

Menulis Akademis Bantuan AI: Mula-mula Diantisipasi Tetapi Sekarang Sudah Tidak Terbendung, Akhirnya Menyisakan Wilayah Abu-abu

Wahyudin Darmalaksana
Kelas Menulis Sentra Publikasi Indonesia
yudi_darma@uinsgd.ac.id

Abstract

The presence of AI in academic writing was initially anticipated, but its progress has been unstoppable, ultimately creating a gray area. This study aims to answer the question of how AI can be utilized in academic writing without violating ethics, security, and academic integrity. This study uses a qualitative approach by applying a literature review method. The results show that AI is capable of producing academic writing independently that can be verified as equivalent to the capabilities of human scientists. This paper has discussed the conflict between advocates of AI ethics and AI safety and AI accelerators, including the AI dilemma and gray areas. The conclusion of this study is that understanding AI-generated text can maintain academic integrity.

Keywords: *Academic Integrity, Artificial Intelligence, Ethics, Understanding.*

Abstrak

Mula-mula kehadiran AI diantisipasi dalam penulisan akademis, tetapi kemajuannya tidak terbendung, akhirnya tercipta wilayah abu-abu. Penelitian ini bertujuan menjawab pertanyaan bagaimana pemanfaatan AI dalam penulisan akademis tanpa pelanggaran terhadap etika, keamanan, dan integritas akademik. Penelitian ini menggunakan pendekatan kualitatif dengan menerapkan metode tinjauan literatur. Hasil penelitian menunjukkan bahwa AI mampu membuat tulisan akademis secara mandiri yang dapat diverifikasi setara kemampuan manusia ilmunan. Tulisan ini telah membahas pertentangan antara penganjur etika AI dan keamanan AI dengan para akselerator AI, termasuk dilema AI, dan wilayah abu-abu. Simpulan penelitian ini adalah pemahaman terhadap teks yang dihasilkan AI dapat menjaga integritas akademik.

Kata Kunci: *Artificial Intelligence, Etika, Integritas Akademik, Pemahaman.*

Pendahuluan

AI terbukti dapat menulis akademis secara mandiri. Dunia akademis mempertanyakan integritas (Ncube, Dzvapatsva, Matobobo, & Ranga, 2026), namun para ahli etika AI kewalahan menyaksikan kemajuan AI dari para akselerator AI (Häggström, 2026). Istilah “Etika AI” dan “Keamanan AI” memiliki batasan yang kabur (Häggström, 2026) sekarang ini.

Hasil-hasil penelitian tentang AI selalu menemukan dilema (Yang, Restrepo-Amariles, Allen, & Troussel, 2026). Di satu sisi, ada kekhawatiran terhadap AI (Gatlin & Middleton, 2026; Yang et al., 2026). Ditemukan bahwa AI menciptakan ketergantungan (Damodaran, Aswani, & Sajith, 2026), meningkatnya disinformasi termasuk *deepfake*, kekhawatiran tentang hak kekayaan intelektual, risiko keamanan siber, bias, dan diskriminasi (Uddin et al., 2026), kesalahpahaman, kepercayaan, dan tanggung jawab (Rismanchian, Babar, & Doroudi, 2026). Di bidang karya ilmiah, peneliti menemukan bahwa AI tidak menunjukkan variasi yang sebanding dengan tulisan manusia (El-Dakhs, Afzaal, & Siyanova-Chanturia, 2026). Namun, di sisi lain, AI diakui sedang mentransformasi fondasi epistemik (Chalmers, Hunt, Pachidi, Potočník, & Townsend, 2026). AI berperan meningkatkan kepercayaan diri (Damodaran et al., 2026), AI telah menguasai pengalaman sehari-hari (Gatlin & Middleton, 2026), terbukti efektif dalam pembuatan konten dan peningkatan produktivitas (Yang et al., 2026), meningkatkan orisinalitas karya dan fleksibilitas (Pérez-Jorge, Pirrone, González-Herrera, Martínez-Murciano, & Raya, 2026), berperan menghasilkan tulisan standar (Xiao, Yuan, Pei, Xue, & Cai, 2026), dan AI mampu mendekati penilaian manusia (Huang, Palermo, & Wilson, 2026).

