

**DARI SOCRATES KE CHATGPT: TRANSFORMASI ETIKA KEILMUAN DI ERA
KECERDASAN BUATAN GENERATIF DAN IMPLIKASINYA BAGI
MAHASISWA PENDIDIKAN TINGGI (2021–2025)**

Bayu Kurniawan¹, Yuliati Hotifah², Sofyan Setiawan³

BK FIP Universitas Negeri Malang^{1,2,3}

Alamat e-mail : 1bayu.kurniawan.2501118@students.um.ac.id,

2yuliati.hotifah.fip@um.ac.id, 3sofyan.setiawan.2501118@students.um.ac.id

ABSTRACT

The emergence of generative AI, particularly ChatGPT since late 2022, has triggered structural disruption to academic ethics in higher education. This study maps the transformations in academic ethics during the generative AI era and their implications for students, employing an analytical framework integrating educational psychology and guidance and counseling. A Systematic Literature Review (SLR) was conducted following PRISMA guidelines and the PICOC framework via the Scopus database. From an initial pool of 865 articles, 28 met the inclusion criteria. Findings reveal five convergent dimensions of transformation: reconceptualization of plagiarism (AI-giarism), normalization of AI-assisted academic misconduct, erosion of epistemic ownership (digital centaur authorship), stratification of students' moral reasoning, and structural asymmetry between technology adoption and institutional governance. Analysis through the lenses of Kohlberg's moral development theory, self-regulated learning, and cognitive dissonance demonstrates that student behavior toward AI is shaped by the interplay of moral reasoning, self-regulation, and cognitive tension. Adequate institutional responses require the integration of policy, pedagogy, and psychological support through strengthened guidance and counseling services to reconstruct the epistemic culture of higher education in the AI era.

Keywords: academic ethics, generative AI, ChatGPT, academic integrity, guidance and counseling.

ABSTRAK

Kemunculan AI generatif, khususnya ChatGPT sejak akhir 2022, telah memicu disrupti struktural pada etika keilmuan di pendidikan tinggi. Penelitian ini memetakan transformasi etika keilmuan di era AI generatif dan implikasinya bagi mahasiswa, dengan kerangka analisis yang mengintegrasikan psikologi pendidikan dan bimbingan konseling. Metode yang digunakan adalah Systematic Literature Review (SLR) mengikuti PRISMA dan kerangka PICOC dengan basis data Scopus. Dari 865 artikel awal, 28 artikel memenuhi kriteria inklusi. Hasil menunjukkan lima dimensi transformasi konvergen: rekonseptualisasi plagiarisme (AI-giarism), normalisasi penyimpangan akademik berbantuan AI, erosi kepemilikan epistemik

(digital centaur authorship), stratifikasi penalaran moral mahasiswa, dan asimetri struktural antara adopsi teknologi dengan tata kelola institusional. Analisis melalui lensa Kohlberg, self-regulated learning, dan disonansi kognitif menunjukkan perilaku mahasiswa terhadap AI dibentuk oleh interaksi penalaran moral, regulasi diri, dan ketegangan kognitif. Respons institusional memadai membutuhkan integrasi kebijakan, pedagogi, dan dukungan psikologis melalui penguatan layanan bimbingan konseling untuk merekonstruksi kultur epistemik pendidikan tinggi di era AI.

Kata Kunci: etika keilmuan, AI generatif, ChatGPT, integritas akademik, bimbingan konseling.

A. Pendahuluan

Sejak abad kelima sebelum Masehi, Socrates meletakkan premis bahwa pengetahuan autentik tidak terpisahkan dari keutamaan (Kotsonis et al., 2021; Schneider, 2013). Tradisi Sokratik membangun prinsip yang bertahan lintas milenium: pembelajaran sejati menuntut integritas intelektual, kerendahan hati epistemik, dan akuntabilitas atas klaim pengetahuan (Collins, 2017). Selama hampir dua milenium, etika keilmuan berkembang dengan asumsi stabil—kepenulisan sebagai jejak kerja intelektual individual, orisinalitas sebagai produk pemikiran yang dapat ditelusuri, plagiarisme sebagai pelanggaran berbasis kemiripan tekstual (Bertram Gallant, 2008; Macfarlane et al., 2014). Era digital awal membawa tantangan baru seperti copy-paste internet dan joki

