

ACADEMIC PERMISSIBILITY IN LLM-ASSISTED EFL WRITING: HOW CONTEXTUAL JUSTIFICATIONS SHAPE STUDENTS' MORAL JUDGMENTS IN A WITHIN-SUBJECTS VIGNETTE SURVEY

(Research article)

Mohamed Seddiki ^{a *}, Souhila Korichi ^b

^{a,b} Amar Telidji University, AILE Laboratory, Department of English, Laghouat, Algeria

Received: 27.03.2026

Revised version received: 30.06.2026

Accepted: 03.07.2026

Abstract

This paper investigates EFL students' perceptions of ethical acceptability judgments of using Large Language Models (LLMs) into academic writing. In contrast to the simplistic view of acceptable/unacceptable use of LLMs, the present study models how specific contextual justifications shape students' moral evaluations of LLM-assisted writing. A cross-sectional within-subject design was used with 220 third-year EFL students at three public universities in Laghouat, Algeria, to rate the ethical acceptability of LLMs use to complete four writing assignments. Under a neutral baseline condition, participants assessed the use of LLM for completing the four assignments, as well as four single condition contexts (disclosure, accuracy verification, syllabus permission, and learning intent). Their ratings were analyzed using Δ -effect scores (conditional minus baseline) to quantify condition effects above baseline and to describe task-level differences in ethical acceptability. Contextual conditions increased ethical acceptability to different extents, with learning intent producing the largest positive shift ($\Delta M \approx 1.6$, $d \approx 0.7$) and disclosure exerting only a small, non-robust effect ($\Delta M \approx 0.2$, $d \approx 0.1$). The baseline ratings also followed a clear gradient, which saw AI-assisted proofreading as the most acceptable while paraphrasing was consistently least acceptable. Theoretically, the study adds value by supporting a conditional-ethics perspective because it demonstrates how students' moral considerations of LLMs use are dependent on learning-oriented and epistemically responsible frameworks rather than procedural cues like disclosure or syllabus permissions with no behavioral consequences. Practically, the paper advises that AI policies and pedagogy should be designed to encourage learning intent and verification activities as opposed to disclosure as a standalone requirement.

Keywords: Academic integrity; large language models (LLMs); EFL writing; conditional ethics

© 2021 IJETS. Published by *International Journal of Education Technology and Science (IJETS)*. Copyright for this article is granted to the Journal. This is an open-access article distributed under the terms and conditions of the Creative Commons Attribution license (CC BY-NC-ND) (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

* Corresponding author: Mohamed Seddiki. ORCID ID.: <https://orcid.org/0009-0001-3922-1535>

E-mail: med.seddiki@lagh-univ.dz

DOI: <https://doi.org/10.5281/zenodo.22230040>

1. Introduction

1.1. *The problem and the purpose of the study*

The adoption of Generative Artificial Intelligence (GenAI) in higher education has outpaced the development of responsive policies and practices to incorporate it into teaching and learning. Empirical work reveals that the majority of university students are using GenAI tools for academic tasks, including writing, with implications for pedagogy and assessment in higher education contexts (Stephenson & Armstrong, 2026; Digital Education Council, 2024). For second-language (L2) writers, this technological shift is especially consequential: linguistic challenges increase the attractiveness of LLM-assisted writing (Yao, 2026), even as institutional guidance on ethically acceptable use remains fragmented, inconsistent, or sometimes, no guidance at all (Azevedo et al., 2026; Aslan, 2026). In this respect, GenAI-mediated academic writing directly shapes the design of equitable, integrity-preserving learning environments and critically problematizes the very notion of “quality education” for AI-rich higher education contexts (United Nations, 2019).

The main problem that this article is trying to solve is that most of higher-education policy and practice relating to courses of instruction define the incorporation of large language models in academic writing as an either/or proposition. Such binary framings are increasingly at odds with how students actually interact with AI-mediated writing tools, where judgments about permissibility are nuanced and highly contingent on particular affordances (e.g. outlining vs paraphrasing), motivations (e.g. learning vs shortcutting), and institutional contexts (e.g. permission, disclosure requirements) (Cotton et al., 2024; Liang et al., 2025; Mo & Crosthwaite, 2025). The persistence of this “all-or-nothing” discourse has critical implications for policy and pedagogy. Higher education institutions are revising academic integrity codes, implementing disclosure mandates, and updating course syllabi without empirical evidence about how specific educational-technology practices such as transparency expectations, verification of AI-generated content, institutional permission, and learning-oriented use shape students’ ethical evaluations and likely compliance (Plata et al., 2023; Jin et al., 2025).

In contexts like Algeria, where EFL students already utilize LLM-based tools extensively for course assignments despite a lack of comprehensive institutional frameworks for AI-enabled education, learners are forced to independently negotiate the ethics of AI-mediated writing, often in high-stakes assessment situations and with limited alignment between institutional messages, course-level expectations, and available educational technologies (Boukhelkhal, 2025). For teacher-educators, program leaders, and policymakers, this raises urgent questions about how to design AI-aware writing curricula and institutional policies that are both pedagogically meaningful and realistically aligned with learners’ existing technology practices.