Dari dilema di atas, tercipta wilayah abu-abu. Suatu hal yang ironis teks sendiri yang ditulis rapi, bukan dari AI, terkena AI detector (Bassett et al., 2026). Tulisan AI dari mesin dihumanisasi tetapi oleh AI (mesin) pula, hal ini berarti AI berusaha memanusiakan manusia (Giray, 2026). AI bisa “berpura-pura” (Battista, Lanciano, Ribatti, & Curci, 2026), bantuan AI ada di wilayah benar dan salah, pengguna ada di dua sisi antara boleh dan tidak boleh (Damodaran et al., 2026; Gatlin & Middleton, 2026; Rismanchian et al., 2026; Uddin et al., 2026; Yang et al., 2026), dan banyak lagi wilayah abu-abu yang samar.

Tulisan ini berusaha menyoroti wilayah abu-abu pemanfaatan AI dalam penulisan akademis, seperti tugas akhir, karya ilmiah, dan artikel jurnal (Hao, Xu, Li, & Evans, 2026; Suchikova, Tsybuliak, Teixeira da Silva, & Nazarovets, 2026). Tulisan ini diharapkan dapat menjernihkan yang abu-abu tersebut tanpa membuat pelanggaran terhadap etika, keamanan, dan tanggung jawab integritas akademik.

Metode Penelitian

Studi ini menggunakan pendekatan kualitatif dengan menerapkan metode tinjauan literatur (Passmore, Olafsson, & Tee, 2026). Sumber data ditelusuri dari artikel-artikel jurnal yang terbit tahun 2026. Pengumpulan data dilakukan dengan teknik dokumentasi. Teknis analisis data mencakup inventarisasi, klasifikasi, dan interpretasi.

Hasil Penelitian

Bukan boleh atau tidak boleh. Seperti telah ditegaskan terdahulu bahwa AI saat ini sedang mentransformasi fondasi epistemik (Chalmers et al., 2026). Revolusi AI saat ini telah menginspirasi beberapa sarjana untuk mengusulkan pembentukan sistem AI yang supercerdas sehingga akan mengungguli dan menggantikan ilmuwan manusia (Spiegel, 2026). Di dalam menulis akademis, tidak terbantahkan bahwa bantuan AI berperan menghasilkan struktur standar (Xiao et al., 2026). Jika sebuah teks tampak terstruktur, maka pembaca mudah menangkap ide teks. Keterbacaan teks diperlukan dalam karya akademis. Penggunaan AI detector mungkin tidak perlu, karena AI detector berbeda dengan alat pengecekan plagiarisme (Bassett et al., 2026). Juga mungkin tidak perlu humanisasi karena berarti mesin diatasi oleh mesin (Giray, 2026). Bukan karena tulisan berbasis AI dapat mendekati penilaian manusia (Huang et al., 2026), melainkan AI terbukti dapat menulis akademis secara independen.

Pembahasan

Integritas akademik sering menjadi sorotan (Ncube et al., 2026). Di sini, perlu agenda literasi AI (Rismanchian et al., 2026) untuk mengurangi pelanggaran akademik (Ncube et al., 2026), dan risiko hukum (Uddin et al., 2026; Yang et al., 2026). Juga manusia terlibat untuk penilaian kreatif (Xiao et al., 2026). Menurut Galtin dan Middleton (2026), sebagian besar yang ditulis tentang AI masih sangat teoritis, karenanya tidak boleh dianggap sebagai pengganti manusia di ruang kelas, terkait kebijaksanaan (Gatlin & Middleton, 2026). Meskipun diakui AI dapat mendekati penilaian manusia (Huang et al., 2026), namun AI tidak menunjukkan variasi yang sebanding dengan tulisan manusia (El-Dakhs et al., 2026). Memang AI meningkatkan, tetapi bukan menggantikan, upaya siswa (Ncube et al., 2026). Keterbatasan dan risiko implementasinya dari pengalaman AI perlu diwaspadai (Pérez-Jorge et al., 2026). Hindari kurangnya kemahiran dalam pemahaman saat AI digunakan (Elim, 2026). Menurut Haggstrom, jika memang ada konflik di bidang ini, akan jauh lebih masuk akal jika para ahli etika AI dan para pendukung keamanan AI bergabung melawan apa yang disebut sebagai para akselerator AI (Häggström, 2026).