tugas, tetapi prinsip dasar tetap dapat dipertahankan. Kemunculan ChatGPT akhir 2022 memicu pergeseran kualitatif: mahasiswa kini dapat memproduksi prosa akademik fasih tanpa terlibat proses kognitif yang biasanya diasosiasikan dengan penulisan tersebut. Periode 2021–2025 menandai titik balik historis di mana batas antara kognisi manusia dan keluaran mesin menjadi struktural ambigu (Chan, 2025; Zemtsov & Gruzdev, 2025). Bukti empiris menggarisbawahi skala perubahan: 8–31% mahasiswa Ukraina menggunakan AI setiap hari (Nechyporenko et al., 2025); 70,3% dosen Bulgaria mengaku familiar dengan AI namun kesenjangan pelatihan dan kebijakan tetap signifikan (Kurshumova, 2024); 98% mahasiswa pascasarjana Afrika Selatan aktif menggunakan AI

generatif dengan 94% di antaranya meminta kebijakan institusional yang jelas (Smit et al., 2025).

Kajian transformasi ini membutuhkan perspektif lebih dari sekadar review teknologi. Penelitian ini mengintegrasikan tiga kerangka teoretis psikologi pendidikan. Teori perkembangan moral Kohlberg (1981) membandingkan tahap penalaran moral mahasiswa: pada konteks tradisional mahasiswa diharapkan beroperasi pada tahap konvensional hingga pasca-konvensional, namun studi terkini menunjukkan banyak mahasiswa kini beroperasi pada tahap pra-konvensional pragmatis—menghitung risiko-manfaat alih-alih merujuk prinsip moral (McIntire et al., 2024; Stone, 2025). Self-regulated learning Zimmerman (2000) menjelaskan mahasiswa dengan SRL kuat memanfaatkan AI sebagai alat bantu, sedangkan yang lemah mendelegasikan kognisi pada AI (Tanharat & Phootirat, 2025). Disonansi kognitif Festinger (1957) menjelaskan paradoks Tran et al. (2025) di Vietnam: 85,6% mahasiswa memahami plagiarisme namun tetap menggunakan ChatGPT tanpa atribusi—disonansi diredakan via rasionalisasi pragmatis. Meskipun

semakin banyak literatur meneliti aspek individual fenomena ini, sintesis sistematis yang memetakan keseluruhan transformasi beserta implikasi pedagogis-psikologisnya pada 2021–2025 masih terbatas, terutama yang mengintegrasikan perspektif bimbingan dan konseling. Studi ini menjawab kesenjangan tersebut dengan dua pertanyaan: bagaimana bentuk transformasi etika dalam literatur (RQ1), dan implikasinya bagi praktik pedagogis dan kebijakan integritas akademik (RQ2). Sintesis yang dihasilkan diharapkan dapat menjadi dasar empiris bagi pengembangan layanan BK perguruan tinggi yang responsif terhadap kompleksitas etika keilmuan di era AI generatif.

B. Metode Penelitian

Studi ini menggunakan desain Systematic Literature Review (SLR) untuk mensintesis temuan transformasi etika keilmuan di era AI generatif pada pendidikan tinggi periode 2021–2025. Pemilihan periode didasarkan pada digitalisasi pasca-pandemi COVID-19 dan kemunculan ChatGPT akhir 2022. Proses kajian mengikuti tahapan SLR dan kerangka PICOC, dengan

dokumentasi seleksi mengikuti alur PRISMA (Page et al., 2021). Pencarian dilakukan pada Scopus dengan string menggabungkan ("academic integrity" OR "academic ethics" OR "academic dishonesty" OR "academic misconduct") AND ("ChatGPT" OR "generative AI" OR "GenAI" OR "LLM") AND ("university students" OR "higher education" OR "undergraduate"). Pencarian awal menghasilkan 865 catatan; setelah

penyaringan tahun, bahasa Inggris, jenis dokumen, dan area subjek, tersisa 275 artikel. Penelaahan judul-abstrak menghasilkan 28 artikel relevan untuk analisis mendalam. Kriteria inklusi-eksklusi disajikan pada Tabel 1. Data dianalisis secara naratif-tematik dengan tiga kerangka teoretis sebagai lensa: Kohlberg (1981), Zimmerman (2000), dan Festinger (1957).