Existing research on GenAI in higher education has begun to map the parameters that make AI-mediated academic writing more or less acceptable to students and instructors. One stream of research focuses on contextual factors, including the disclosure requirements (Cleland et al.,

2026; Resnik & Hosseini, 2026; Spirgi et al., 2024), accuracy verification (Danyaro et al., 2025; Mustaffa et al., 2025; Martín-Moncunill & Alonso Martínez, 2025), and institutional AI policies (Ahmed et al., 2026; Saad & Zolkifli, 2026; McDonald et al., 2025), as key drivers of academic integrity in technology-mediated assessments environments. Another body of literature highlights educational intent, revealing that students make a clear ethical distinction between using AI as a scaffold for understanding, feedback, or revision versus using it to bypass cognitive effort (Fathi & Rahimi, 2026; Anani et al., 2025; Yao, 2026). The writing process also emerges as an integrity assessment factor in studies, demonstrating that stakeholders evaluate AI assistance differently across tasks. In this context, the least ethically questionable tasks are proofreading and editing whereas paraphrasing and content creation are associated with higher risks of authorship dilution and plagiarism (Cheng et al., 2025; Malik et al., 2023; Spirgi et al., 2024). Taken together, these strands point to the need for models that link concrete educational-technology practices to students' moral evaluations in specific teaching and learning contexts. Yet three critical gaps persist. First, no framework has operationalized conditional ethics to quantify how specific variables shift perceptions of acceptability. Second, no within-subjects design has systematically crossed task types with underlying justificatory conditions. Third, empirical evidence from the Global South contexts remains conspicuously absent, especially from those navigating institutional policy vacuums.

The current study tackles these issues by conceptualizing the ethical acceptability of LLM-assisted academic writing as a conditional, task-dependent judgment, examined within a fully crossed 4 (writing task: outlining, summarizing, paraphrasing, proofreading) $\times$ 4 (contextual condition: disclosure, accuracy verification, syllabus permission, learning intent) within-subjects vignette design. Focusing on third-year undergraduate EFL students from three Algerian public institutions, we model the influence of specific pedagogical justifications and institutional signals embedded in AI-mediated writing scenarios on students' ethical evaluations as compared to an unconditional baseline. For each Task $\times$ Condition combination, we estimate Δ scores that capture intra-individual changes in ethical acceptability perform repeated measures analysis to assess both the relative impact of each Condition and the baseline differences between Tasks.

2. Literature Review and Hypotheses Development

2.1. Theoretical Anchors

The present study is based on two different conceptualizations of students' moral judgment about utilizing AI. First, according to Rest's (1986) four-component model, moral judgment is differentiated from moral sensitivity, motivation, and character. Specifically, moral judgment is defined as the individual's ability to make a decision about what is right in a given situation. In the present study, ethical acceptability ratings for different forms of LLM assisted writing are treated as direct expressions of such moral judgments. Most recently, Rest's model has been utilized in studies concerned with ethical issues related to using AI. For instance, Masrek et al.

(2026) found that higher ethical decision-making scores predict more responsible and “astute” engagement with AI tools among university students. Noddings’ Care Ethics (2013) builds upon this by stating that each of these judgements has relational and learning-focused intentions. That is, what matters ethically is whether LLM use supports the learner’s intellectual development and the integrity of the educational relationship, rather than whether AI is used in the abstract. Together, these perspectives provide a conditional ethics lens for the present study, in which task- and condition-specific changes in acceptability (Δ scores) can be seen as task- and condition-dependent modifications to their ethical judgement of AI-assisted academic writing.

2.2. LLMs Use in Academic Writing

The rapid rise of LLMs such as ChatGPT and Gemini has seen academic writing involving AI become commonplace (Meyer et al., 2023; Bao et al., 2025). This has triggered debates about academic integrity and responsible use of LLMs in higher education (Kasneci et al., 2023; Cotton et al., 2024). Initial studies on the ethics of AI-generated writing tended to adopt a binary view of the issue, in which using LLMs for writing constitutes academic misconduct or it does not, with undisclosed AI authorship taken as the core violation (Cotton et al., 2024; Cleland et al., 2026), with low-risk applications (e.g., grammatical correction) being implicitly tolerated (Allen & Mizumoto, 2026). Meanwhile, more contemporary research highlights that this binary is not informative because students’ ethical judgments are contingent upon the tasks they request and the circumstances in which they do so (Rowland, 2023; Mo & Crosthwaite, 2025). Qualitative evidence identifies recurring benefits (efficiency, language support) and risks (plagiarism, over-dependence) but offer little guidance on how moral judgments might vary depending upon specific tasks and circumstances (Liang et al., 2025). The present study responds to this gap by explicitly conceptualizing ethical acceptability as conditional on task type and justification, and by operationalizing these condition-sensitive judgments through within-subject Δ scores.

2.3. Conditions That Shift Ethical Acceptability

Research has shown that there are specific contextual justifications that can transform morally ambiguous AI use into ethically acceptable practice. The following four have been selected as experimental conditions in the present study on the basis of their theoretical salience and empirical prevalence in the literature.

Disclosure. Disclosure, explicit acknowledgment of the use of AI, is the most discussed prerequisite for ethically acceptable use of LLMs in academic writing, with institutional policies and expert guidance alike emphasizing transparency as a foundational principle of honesty, authorship, and research integrity (Cleland et al., 2026; Resnik & Hosseini, 2026). It has been found through survey-based study that non-disclosure is the defining feature of AI-assisted academic dishonesty: the key ethical issue is not the support provided but the misrepresentation of authorship (Spirgi et al., 2024). A scoping review similarly argues that integrity concerns in GenAI-mediated assessment are increasingly framed around transparency, attribution, and

declared responsibility, rather than the mere presence of AI support (Pérez-Pérez et al., 2026). In this regard, transparency emerges as a critical ethical consideration, which transforms “a covert act (cheating)” into an open act of “assisted productivity,” and, consequently, increases the perceived level of integrity. By contrast, the structural barriers to such transparency may severely limit honest disclosures. For example, students may refrain from declaring the use of AI due to institutional policy ambiguity (Smit et al., 2025), fears of self-incrimination (Gonsalves, 2025), and concerns about grade reduction (Spirgi et al., 2024). This normative practical tension motivates **H1**: the disclosure condition (C1) will produce a significant positive Δ Effect above baseline across writing tasks.