Tidak perlu dilematis. Paling utama penulis paham terhadap teks saat AI digunakan (El-Dakhs et al., 2026; Elim, 2026; Gatlin & Middleton, 2026; Huang et al., 2026; Ncube et al., 2026; Pérez-Jorge et al., 2026; Rismanchian et al., 2026; Uddin et al., 2026; Yang et al., 2026). Mula-mula struktur dari AI (Xiao et al., 2026), lalu pemahaman penulis (Rismanchian et al., 2026). Jika tugas akhir, maka pemahaman dilihat saat sidang dan penilaian penguji. Jika konferensi, maka pemahaman, selain prompt, dinilai saat presentasi dan tanggapan juri. Institusi akademik atau pendidikan tinggi dapat saja mengurangi skor penilaian teks tetapi di saat yang sama meningkatkan skor penilaian pemahaman. Konsekuensinya, peserta sidang bisa tidak lulus meskipun tugas akhirnya rapi, namun tidak menguasai pemahaman terhadap teks. Dengan demikian, seimbang antara penggunaan bantuan AI dan penilaian. Supaya tidak abu-abu, maka silakan gunakan AI dengan prompt terbaik (Shah, 2026), tetapi hindari pabrikasi, seperti mengambil data tanpa validasi, sebab hal itu pelanggaran terhadap etika dan integritas akademik, dan pastikan menguasai pemahaman teks (Heimberger, Horvat, & Schultmann, 2026).

Simpulan

AI terbukti mampu membuat tulisan akademis secara mandiri, hal ini dapat diverifikasi setara kemampuan manusia. Etika, keamanan, dan integritas akademik pemanfaatan AI telah dibahas agar pengguna tidak dilematis dan tidak terjebak di dalam wilayah abu-abu. Keterbatasan penelitian ini hanya menelusuri literatur terbitan tahun 2026. Penelitian ini diharapkan memiliki implikasi manfaat bagi para akademisi dalam penulisan akademis di era AI. Penelitian ini memberi peluang terhadap studi di masa depan apakah akan terjadi ledakan AI. Tulisan ini merekomendasikan kepada pengguna AI pada penguasaan pemahaman terhadap bantuan AI dalam menulis akademis, seperti karya ilmiah, tugas akhir, dan artikel jurnal, demi terpeliharanya integritas akademik.

Daftar Pustaka

- Bassett, Mark Andrew, Bradshaw, Wayne, Bornsztejn, Hannah, Hogg, Alyce, Murdoch, Kane, Pearce, Bridget, & Webber, Colin. (2026). Heads we win, tails you lose: AI detectors in education. *Journal of Higher Education Policy and Management*, 1-16.
- Battista, Fabiana, Lanciano, Tiziana, Ribatti, Raffaella Maria, & Curci, Antonietta. (2026). Malingered depression: a comparative study of human and GPT-3.5 performance. *Current Psychology*, 45(4), 379.