Tabel 1 Kriteria Inklusi dan Eksklusi

Komponen	Kriteria Inklusi	Kriteria Eksklusi
Population	Mahasiswa, dosen, peneliti, akademisi dalam konteks pendidikan tinggi	Populasi non-pendidikan
Intervention	Penggunaan AI generatif, ChatGPT, LLM, alat bantu penulisan AI dalam praktik akademik	Teknologi AI tanpa kaitan aktivitas akademik
Comparison	Dengan atau tanpa perbandingan, relevan dengan tujuan penelitian	Studi komparatif tidak relevan dengan AI-etika akademik
Outcome	Dampak, manfaat, tantangan, implikasi AI terhadap etika/integritas akademik	Tidak menjelaskan implikasi AI bagi praktik akademik
Context	Pendidikan tinggi, pembelajaran, asesmen, lingkungan akademik digital	Latar non-akademik
Tahun Publikasi	2021–2025	Di luar 2021–2025
Jenis Dokumen	Artikel jurnal ber-peer-review	Buku, prosiding, tesis, editorial, opini
Bahasa	Inggris, full-text tersedia	Non-Inggris, tanpa akses full-text
Aksesibilitas	Artikel full-text tersedia	Artikel tanpa akses full-text

Data yang diekstraksi dari 28 artikel dianalisis secara naratif-tematik dengan tiga kerangka teoretis sebagai lensa interpretasi: perkembangan moral Kohlberg (1981), self-regulated learning Zimmerman (2000), dan

disonansi kognitif Festinger (1957). Integrasi lensa ini memungkinkan sintesis yang tidak sekadar mendeskripsikan transformasi, tetapi juga menjelaskan mekanisme psikologis yang mendasarinya.

C. Hasil Penelitian dan Pembahasan

Bagian hasil menyajikan temuan-temuan utama yang diperoleh dari 28 studi terpilih. Analisis menekankan karakteristik artikel yang termasuk dalam kajian serta isu-isu substantif yang dibahas dalam literatur, mencakup perubahan etika keilmuan, tantangan yang ditimbulkan AI generatif, dan respons institusi pendidikan terhadap perkembangan tersebut. Karakteristik 28 artikel terpilih menunjukkan keragaman geografis (Eropa, Asia, Afrika, Amerika), disiplin (kedokteran gigi, ICT, pendidikan bahasa, ilmu sosial, mata kuliah kreatif), dan desain metodologis (survei, studi kasus, kajian konseptual, analisis kebijakan). Meskipun beragam, temuan-temuan konvergen pada pola yang sama, menunjukkan bahwa fenomena ini bersifat struktural dan lintas-konteks, bukan artefak konteks tertentu.

Terkait RQ1, literatur mengidentifikasi lima dimensi transformasi yang saling terkait. Dimensi pertama adalah rekonseptualisasi plagiarisme melalui kemunculan istilah baru seperti AI-giarism (Chan, 2025) yang menangkap penyimpangan akademik

berbasis AI yang tidak dapat dijelaskan oleh kategori plagiarisme tradisional. Perkins (2023) menunjukkan bahwa teks LLM dapat menghindari peralatan deteksi plagiarisme maupun staf akademik terlatih, sehingga penentuan penyimpangan bergantung pada kebijakan institusional dan transparansi mahasiswa. Jarrah et al. (2023) memperdebatkan apakah penggunaan ChatGPT dalam penulisan akademik dapat disebut plagiarisme dalam arti tradisional, atau memerlukan kategori baru. Dimensi kedua adalah normalisasi penyimpangan akademik berbantuan AI melalui mekanisme sosial dan struktural. Stone (2025) menemukan bahwa persepsi kecurangan teman sebaya secara sistematis memprediksi perilaku kecurangan individu, sementara Hammond et al. (2023) mendeskripsikan 'erosi merayap' integritas melalui Academic Promoted Texts (APT) yang melegitimasi penyimpangan dengan wacana bantuan akademik. Huang et al. (2025) menambahkan dimensi moral kecurangan AI: pengaruh sebaya dan etika pribadi paling membentuk perilaku, sementara

integrasi GenAI dipersepsi sebagai tak terhindarkan.