Accuracy Verification. Accuracy verification is particularly crucial for summarizing and other content-generation tasks, where hallucination (the production of plausible but factually incorrect or incomplete content) constitutes a primary epistemic risk for knowledge reliability and academic credibility (Danyaro et al., 2025; Virvou et al., 2025). Higher-education instructors conditionally endorse AI assistance only when students actively check the correctness, completeness, and sourcing of generated text (Lien et al., 2025). At the same time, Al Murshidi et al., 2026 show that AI-powered proofreading and summarizing tools can improve students’ performance; however, relying on unverified information may worsen their academic skills. In practice, however, most students’ verification behavior falls short of this standard (Martín-Moncunill & Alonso Martínez, 2025). Active verification, from the point of view of current conditional ethics, is a substantive ethical justification, a signal of epistemic responsibility, and not a mere procedural step. This leads to **H2**: the accuracy verification condition (C2) will produce a significant positive Δ Effect above baseline.

Syllabus Permission. Institutional policies are another important consideration in students’ ethical assessments of AI usage. In survey research for different higher-education contexts, Saad and Zolkifli (2026) show that students particularly want clear AI-use guidelines and experience policy ambiguity is a source of anxiety and ethical confusion. This result agrees with the study of Ahmed et al. (2026). The authors emphasize the need for ethical guidelines and institutional support to steer students toward responsible AI use. Policy analyzes show that universities are increasingly publishing governance level AI statements, but concrete course-level expectations are often delegated to individual instructors, who are expected to specify permissible use in syllabi with limited training or resources (McDonald et al., 2025; Azevedo et al., 2026; Tong et al., 2025). This discrepancy can be resolved by moving away from the rhetoric of abstract integrity and toward clear frameworks that specify the acceptable use of tasks on a case-by-case basis (Kadwa, 2025; Amini et al., 2026). Therefore, the explicit syllabus permission functions as a norm setting tool that legitimizes some usage patterns and reduces the uncertainty, rendering the AI assistance within the scope of the defined bounds morally acceptable instead of covert. These findings motivate **H3**: The syllabus permission condition (C3) will produce a stronger positive Δ effect than disclosure (C1). Its relative magnitude compared with accuracy verification (C2) and learning intent (C4) will be examined through secondary comparisons,

reflecting the ethical importance students attach to explicit institutional guidance on permissible AI use.

Learning Intent. Learning intent has proved to be one of the most powerful ethical justifications in both quantitative and qualitative research under all possible conditions. Recent studies reliably demonstrate that students are more willing to use AI tools for activities they consider to be enhancing their learning (e.g., clarification of concepts, feedback, practice of genre appropriate writing) than for ones that simply avoid work or disguise authorship (Cho et al., 2026; Cui, 2025; Nelson et al., 2025; Karani & Mwancha, 2025). This dichotomy reflects the relationship between in-depth and superficial approaches to learning. Undergraduates utilize AI as a cognitive guide to gain an in-depth understanding of complex theories while acknowledging its potential to erode their critical thinking and originality (Anani et al., 2025; Karani & Mwancha, 2025). Fathi & Rahimi (2024) characterize GenAI as an effective scaffold when it mediates planning, feedback, and revision when integrated into dialogic frameworks, but its use for copy-pasting undermines students' agency, creativity, and strategic use of resources. Therefore, when interpreted through the lens of Care Ethics, the decision to use GenAI is evaluated based on one's intention to benefit from the technology (Noddings, 2013). These converging findings motivate **H4**: the learning intent condition (C4) will produce the largest positive Δ Effect among all four conditions, significantly increasing ethical acceptability ratings above baseline.

2.4. Writing Task Types as Contextual Backdrop

Research on perceptions of AI assistance in academic writing indicates that students and teachers often make nuanced ethical distinctions between different kinds of writing tasks and view specific stages of the writing processes more favorably than others (Cheng et al., 2025; Spirgi et al., 2024). AI-assisted proofreading and grammar checking are widely regarded as the least ethically contentious form of support (Malik et al., 2023). Outlining and summarizing occupy an intermediate zone, where students are generally comfortable with AI-generated structures but exercise greater scrutiny once AI output starts shaping content (Utami et al., 2023; Rosalina et al., 2026; Barrett & Pack, 2023). AI-assisted paraphrasing, by contrast, attracts the most cautious evaluations, as students worry about loss of originality, difficulty distinguishing AI-reworded text from their own ideas, and the risk of covert substitution of authorial voice (Malik et al., 2023; Subaveerapandiyana et al., 2025; Xu et al., 2026). These task-specific baseline differences provide important contextual variation against which the four conditions in the present study are tested, though no a priori directional predictions about the baseline gradient are made here.

2.5. Research Gaps and Study Rationale

Three gaps in the existing literature directly motivate the present study.

- **Gap 1 (Theoretical):** No published framework operationalizes conditional ethical acceptability — the degree to which specific contextual conditions transform a morally

ambiguous action into an ethically acceptable one. Existing research confirms that studies identify risks and benefits but fail to model the conditions under which ethical judgments change.

- **Gap 2 (Empirical):** No study has crossed all four writing tasks (outlining, summarizing, paraphrasing, proofreading) with all four conditions (disclosure, accuracy verification, syllabus permission, learning intent) within a single within-subjects instrument. The present study eliminates between-sample confounds and yields the first full 4×4 Task $\times$ Condition matrix of Δ Effects.
- **Gap 3 (Contextual):** Quantitative research on students' LLM ethics perceptions is overwhelmingly concentrated in Global North contexts. Algeria — with its combination of EFL learners, accelerating AI adoption, and near-total institutional policy vacuum — represents a critical yet unexamined context in which no baseline ethical acceptability data exist.

