilton, J., Kelton, F., Miller, L., Simens, M., Aspell, A., Welinder, P., Christiano, P. F., Leike, J., & Lowe, R. (2022). Training language models to follow instructions with human feedback. In *Advances in Neural Information Processing Systems* (Vol. 35, pp. 27730–27744). Curran Associates.
- Papineni, K., Roukos, S., Ward, T., & Zhu, W.-J. (2002). BLEU: A method for automatic evaluation of machine translation. In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics* (pp. 311–318). Association for Computational Linguistics. <https://doi.org/10.3115/1073083.1073135>
- Preti, D., Giannone, C., Favalli, A., & Romagnoli, R. (2024). Automatic summarization of legal texts, extractive summarization using LLMs. In S. Di Martino, C. Sansone, E. Masciari, S. Rossi, & M. Gravina (Eds.), *Proceedings of the Ital-IA Intelligenza Artificiale – Thematic Workshops co-located with the 4th CINI National Lab AIIS Conference on Artificial Intelligence (Ital-IA 2024)* (pp. 436–440). CEUR Workshop Proceedings. <https://ceur-ws.org/Vol-3762/501.pdf>
- Putri, E. I. K., Dharmawan, A. H., Hospes, O., Yulian, B. E., Amalia, R., Mardiyarningsih, D. I., Kinseng, R. A., Tonny, F., Pramudya, E. P., Rahmadian, F., & Suradiredja, D. Y. (2022). The oil palm governance: Challenges of sustainability policy in Indonesia. *Sustainability*, 14(3), Article 1820. <https://doi.org/10.3390/su14031820>
- Rani Krishna, K. M., Somasundaram, K., Arulmozhivarman, P., Immanuel, S. A., & Rajkumar, E. R. (2025). Deep learning for text summarization using NLP for automated news digest. *Scientific Reports*, 15(1), Article 36343. <https://doi.org/10.1038/s41598-025-20224-1>

- Russell, S. J., & Norvig, P. (2021). Artificial intelligence: A modern approach (4th ed., Global ed.). Pearson.
- Sahoo, P., Singh, A. K., Saha, S., Jain, V., Mondal, S., & Chadha, A. (2025). A systematic survey of prompt engineering in large language models: Techniques and applications [Preprint]. arXiv. <https://arxiv.org/abs/2402.07927>
- Song, Y., Qin, Y., Huang, R., Chen, Y., & Lin, C. (2025). Legal text summarization via judicial syllogism with large language models. *Journal of King Saud University – Computer and Information Sciences*, 37(5), Article 111. <https://doi.org/10.1007/s44443-025-00113-3>
- Sun, J., Shaib, C., & Wallace, B. C. (2023). Evaluating the zero-shot robustness of instruction-tuned language models [Preprint]. arXiv. <https://arxiv.org/abs/2306.11270>
- Swamy, K., Salgare, O., Sakhare, O., & Das, S. (2025). ValidEase: NLP for simplification and summarization of legal documents. *International Journal of Computer Techniques*, 12(2). <https://ijctjournal.org/wp-content/uploads/2025/04/ValidEase-NLP-for-Simplification-and-Summarization-of-Legal-Documents.pdf>
- Tang, Y., Tuncel, D., Koerner, C., & Runkler, T. (2025). The few-shot dilemma: Over-prompting large language models [Preprint]. arXiv. <https://arxiv.org/abs/2509.13196>
- Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M.-A., Lacroix, T., Rozière, B., Goyal, N., Hambro, E., Azhar, F., Rodriguez, A., Joulin, A., Grave, E., & Lample, G. (2023). LLaMA: Open and efficient foundation language models [Preprint]. arXiv. <https://arxiv.org/abs/2302.13971>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (Vol. 30, pp. 5998–6008). Curran Associates.
- Zhang, T., Kishore, V., Wu, F., Weinberger, K. Q., & Artzi, Y. (2020). BERTScore: Evaluating text generation with BERT. In *Proceedings of the 8th International Conference on Learning Representations (ICLR)*. <https://arxiv.org/abs/1904.09675>
- Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., Min, Y., Zhang, B., Zhang, J., Dong, Z., Du, Y., Yang, C., Chen, Y., Chen, Z., Jiang, J., Ren, R., Li, Y., Tang, X., Liu, Z., ... Wen, J.-R. (2026). A survey of large language models [Preprint]. arXiv. <https://arxiv.org/abs/2303.18223>
- Zhong, H., Xiao, C., Tu, C., Zhang, T., Liu, Z., & Sun, M. (2020). How does NLP benefit legal system: A summary of legal artificial intelligence [Preprint]. arXiv. <https://arxiv.org/abs/2004.12158>