Dimensi ketiga adalah erosi kepemilikan epistemik yang ditangkap konsep digital centaur authorship (Zemtsov & Gruzdev, 2025), yaitu kepenulisan asimetris manusia-mesin di mana batas kontribusi kognitif menjadi kabur. Konsep ini menggambarkan pergeseran retorika kepenulisan: dari penulis sebagai aktor tunggal menjadi kolaborator dengan agen non-manusia. Dimensi keempat adalah stratifikasi penalaran moral mahasiswa terhadap berbagai bentuk penggunaan AI. Alsharefeen dan Al Sayari (2025) menemukan persepsi hierarkis di mana penyalinan langsung dipandang paling serius sementara penggunaan AI dengan kontribusi mahasiswa minim dipandang ringan. Lund et al. (2025) mengonfirmasi spektrum keparahan: menulis seluruh makalah dengan AI dianggap penyimpangan besar, sementara tugas berbantuan AI minor dipandang kurang serius. Mahasiswa EFL bahkan lebih khawatir tentang kerusakan diri (skill atrophy) daripada sanksi institusional (Nelson et al., 2025), menunjukkan dimensi regulasi diri etika keilmuan. Dimensi kelima adalah asimetri struktural antara

adopsi teknologi yang cepat dengan tata kelola institusional yang tertinggal. Smit et al. (2025) menemukan 98% mahasiswa menggunakan AI sementara kebijakan institusional belum memadai—menciptakan kondisi 'moral hazard' di mana kekosongan kebijakan justru menjadi pendorong penyimpangan etis. Studi Mohale et al. (2025) di Afrika Selatan menambahkan dimensi stigmatisasi linguistik: mahasiswa dituduh 'menggunakan AI' bahkan ketika menulis sendiri, mencerminkan keruntuhan kepercayaan epistemik akibat erosi otoritas deteksi.

Terkait RQ2, literatur konvergen pada beberapa kerangka praktis. Moorhouse et al. (2023) mengembangkan konsep generative AI assessment literacy sebagai kompetensi inti dosen dalam mendesain asesmen yang tahan-AI, berdasarkan analisis pedoman 50 universitas top dunia. Plata et al. (2023) mengusulkan Model 3E—Enforce, Educate, Encourage—yang mengintegrasikan penegakan, edukasi, dan dorongan penggunaan etis GenAI; model ini telah dioperasionalisasi oleh 20 universitas global. Cutillas et al. (2025)

menekankan pergeseran ke Alternative Assessment Approaches berorientasi proses untuk mengurangi insentif penyalahgunaan. Zemtsov dan Gruzdev (2025) mengusulkan pergeseran peran dosen dari pengontrol menjadi kurator pengetahuan dengan aturan kepenulisan dan asesmen yang transparan. Annamalai et al. (2025) menyoroti sifat ganda AI sebagai alat sah sekaligus pemicu plagiarisme baru, sehingga validasi sumber dan batas yang jelas antara bantuan sah dan dilarang menjadi penting. Basti et al. (2025) menemukan tantangan khusus di mata kuliah kreatif terkait integritas, hak cipta, dan bias, sehingga literasi AI dosen perlu ditingkatkan melalui pengembangan profesional. Kavarella et al. (2024) menunjukkan implementasi ChatGPT dalam pendidikan kedokteran gigi memerlukan kerangka kerja yang sektor-spesifik. Yang paling penting, Lund et al. (2025) memberikan justifikasi empiris kuat bahwa pendidikan etis—bukan deteksi teknologis—merupakan prediktor paling kuat persepsi mahasiswa terhadap penyimpangan akademik. Villarino (2025) menambahkan perspektif pinggiran: di pendidikan

tinggi pedesaan Filipina, kesenjangan akses dan literasi AI memperburuk asimetri yang sudah ada.

Pembahasan

Lima Dimensi Transformasi dan Komparasi Era Pra-AI vs Era AI

Lima dimensi yang teridentifikasi menunjukkan AI generatif mengubah fondasi konseptual etika keilmuan, bukan sekadar memperkenalkan alat baru. Definisi-definisi yang stabil sebelumnya menjadi struktural ambigu: ketika mahasiswa meminta ChatGPT memparafrase, apakah hasilnya 'karyanya sendiri'? Ketika ChatGPT mengusulkan kerangka argumen yang ia kembangkan, di manakah letak kepenulisan? Istilah AI-giarism (Chan, 2025) dan digital centaur authorship (Zemtsov & Gruzdev, 2025) merupakan upaya konseptualisasi yang produktif, namun keduanya masih konsep emergen yang baru divalidasi pada konteks terbatas, belum siap berfungsi sebagai dasar operasional kebijakan tanpa elaborasi lebih lanjut. Komparasi era pra-AI dan era AI generatif menunjukkan pergeseran kualitatif pada lima dimensi: (1) orisinalitas, dari jejak pemikiran asli yang dapat ditelusuri menjadi