Guided by Rest's moral-judgment framework and Noddings' care-ethics perspective, these gaps motivate a design that tests whether distinct contextual conditions produce systematically different shifts in students' ethical acceptability judgments, and whether those shifts follow theoretically meaningful patterns across H1–H4. To address these gaps, the present study examined Algerian EFL university students' ethical acceptability judgments of LLM use in academic writing through one primary research question and four directional hypotheses (**H1–H4**) as follows:

- To what extent do the four contextual conditions (disclosure, accuracy verification, syllabus permission, and learning intent) produce a significant positive shift (Δ Effect) in ethical acceptability ratings above the baseline across academic writing tasks?

3. Method

3.1. Research Design

This study employed a quantitative survey design based on a cross-sectional, fully within-subjects framework. Each participant evaluated the ethical acceptability of LLM integration across a 4 (Task: outlining, summarizing, paraphrasing, proofreading) $\times$ 4 (Condition: disclosure, accuracy verification, syllabus permission, learning intent) factorial framework. For each task, participants' levels of ethical acceptability were estimated in two successively completed vignettes that involved either an unconditional and a conditional scenario. This design resulted in the formation of 4×4 Δ -score matrix per participant to operationalize the **Conditional Ethics Δ Effect**, which reflected participants' conditional ethics scores in comparison to their unconditional ethics score within the same task. Because all measures were collected from the same individuals under all conditions, between-subject confounds are eliminated and statistical power is maximized for detecting condition-level shifts. While this

fully within-subjects vignette design is well suited to isolating condition-specific shifts in ethical acceptability, respondents' judgments concern hypothetical behaviors rather than actual conduct in the real world.

3.2. Participants and Sampling

A non-probability, criterion-based convenience sample of $N = 220$ third-year undergraduate EFL students was recruited across three public higher education institutions in Laghouat Province, Algeria: Amar Telidji University (42.7%, $n = 94$), the École Normale Supérieure de Laghouat — ENSL (33.2%, $n = 73$), and the University Center of Aflou — UCA (24.1%, $n = 53$). Eligibility was restricted to third-year undergraduate EFL students enrolled in an advanced academic writing curriculum and exposed to institutional academic-integrity expectations. Recruitment across the three institutions provided some institutional variation within a relatively homogeneous EFL disciplinary context; however, the sample should not be considered statistically representative of all Algerian EFL students. A post-hoc sensitivity analysis confirmed that, given $\alpha = .05$ and power = .80, a sample of $N = 220$ provides sufficient statistical power to detect small-to-medium effect sizes in the planned one-sample t-tests and repeated-measures ANOVAs (Cohen, 2009).

3.3. Ethical Considerations and Voluntary Participation

At the time of data collection, a formal Institutional Review Board (IRB) or dedicated Research Ethics Committee for non-medical social science research was unavailable at the three participating institutions. To guarantee the participants' safety, the research adhered to ethical standards and criteria established by international educational research assessment agencies, including the British Educational Research Association's (BERA, 2024). Particularly, the requirement of the study's design involved the research being non-sensitive, thus only collecting general data on participants' views on language learning. Participation was entirely voluntary, uncompensated, and strictly anonymous, with no collection of personally identifiable information (PII) or network identifiers. Prior to accessing the questionnaire, participants were presented with a mandatory digital informed consent page that details the study's purpose and data privacy measures. An explicit consent was required via a mandatory checkbox before the survey items could be accessed, and participants retained the right to withdraw at any time by closing the window.

3.4. Instrument

The survey instrument comprised five sections:

Section A - Demographic Profile. Six items collected: gender, age, university, English proficiency level, self-reported frequency of LLM use, and the number of months the participant has you been using LLMs in academic writing.

Sections B– E- Ethical Acceptability Ratings. Four parallel sections addressed outlining (B), summarizing (C), paraphrasing (D), and proofreading (E). Within each section, participants responded to:

- **One baseline item (BASE):** " Rate how ethically acceptable you believe the use of LLMs (ChatGPT, Gemini...) is in [task], without any qualifying conditions?" (0–10 visual analogue scale anchored at Completely Unacceptable and Completely Acceptable).
- **Four conditional items** rated on the same scale, each introducing a single qualifying condition:
 - **C1 - Disclosure:** "... IF the student discloses its use to the instructor."
 - **C2 - Accuracy Verification:** "... IF the student verifies the generated written output for accuracy before submitting."
 - **C3 - Syllabus Permission:** "... IF the course syllabus explicitly permits the use of LLMs for this task."
 - **C4 - Learning Intent:** "... IF the student uses LLMs as a learning aid rather than a replacement for their own thinking."

3.4.1. Item sequence and randomization

Within each task section, participants first rated an unconditional baseline item followed by the four conditional items in the same order: Disclosure, Accuracy Verification, Syllabus Permission, and Learning Intent. Similarly, all respondents saw the four tasks (outlining, summarizing, paraphrasing, and proofreading) in the same order. No item-level randomization or counterbalancing was implemented. To reduce potential carry-over effects, the introduction to the task explicitly stated that participants should treat each prompt individually and base their ratings on the question as presented as opposed to prior questions.