- Chalmers, Dominic, Hunt, Richard 'Rick,' Pachidi, Stella, Potočník, Kristina, & Townsend, David. (2026). The acceleration of artificial intelligence: Rethinking organization and work in an era of rapid technological change. *Journal of Management Studies*, 63(2), 285–314.
- Damodaran, Akhil, Aswani, R. S., & Sajith, Shambhu. (2026). Integrating GPT in undergraduate management education and its impact on performance and perception. *Discover Sustainability*.
- El-Dakhs, Dina Abdel Salam, Afzaal, Muhammed, & Siyanova-Chanturia, Anna. (2026). A genre-based comparison of Chat-GPT-generated abstracts versus human-authored abstracts: Focus on applied linguistics research articles. *Corpus Pragmatics*, 10(1), 1.
- Elim, Emily Hui Sein Yue. (2026). Promoting cognitive skills in AI-supported learning environments: the integration of bloom's taxonomy. *Education 3-13*, 54(3), 612–622.
- Gatlin, Melissa Dawn, & Middleton, Kayla D. (2026). My Friend Chat GPT: AI in the Daily Life of an Early Childhood Teacher. *Childhood Education*, 102(2), 62–67.
- Giray, Louie. (2026). AI humanizers and the crisis of information integrity: implications for scientific writing. *Naunyn-Schmiedeberg's Archives of Pharmacology*, 399(8), 11181–11193.
- Häggström, Olle. (2026). On the troubled relation between AI ethics and AI safety. *Contemporary Debates in the Ethics of Artificial Intelligence*, 295–308.
- Hao, Qian Yue, Xu, Fengli, Li, Yong, & Evans, James. (2026). Artificial intelligence tools expand scientists' impact but contract science's focus. *Nature*, 1–7.
- Heimberger, Heidi, Horvat, Djerdj, & Schultmann, Frank. (2026). Exploring the factors driving AI adoption in production: a systematic literature review and future research agenda. *Information Technology and Management*, 27(1), 53–69.
- Huang, Yue, Palermo, Corey, & Wilson, Joshua. (2026). Accuracy and fairness of generative AI in automated essay scoring: Comparing GPT-4o, feature-based models, and human raters. *Assessing Writing*, 69, 101047.
- Ncube, Prince D. N., Dzvapatsva, Godwin P., Matobobo, Courage, & Ranga, Memory M. (2026). Redefining student assessment in AI-infused learning environments: a systematic review of challenges and strategies for academic integrity. *AI and Ethics*, 6(1), 68.
- Passmore, Jonathan, Olafsson, Bergsveinn, & Tee, David. (2026). A systematic literature review of artificial intelligence (AI) in coaching:

- insights for future research and product development. *Journal of Work-Applied Management*, 18(1), 110–129.
- Pérez-Jorge, David, Pirrone, Concetta, González-Herrera, Ana Isabel, Martínez-Murciano, María Carmen, & Raya, Elena Olmos. (2026). Nurturing creative learning through generative AI: A systematic review. *Generative AI in Education: Theories, Applications, and Ethical Frontiers*, 371–396.
- Rismanchian, Sina, Babar, Eesha Tur Razia, & Doroudi, Shayan. (2026). What Undergraduate Students Need to Know and Actually Know About Generative AI. *Computers and Education: Artificial Intelligence*, 100554.
- Shah, Syed Jawad Hussain. (2026). NLP and Prompt Engineering for AI Chatbots. In *Unleashing User Innovation* (pp. 254–268). Auerbach Publications.
- Spiegel, Thomas J. (2026). Can we automate philosophy through AI? And should we want to? *AI and Ethics*, 6(1), 122.
- Suchikova, Yana, Tsybuliak, Natalia, Teixeira da Silva, Jaime A., & Nazarovets, Serhii. (2026). GAIDeT (Generative AI Delegation Taxonomy): A taxonomy for humans to delegate tasks to generative artificial intelligence in scientific research and publishing. *Accountability in Research*, 33(3), 2544331.
- Uddin, Mueen, Arfeen, Shams Ul, Alanazi, Fuhid, Hussain, Saddam, Mazhar, Tehseen, & Arafatur Rahman, Md. (2026). A critical analysis of generative AI: challenges, opportunities, and future research directions. *Archives of Computational Methods in Engineering*, 33(2), 1763–1793.
- Xiao, Nan, Yuan, ChunHong, Pei, YuTing, Xue, Wang, & Cai, Yule. (2026). A study of artificial intelligence in writing assessment for secondary school students: a comparative analysis based on the GPT-4 and human raters. *Educational Studies*, 52(4), 504–526.
- Yang, Cathy, Restrepo-Amariles, David, Allen, Leo Pick, & Troussel, Aurore Clément. (2026). GPT Adoption Dilemma and the Impact of Disclosure Policies. In *Compliance for Artificial Intelligence Systems: Strategies, Principles and Methods* (pp. 96–125). Springer.