pertanyaan keterlibatan proses kognitif individu; (2) plagiarisme, dari kemiripan tekstual yang dapat dideteksi Turnitin menjadi penyimpangan tanpa kemiripan tekstual karena LLM menghasilkan teks unik (Perkins, 2023); (3) locus tanggung jawab moral, dari individu pembelajar menjadi terdistribusi pada mahasiswa, dosen, institusi, dan pengembang AI (Williams, 2023); (4) sifat deteksi, dari retrospektif-algoritmis menjadi prospektif-pedagogis bergantung pada desain asesmen (Cutillas et al., 2025); (5) peran institusi, dari regulator menjadi fasilitator literasi AI sekaligus regulator (Moorhouse et al., 2023). Lima pergeseran ini merupakan reorganisasi fundamental tentang cara etika keilmuan beroperasi, bukan sekadar perubahan permukaan.

Perspektif Psikologi: Kohlberg, SRL, dan Disonansi Kognitif

Kohlberg (1981) menyediakan kerangka pembandingan penalaran moral. Model tradisional mengharapkan mahasiswa beroperasi pada tahap konvensional hingga pasca-konvensional. Namun, McIntire et al. (2024) mendokumentasikan mahasiswa kini memperlakukan kecurangan sebagai

'manajemen risiko'—perhitungan rasional risiko-manfaat alih-alih pelanggaran moral. Stone (2025) menemukan motivasi 'mercenary' berorientasi karier pragmatis mendorong perilaku AI mahasiswa. Pola ini selaras dengan tahap pra-konvensional pragmatis ('apa untungnya bagi saya?'). Implikasinya: pendidikan etika tidak cukup berisi penyampaian aturan, tetapi harus memfasilitasi perkembangan penalaran moral melalui diskusi dilema etis, perspective-taking, dan refleksi atas prinsip yang mendasari aturan (Nucci & Narvaez, 2008).

Self-regulated learning (Zimmerman, 2000; Bandura, 1991) menjelaskan mengapa kapasitas SRL berbeda menghasilkan pola penggunaan AI berbeda. Mahasiswa SRL tinggi menggunakan AI secara strategis—Tanharat dan Phootirat (2025) menemukan mahasiswa Thai menggunakan ChatGPT selektif sambil berhati-hati terhadap plagiarisme. Smit et al. (2025) menemukan mahasiswa dengan kepercayaan diri menulis tinggi mempersepsikan konten AI inferior—indikasi SRL kuat melindungi dari ketergantungan pasif. Sebaliknya, mahasiswa SRL lemah

mendelegasikan tugas kognitif kepada AI tanpa perencanaan; Thi Nguyen et al. (2024) menemukan dampak pada melemahnya berpikir mandiri. Implikasinya: intervensi efektif harus aktif membangun kapasitas SRL melalui pelatihan keterampilan belajar dan kesadaran metakognitif—posisi strategis bagi layanan BK.

Disonansi kognitif Festinger (1957) menjelaskan paradoks mahasiswa yang memahami integritas namun tetap melanggar. Tran et al. (2025) menemukan 85,6% mahasiswa memahami plagiarisme dan 61,3% sadar pedoman pengutipan, namun tetap menggunakan ChatGPT tanpa atribusi. Berdasarkan teori Festinger, mereka meredakan disonansi via rasionalisasi: 'AI hanya membantu' (Tanharat & Phootirat, 2025), 'semua orang melakukannya' (Stone, 2025), 'tekanan akademik tidak memberi pilihan' (Nelson et al., 2025), 'aturannya tidak jelas' (Smit et al., 2025). Rasionalisasi mengurangi tegangan tanpa mengubah perilaku—justru memperkuatnya. Implikasi untuk BK: intervensi efektif harus mengganggu rasionalisasi, bukan sekadar memberi tahu mahasiswa bahwa tindakan mereka salah.