3.4.2. Pilot Study and Instrument Validation

Prior to piloting, the survey was subjected to expert review by three university academics (two in language didactics and one in educational statistics) to confirm content validity and item alignment. The instrument was then pilot tested among an independent sample of $N = 26$ third-year undergraduate EFL students. A post-survey face-to-face debriefing session with a purposive sub-sample of eight participants reported no meaningful challenges during item comprehension, only minor syntactic revisions were made to eliminate ambiguity. Item-level analysis did not reveal any floor or ceiling effects, and corrected item-total correlations were above .40 for all items. Scale reliability was exceptionally high: Cronbach's alpha coefficients ranged from $\alpha = .873$ (Outlining conditional scale) to $\alpha = .973$ (Proofreading conditional scale), with the complete 16-item conditional index achieving an overall $\alpha = .972$ (George & Mallery, 2016). Preliminary baseline descriptive statistics within the pilot sample directionally supported

the task-based variation (Proofreading: $M = 7.27$, $SD = 3.35$; Paraphrasing: $M = 4.92$, $SD = 3.75$). As a result, all original items were retained for the data collection stage.

Taken together, the expert review and pilot debriefing provide content and face validity evidence that the vignette scenarios and rating items were clear and captured the intended constructs of ethical acceptability. Similarly, the high internal consistency coefficients and strong corrected item–total correlations for each of the four conditional scales lend support to the internal-structure validity of the vignette rating items.

3.5. Data collection

The described study was conducted using Google Forms as an online survey tool, which was made available to the target population during the 2026 academic term. A standardized introductory screen defined "LLM use" for participants, provided a brief neutral description of the four writing tasks, and clarified the difference between baseline and conditional rating formats. Participants were advised to review each item on its own merits rather than adjust their responses based on preceding items to avoid influencing one another. The average time required to complete the survey was between 6 and 9 minutes, and no compensation was offered. Additionally, the anonymity was guaranteed, and participants were free to withdraw at any point without consequence.

3.6. Data Analysis

All analyses were conducted in IBM SPSS Statistics. The analytic sequence corresponded to RQ1 and H1–H4 as follows:

H1–H4 — Condition Δ Effects. Four composite Δ variables were computed by averaging each condition's Δ score across the four writing tasks (D_COND_DISC, D_COND_VER, D_COND_SYL, D_COND_INT). Four one-sample t-tests were then conducted — one per composite Δ score — to test whether each mean Δ differed significantly from zero ($\mu_0 = 0$), confirming that a given condition raised ethical acceptability above baseline. Effect sizes were quantified as Cohen's $d = M/SD$ (one-sample formulation). Regarding **H3**, the planned comparative expectation ($\Delta_{\text{SYL}} > \Delta_{\text{DISC}}$) was evaluated using a paired-samples t-test. Comparisons of Δ_{SYL} with Δ_{VER} and Δ_{INT} were retained as secondary comparisons between composite Δ scores. A Bonferroni-corrected α of .0125 (.05/4) was applied across the four one-sample tests.

In addition, for each of the four writing tasks, the baseline ethical acceptability ratings are descriptive to put the condition effects in context, as there was no a priori hypothesis about these effects.

Note: Variable name conventions used throughout this paper: suffix **_BASE** indicates an unconditional baseline rating; prefix **D_COND_** is a composite Δ score for a particular condition (DISC

= Disclosure, VER = Accuracy Verification, SYL = Syllabus Permission, INT = Learning Intent). Abbreviations for the tasks are OUT = Outlining, SUM = Summarizing, PAR = Paraphrasing, PRF = Proofreading.

4. Results

This section reports findings for RQ1 and H1–H4 in sequence. All Δ scores represent the difference between conditional and baseline ethical acceptability ratings (Conditional – Baseline) on an 11-point scale (0–10). One-sample t-tests (Bonferroni-corrected $\alpha = .0125$ across four tests) were used to evaluate the positive condition effects, while paired-samples t-tests (Bonferroni-corrected $\alpha = .017$) were used for the planned H3 contrast and secondary comparisons among conditions. All means and standard deviations refer to Δ scores unless stated otherwise.

4.1. Demographic characteristics of participants

A total of 220 third-year EFL undergraduates were involved in this research. The sample consisted of 33 male students (15%) and 187 female students (85%). Most participants were between 18 and 20 years old ($n = 119$, 54.1%), followed by those aged 21–22 ($n = 78$, 35.5%), 23–25 ($n = 17$, 7.7%), and 26 or older ($n = 6$, 2.7%). Students were drawn from three universities in Algeria, namely Laghouat University ($n = 94$, 42.7%), ENSL ($n = 73$, 33.2%), and UCA ($n = 53$, 24.1%). In terms of self-reported English proficiency, 34 students (15.5%) were beginners, 145 (65.9%) were intermediate, and 41 (18.6%) were advanced. Regarding LLM use for academic writing, 10 participants reported using LLMs rarely (4.5%), 117 sometimes (53.2%), 65 often (29.5%), and 28 always (12.7%). Consistent with this, 30 students reported 1–3 months of LLM experience (13.6%), 58 reported 4–6 months (26.4%), 39 reported 7–12 months (17.7%), and 93 reported more than 12 months of use (42.3%).

4.2. Condition Effects on Ethical Acceptability (H1–H4)

To answer RQ1, four composite Δ variables were computed by averaging each condition's Δ score across the four writing tasks (D_COND_DISC, D_COND_VER, D_COND_SYL, D_COND_INT). Then one-sample t-tests were conducted to compare the mean of each Δ variable with a test value of 0 to determine if each condition increased the ethical acceptability.