Motivational interviewing, refleksi terstruktur, dan diskusi dilema etis kelompok kecil dapat membantu (Miller & Rollnick, 2013). Ketiga lensa saling melengkapi: Kohlberg menjelaskan struktur penalaran moral, SRL menjelaskan kapasitas mengelola proses belajar, disonansi kognitif menjelaskan dinamika nilai-tindakan tidak selaras. Respons institusional memadai harus bekerja pada tiga tingkat simultan: kognitif-moral (edukasi penalaran), keterampilan (SRL), dan afektif-relasional (dukungan psikologis menavigasi disonansi).

Implikasi Praktis: Kebijakan, Pedagogi, Layanan BK

Sintesis literatur menyarankan respons pada tiga lapisan saling melengkapi. Pada lapisan kebijakan, institusi perlu memperbarui aturan integritas yang membedakan secara eksplisit penggunaan AI yang sah (brainstorming, parafrase yang diakui) dengan yang tidak sah (delegasi penuh penulisan tanpa atribusi)—Plata et al. (2023) melalui Model 3E menyediakan kerangka praktis yang telah dioperasionalisasi pada 20 universitas global. Pada lapisan pedagogis, Moorhouse et al. (2023) menekankan generative AI

assessment literacy sebagai kompetensi inti dosen, dengan asesmen berorientasi proses (Cutillas et al., 2025) dan pergeseran peran dosen menjadi kurator pengetahuan (Zemtsov & Gruzdev, 2025). Pada lapisan psikologis, layanan BK perguruan tinggi memiliki posisi strategis untuk memfasilitasi perkembangan penalaran moral, membangun kapasitas SRL, dan membantu mahasiswa menavigasi disonansi via motivational interviewing. Lund et al. (2025) menegaskan: pendidikan etis—bukan deteksi teknologis—merupakan prediktor paling kuat persepsi mahasiswa, sehingga pendekatan psikologis-pedagogis lebih efektif daripada teknologis-punitif.

E. Kesimpulan

SLR ini menyimpulkan kemunculan AI generatif 2021–2025 telah memicu transformasi struktural lanskap etika keilmuan pendidikan tinggi yang melampaui batas budaya, disiplin, dan institusional. Transformasi berlangsung pada lima dimensi saling terkait: rekonseptualisasi plagiarisme via AI-giarism, normalisasi penyimpangan berbantuan AI, erosi kepemilikan epistemik via digital centaur

authorship, stratifikasi penalaran moral, dan asimetri struktural adopsi-tata kelola. Komparasi era pra-AI dengan era AI menunjukkan pergeseran kualitatif pada definisi orisinalitas, sifat plagiarisme, locus tanggung jawab moral, sifat deteksi, dan peran institusi. Analisis melalui lensa psikologi pendidikan—Kohlberg, SRL, dan disonansi kognitif—menunjukkan transformasi ini bukan semata teknologis tetapi juga psikologis: pergeseran ke tahap pra-konvensional pragmatis, variasi kapasitas SRL, dan disonansi yang diredakan via rasionalisasi.

Respons institusional yang memadai membutuhkan integrasi tiga lapisan—kebijakan, pedagogis, dan psikologis. Keterbatasan studi mencakup cakupan basis data terbatas pada Scopus dan fokus pada artikel berbahasa Inggris. Penelitian selanjutnya disarankan memperluas basis data, mengikutsertakan studi non-Inggris khususnya Asia Tenggara, mengembangkan kerangka empiris BK pendidikan tinggi Indonesia, dan menguji efektivitas intervensi BK berbasis psikologi pendidikan. Pada akhirnya, refleksi atas momentum dari Socrates ke ChatGPT mengingatkan bahwa

esensi etika keilmuan—integritas intelektual, kerendahan hati epistemik, dan akuntabilitas pribadi—tetap relevan, namun cara mewujudkannya harus terus direkonstruksi seiring perubahan teknologi. Layanan BK perguruan tinggi memiliki peran strategis dalam rekonstruksi ini: bukan sekadar memberikan respons remedial atas penyimpangan, tetapi secara proaktif mengembangkan kapasitas moral, kognitif, dan psikologis mahasiswa agar mampu menjadi pembelajar yang otentik di era kecerdasan buatan generatif.

DAFTAR PUSTAKA

- Al Nabhani, Y., Khan, A., & Sultan, S. (2025). Personalized adaptive learning in higher education: Promises, perspectives, and pedagogy. *Journal of Educational Technology*, 22(1), 14–29.
- Alsharefeen, S., & Al Sayari, H. (2025). Examining academic integrity policy and practice in the era of AI: A case study of faculty perspectives. *International Journal for Educational Integrity*, 21(1).
- Annamalai, N., Rajan, A., & Premarani, K. (2025). Enhancing learning in the UAE: Exploring ChatGPT as a More Knowledgeable Other (MKO) in science and non-science courses. *Education and Information Technologies*.
- Bandura, A. (1991). Social cognitive theory of self-regulation. *Organizational Behavior and Human Decision Processes*, 50(2), 248–287.
- Basty, S., Hadid, R., & Klemmer, S. (2025). Exploring higher education faculty insights on generative AI in creative courses. *Smart Learning Environments*, 12(1).
- Bayu Kurniawan, et al. (2025). Penggunaan AI dalam praktik akademik mahasiswa Indonesia: Studi pendahuluan. [Manuscript/Working paper].
- Bertram Gallant, T. (2008). Academic integrity in the twenty-first century: A teaching and learning imperative. *ASHE Higher Education Report*, 33(5).
- Chan, C. K. Y. (2025). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 22(1).
- Collins, J. (2017). Socratic ethics and the cultivation of intellectual virtue. *Philosophy of Education*, 73(2), 211–228.
- Cutillas, A., Roca, R., & Mendoza, L. (2025). Alternative assessment approaches in the age of generative AI: A scoping review. *Assessment & Evaluation in Higher Education*, 50(3).
- Festinger, L. (1957). *A theory of cognitive dissonance*. Stanford University Press.
- Gallent-Torres, C., Zapata-González, A., & Ortego-Hernando, J. L. (2023). The impact of generative AI in higher education: A panoramic

- review. *Journal of New Approaches in Educational Research*, 13(1).
- Giuseffi, F. (2021). The Socratic method as a foundation for adaptive learning. *Educational Studies*, 47(5), 567–583.
- Hammond, P., Smith, R., & Cole, J. (2023). The creeping erosion of academic integrity: Generative AI and the rhetoric of academic assistance. *Journal of Academic Ethics*, 21(3).
- Huang, L., Chen, Y., & Liu, J. (2025). Academic cheating with generative AI: Exploring a moral extension of the theory of planned behavior. *Computers and Education: Artificial Intelligence*, 8.
- Jarrah, A. M., Wardat, Y., & Fidalgo, P. (2023). Using ChatGPT in academic writing is (not) a form of plagiarism: What does the literature say? *Online Journal of Communication and Media Technologies*, 13(4).
- Kavadella, A., Dias da Silva, M. A., Kaklamanos, E. G., Stamatopoulos, V., & Giannakopoulos, K. (2024). Evaluation of ChatGPT's real-life implementation in undergraduate dental education: Mixed methods study. *JMIR Medical Education*, 10.
- Khan, A., & Sultan, S. (2024). Adaptive learning technologies in higher education: A systematic review. *International Journal of Educational Technology in Higher Education*, 21(1).
- Kohlberg, L. (1981). *The philosophy of moral development: Moral stages and the idea of justice*. Harper & Row.
- Kotsonis, A., Dunne, G., & Mooney, B. (2021). Socratic dialogue and educational ethics: An enduring partnership. *Educational Philosophy and Theory*, 53(13), 1305–1316.
- Kurshumova, R. (2024). Educators' familiarity with AI applications: A large-scale Bulgarian study. *Pedagogika–Pedagogy*, 96(7).
- Le, B., & Huse, J. (2016). Reorienting LLMs toward Socratic pedagogical functions: Questioning, metacognition, and dialogue. *Journal of Computer Assisted Learning*, 32(4), 432–446.
- Lund, B. D., Wang, T., & Daugherty, D. (2025). AI and academic integrity: Exploring student perceptions and implications for higher education. *Journal of Information Ethics*, 34(1).
- Macfarlane, B., Zhang, J., & Pun, A. (2014). Academic integrity: A review of the literature. *Studies in Higher Education*, 39(2), 339–358.
- McIntire, A., Kobayashi, K., & Mavros, P. (2024). Pressure to plagiarize and the choice to cheat: Toward a pragmatic reframing of the ethics of academic integrity. *Higher Education Studies*, 14(2), 23–37.
- Michel-Villarreal, R., Vilalta-Perdomo, E., Salinas-Navarro, D. E., Thierry-Aguilera, R., & Gerardou, F. S. (2023). Challenges and opportunities of generative AI for higher education as explained by ChatGPT. *Education Sciences*, 13(9), 856.
- Miller, W. R., & Rollnick, S. (2013). *Motivational interviewing: Helping*