Table 1. One-Sample t-Test Results for Condition Δ Effects (N = 220)

Condition	Mean Δ	SD	95% CI	$t(219)$	p (one-tailed)	Cohen's d	Decision
C1: Disclosure	0.21	1.63	[−0.00, 0.43]	1.93	.027	0.13	H1: Not supported
C2: Accuracy Verification	0.64	1.71	[0.41, 0.86]	5.53	<.001	0.37	H2: Supported
C3: Syllabus Permission	0.82	2.19	[0.52, 1.13]	5.57	<.001	0.37	See H3 below
C4: Learning Intent	1.62	2.22	[1.33, 1.92]	10.82	<.001	0.73	H4: Supported

Note: Bonferroni-adjusted $\alpha = .0125$ (four simultaneous tests). One-tailed p -values reported (directional hypotheses). Cohen's d = Mean Δ / SD.

4.2.1. H1 – Disclosure Condition

The Disclosure condition produced a small, positive mean Δ ($M = 0.21$, $SD = 1.63$), $t(219) = 1.93$, $p_{\text{one-tailed}} = .027$, 95% CI [−0.00, 0.43], $d = 0.13$. The direction of the effect was consistent with H1, but the p value did not reach the Bonferroni adjusted threshold of .0125 and the confidence interval just included zero. **H1 was therefore not supported.** This indicates that disclosure alone exerted only a weak and unreliable effect on ethical acceptability.

4.2.2. H2 – Accuracy Verification Condition

The Accuracy Verification condition produced a moderate and statistically significant positive Δ ($M = 0.64$, $SD = 1.71$), $t(219) = 5.53$, $p_{\text{one-tailed}} < .001$, 95% CI [0.41, 0.86], $d = 0.37$. The confidence interval was entirely above zero, providing strong evidence that accuracy verification caused students to report higher ethical acceptability than the control condition. Thus, **H2 was supported.**

4.2.3. H3 – Syllabus Permission (Comparative Claim)

H3 predicted that the Syllabus Permission condition would produce a larger Δ than Disclosure. In Table 1, the one-sample test first confirmed that the Syllabus Δ was significantly positive ($M = 0.82$, $SD = 2.19$), $t(219) = 5.57$, $p_{\text{one-tailed}} < .001$, 95% CI [0.52, 1.13], $d = 0.37$. Three paired-samples t -tests then tested the planned Syllabus–Disclosure contrast and the two secondary comparisons (Bonferroni-adjusted $\alpha = .017$). As shown in Table 2, Syllabus permission produced a significantly larger Δ than Disclosure (mean difference = 0.61, $t(219) = 4.59$, $p < .001$), but did not differ significantly from Accuracy Verification (mean difference = 0.19, $t(219) = 1.46$, $p = .146$). In the secondary comparison, Learning Intent yielded a significantly larger Δ than Syllabus Permission (mean difference = −0.80, $t(219) = −5.97$, p

< .001). Thus, **H3 was supported**: the Syllabus Permission condition was significantly better than Disclosure. But the secondary comparisons showed that Syllabus Permission was no different than Accuracy Verification and was not the most potent condition overall, with Learning Intent being the leading contextual justification.

Table 2. Paired-Samples t-Test Results: Syllabus Permission vs. Other Conditions

Comparison	Mean Difference	SD	<i>t</i> (219)	<i>p</i> (two-tailed)	Decision
Syllabus – Disclosure	+0.61	1.97	4.59	.001	Syllabus > Disclosure ✓
Syllabus – Verification	+0.19	1.88	1.46	.146	Not significant X
Syllabus – Learning Intent	–0.80	1.98	–5.97	<.001	Syllabus < Learning Intent X

4.2.4. H4– Learning Intent Condition

The condition of the Learning Intent produced the largest Δ Effect among all four conditions ($M = 1.62$, $SD = 2.22$), $t(219) = 10.82$, $p_{\text{one-tailed}} < .001$, 95% CI [1.33, 1.92], $d = 0.73$. This result indicates that framing LLM use as explicitly oriented toward learning substantially elevated students' ethical acceptability ratings relative to baseline. **H4 was strongly supported**.

4.3. Descriptive Overview of Baseline Task Ratings

Table 3 provides descriptive statistics for baseline ethical acceptability ratings for the four writing tasks, to provide context for the condition effects reported above. Estimated marginal means showed Proofreading to be the highest ($M = 6.79$), followed by Summarizing ($M = 5.89$) and Outlining ($M = 5.61$), with Paraphrasing being the lowest ($M = 5.20$).

Table 3. Baseline ethical acceptability ratings of LLM use across writing tasks

Writing task	M	SD
Outlining	5.60	2.95
Summarizing	5.89	2.87
Paraphrasing	5.20	2.84
Proofreading	6.79	2.86

5. Discussion

This study explored how contextual factors influence students' moral perceptions of using LLMs to complete writing tasks. The findings are based on vignette-based ratings and thus reflect students reported moral judgments of hypothetical situations, rather than their actual use of LLMs or adherence to institutional policies. The findings demonstrate that student moral evaluations were not randomly varying but systematically responded to differences in the rationalizations of the proposed use of AI and the demands of the writing task. Regarding RQ1, the four contextual conditions produced distinct, hierarchical shifts above baseline (Learning Intent > Syllabus Permission $\approx$ Accuracy Verification >> Disclosure), directly demonstrating that the type of justification offered for LLM use constitutes a primary driver of students' conditional moral judgments. These results have implications for academic integrity policy and writing pedagogy in that they demonstrate what types of AI use students feel are ethically appropriate versus unacceptable.

5.1. *The Disclosure Paradox: Why Transparency Alone Is Insufficient*

The current study finds that the effect of disclosure is weak and unreliable (small positive Δ , $d = 0.13$, non-significant after correction) which is a stark contrast to normative frameworks that position transparency as a central precondition for ethical AI use in academic writing (Cleland et al., 2026; Resnik & Hosseini, 2026). This aligns with empirical evidence demonstrating that policy awareness fails to meaningfully influence student ethical judgments or actual AI use behaviors (Lund et al., 2025). Rather than being a motivating factor, disclosure regimes tend to have the opposite effect, as students are incentivized to conceal information under the assumption that any statement will be used against them (Gonsalves, 2025; Pérez-Pérez et al., 2026). These results extend current literature by experimentally showing that transparent, explicit disclosure only leads to a trivial increase in acceptability, with transparency functioning as a weak auxiliary cue rather than the decisive ethical lever assumed by institutional policy discourse.

5.2. *Accuracy Verification as a Credible Ethical Anchor*

Accuracy verification produced a significant positive shift in ethical acceptability ($d = 0.37$). This reveals that student moral frameworks prioritize epistemic accountability over mere procedural transparency. This pattern converges with findings which suggest that students value fact checking, revision, and personalization of AI generated texts as key components of using AI responsibly, and a means by which they can maintain authorship integrity and avoid plagiarism. (Sabbaghan & Eaton, 2025; Massoud & Zhang, 2025; Al Murshidi et al., 2026). Given that LLM hallucinations and fabricated references pose severe threats to scholarly integrity (Mustaffa et al., 2025), manual cross-checking against authoritative sources acts as a practical expression of ethical responsibility. Conversely, a lack of deliberate verification

frequently results in the uncritical submission of factually flawed outputs (Martín-Moncunill & Alonso Martínez, 2025). When viewed through the lens of ethical decision-making theory (Rest, 1986), accuracy verification functions as a form of epistemic accountability at the moral-judgment stage. In this regard, students appear to accept LLM assistance when they retain critical control over content and assume responsibility for its truthfulness (Zhu et al., 2025; Ahmed et al., 2026).

5.3. Syllabus Permission: Institutional Authority and Its Limits

Syllabus permission (C3) yielded a significantly larger Δ shift than disclosure (C1). However, it did not differ from accuracy verification (C2) and was surpassed by learning intent (C4). At the level of moral judgment, this reveals the limits of institutional authority as an ethical source. Accordingly, students appear to treat syllabus permission as a necessary but insufficient condition for ethical acceptability, one that legitimizes LLM use procedurally without resolving the deeper question of whether such use is epistemically sound or genuinely learning-oriented. This outcome reinforces arguments that higher education must move beyond simplistic binaries of prohibition or permission (Kadwa, 2025). The strategies for declaration and permission are ineffective in ensuring academic integrity because they do not specify task constraints and evaluate behavioral compliance (Nguyen, 2025; Amini et al., 2026). This is further exacerbated by the fact that students are more likely to be bound by the rules when they are facing deadlines or when they believe that they will benefit from their actions (Ahmed et al., 2026). Therefore, beliefs about ethics and task usefulness are more critical success factors than awareness of rules about AI usage (Lund et al., 2025). These results support H3 and indicate that policy changes alone are insufficient to enhance adoption rates unless they are combined with epistemic responsibility and learning.

5.4. Learning Intent as the Dominant Ethical Justification

Learning intent (C4) produced the strongest ethical premium across all tasks ($d = 0.73$), highlighting the importance of developing framing as a criterion of moral acceptability. This suggests that students think about the use of AI in terms of its functional capacity to promote certain outcomes, not an intrinsic moral value (Tekir, 2026; Vendrell & Johnston, 2026). At L2, users embrace the opportunity to leverage LLM for concept clarification, grammatical accuracy, and textual cohesion, while at the same time rejecting the possibility to resort to tools that undermine their motivation and autonomy (Fathi & Rahimi, 2026; Anani et al., 2025). This aligns with the evidence that full textual outsourcing erodes cognitive ownership and compromises authorial voice (Cao & Abdullah, 2025; Jose et al., 2025). Consequently, LLM integration is judged as morally justifiable only when it promotes intellectual growth rather than short-circuiting the writing assignments (Yao, 2026). Framed through Care Ethics perspective, the core moral consideration is not the technological intervention itself, but whether its

application sustains or hollows out the learner's epistemic engagement with the task (Noddings, 2013).

Synthesizing the results for RQ1, it emerges that student ethical approval tends to be conditioned by particular contextual factors in a hierarchical manner. Indeed, procedural transparency through simple disclosure (C1) proved ineffective, whereas factual verification (C2) and formal syllabus permission (C3) have positive effects, albeit to different degrees. Nevertheless, the key finding is the criticality of the learning intention condition (C4), which ultimately suggests that the ethical framework prioritizes the protection of cognitive sovereignty

5.5. Task Type as Contextual Backdrop for Ethical Evaluation

This gradient from the most mechanically-oriented task (proofreading, $M = 6.79$) to the most content-oriented one (paraphrasing, $M = 5.20$) is consonant with the broader literature on students nuanced ethical distinctions across writing stages (Cheng et al., 2025; Spirgi et al., 2024), and was used as the context for interpreting the four condition effects described above. As a form-focused activity, proofreading does not interfere with conceptual argumentation. Accordingly, leveraging LLMs for surface accuracy automates mechanical editing while fully preserving the author's intellectual input (Al Sawi & Alaa, 2024; Cheng et al., 2025; Wang, 2025). Students and instructors are consistently most comfortable when AI functions as a language-level assistant rather than a content generator, viewing the tool as an editorial aid that complements, rather than replaces, human cognitive effort (Barrett & Pack, 2023; Malik et al., 2023; Kim et al., 2025). AI-assisted paraphrasing, by contrast, sits at the intersection of text comprehension, language production, and identity. Although students value AI rewording for vocabulary acquisition and plagiarism avoidance (Alammar & Amin, 2023; Triana et al., 2025), severe ethical ambivalence persists. In environments that lack explicit institutional parameters, automated paraphrasing blurs the boundary between legitimate editing and misconduct, introducing risks of accidental plagiarism, diminished critical thinking, and the covert substitution of authorial voice (Hammond et al., 2023; Chen et al., 2025; Nik Kamarudin et al., 2025; Rosalina et al., 2026). This positioning of paraphrasing at the bottom of the ethical spectrum reflects a distinct psychological dynamic within L2 writing pedagogy. These EFL writers view the linguistic manipulation of text as a highly sensitive ethical boundary. This aligns with Pecorari's (2010) foundational account of EFL students' acute anxiety over text transformation and patchwriting, and with documented struggles with lexical choice and source-text synthesis in the Algerian higher education context (Haireche & Lamri, 2025).