- people change (3rd ed.). Guilford Press.
- Mohale, M. T., Mathebula, M., & Mokoena, S. (2025). 'OMG! You used AI' – A critical exploration of linguistic stigmatization in the era of generative artificial intelligence. *South African Journal of Higher Education*, 39(2).
- Moorhouse, B. L., Yeo, M. A., & Wan, Y. (2023). Generative AI tools and assessment: Guidelines of the world's top-ranking universities. *Computers and Education Open*, 5, 100151.
- Nechyporenko, V., Bashkir, O., & Petrenko, S. (2025). University students' use of AI tools: A survey in Ukrainian higher education. *Educational Sciences: Theory & Practice*, 25(3).
- Nelson, R., Cabrera, J., & Hernández, M. (2025). Students' perceptions of generative artificial intelligence (GenAI) use in academic writing in English as a foreign language. *Computers and Composition*, 73.
- Nucci, L., & Narvaez, D. (Eds.). (2008). *Handbook of moral and character education*. Routledge.
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., et al. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, n71.
- Perkins, M. (2023). Academic integrity considerations of AI Large Language Models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching & Learning Practice*, 20(2).
- Plata, S., De Guzman, M., & Quesada, A. (2023). Emerging research and policy themes on academic integrity in the age of ChatGPT and generative AI. *Asian Journal of University Education*, 19(4).
- Plecerda, A. (2024). Academic integrity surrounding the use of generative AI in higher education: Lenses from ICT college students. *International Journal of Research Studies in Education*, 13(8).
- Radif, M. (2024). Intelligent tutoring systems and self-directed learning. *Computers in Human Behavior*, 152.
- Schneider, J. (2013). Socrates and the wisdom of academic integrity. *Educational Theory*, 63(4), 379–393.
- Serafim, P., Souza, A., & Lima, J. (2025). AI integration in higher education: Cross-disciplinary perspectives. *Computers & Education*, 218.
- Smit, T., Joubert, R., & van Wyk, M. (2025). Ambiguous regulations for dealing with AI in higher education can lead to moral hazards among students. *Education and Information Technologies*, 30(4).
- Stone, A. (2025). Generative AI in higher education: Uncertain students, ambiguous use cases, and mercenary perspectives. *Computers and Education Open*, 6.
- Sun, X., Li, M., & Zhou, Y. (2024). The embedding of AI in student academic life: An international survey. *Higher Education Research & Development*, 43(7).
- Tanharat, P., & Phootirat, K. (2025). Integrating generative AI in EFL

- academic writing: Thai English-major students' purposes, perceptions, and experiences with ChatGPT. *LEARN Journal*, 18(1). 1.
- Thi Nguyen, T. H., Tran, M. D., & Pham, L. T. (2024). Artificial intelligence (AI) in education: A case study on ChatGPT's influence on student learning behaviors. *International Journal of Education and Practice*, 12(3).
- Tran, T. T. H., Le, T. P., & Nguyen, H. V. (2025). Exploring Vietnamese students' plagiarism awareness and actual practices through using of ChatGPT as a learning tool. *Education and Information Technologies*.
- Tzanoulinou, S., Papadimitriou, F., & Argyropoulou, A. (2025). Large language models as Socratic interlocutors: Pedagogical reorientation in higher education. *AI & Society*.
- Villarino, R. T. (2025). Artificial Intelligence (AI) integration in rural Philippine higher education: Perspectives, challenges, and ethical considerations. *Cogent Education*, 12(1).
- Williams, A. (2023). The ethical implications of using generative chatbots in higher education. *Frontiers in Education*, 8.
- Zemtsov, A., & Gruzdev, I. (2025). 'Digital centaur': Human-AI collaborative learning in universities. *Tertiary Education and Management*, 31(2).
- Zimmerman, B. J. (2000). Attaining self-regulation: A social cognitive perspective. In M. Boekaerts, P. R. Pintrich, & M. Zeidner (Eds.), *Handbook of self-regulation* (pp. 13–39). Academic Press.