5.6. Theoretical Contributions

This study makes two distinct theoretical contributions to the literature on AI ethics in higher education. First, it provides experimental, within-subjects evidence that moral judgments about LLMs are hierarchically ordered across justificatory contexts (Learning Intent > Syllabus Permission $\approx$ Accuracy Verification $\gg$ Disclosure). This hierarchy directly challenges policy

models that treat transparency, institutional permission, and epistemic responsibility as interchangeable safeguards. Framed within Rest's (1986) four-component model of ethical decision-making, these ratings operationalize context-sensitive moral judgment, demonstrating that justifications which foreground the epistemic and developmental consequences of AI use (Learning Intent, Accuracy Verification) elicit stronger moral responses than procedural cues that can be satisfied without deep evaluative engagement (Disclosure). Such findings resonate with Zhu et al.'s (2025) separation of ethical awareness and perceived risk, and with Ahmed et al.'s (2026) account of ethical dissonance between judgment and behavior. Second, the dominant effect of Learning Intent, read alongside the negligible effect of Disclosure, provides controlled experimental evidence that students' moral frameworks are governed by Care Ethics rather than rule-compliance norms. In line with Noddings's (2013) relational framework, these L2 writers assign ethical weight to LLMs based on their orientation toward intellectual development rather than institutional declarations. While policy frameworks treat simple declaration as a sufficient safeguard, the disclosure paradox observed here suggests that such procedural compliance measures do not promote responsible AI use when disconnected from students' deeper, care-centered conceptions of legitimate learning.

6. Conclusions and Limitations

The study results suggest that higher education policies and writing pedagogies should move beyond stand-alone disclosure requirements and give greater weight to learning intent and verification practices when framing acceptable LLM use. Instructors and program leaders should position LLMs as tools for scaffolding students' own thinking rather than replacing it, and writing tasks should require learners to document how AI-generated output was checked, revised, and integrated into their work. Finally, the institutions should provide guidelines clarifying the permissible use of such technology in specific tasks while preserving intellectual authorship, epistemic responsibility, and learning.

This study has three primary limitations that offer pathways for future inquiry. First, the data rely entirely on self-reported ethical acceptability ratings, capturing students' stated moral judgments rather than their actual AI integration behaviors or underlying intentions. Future studies may consider utilizing behavioral data or think-aloud protocols to capture the phenomenon of interest in more detail. Second, the sample is restricted to a homogeneous sample of third-year undergraduate EFL students within a localized Algerian territory. These findings may not generalize to early-stage undergraduates, postgraduate cohorts, or non-linguistic disciplines with differing AI proficiencies and ethical standards. Future investigations should replicate this 4×4 factorial design across multi-institutional and cross-cultural samples to evaluate the stability of the observed Δ -Effect patterns. Third, this study focuses on the judgment aspect of Rest's (1986) model. Longitudinal and mixed-methods designs are necessary to map how moral recognition, intention, and behavior interact over time, revealing the exact conditions under which care-centered moral frameworks translate into sustained,

responsible academic practice among L2 writers. Finally, all baseline and conditional items were presented in a fixed task and condition order, without randomization or counterbalancing. Although the vignette prompts emphasized independent judgments for each item, this design may have introduced order or carry-over effects that future studies could address through randomized or counterbalanced presentation formats. Future research should move beyond vignette-based judgment studies by combining ethical acceptability measures with behavioral evidence of actual LLM use, randomized survey designs, and cross-context replications that test whether the present hierarchy of conditions remains stable across disciplines, institutions, and national settings.

Acknowledgements

The authors gratefully acknowledge the valuable contributions of the participants and institutions involved in this study.

Declaration of Conflicting Interests and Ethics

The authors declare no conflict of interest.

Data Availability Statement

Data will be obtained from the corresponding author upon reasonable request.

AI Use Statement

It is confirmed that the conceptual framework, theoretical arguments, and analysis of the study data of this manuscript are fully generated by the authors. A generative AI tool was utilized only for the purposes of grammatical corrections and fluency of the text body.

References

- Ahmed, T., Rocky, K. H., Rahman, S., & Hasan, M. J. (2026). Ethical use of ChatGPT in higher education: Insights from business students on responsible AI adoption. *Ethics & Behavior*. Advance online publication. <https://doi.org/10.1080/10508422.2026.2615333>
- Al Murshidi, G., Alfahid, M., & Al Zaabi, A. (2026). How do college students weigh the reliability of AI in academic writing? *Cogent Education*, 13(1), Article 2670040. <https://doi.org/10.1080/2331186X.2026.2670040>
- Al Sawi, I., & Alaa, A. (2024). Navigating the impact: A study of editors' and proofreaders' perceptions of AI tools in editing and proofreading. *Discover Artificial Intelligence*, 4(1), Article 23. <https://doi.org/10.1007/s44163-024-00116-5>

- Alammar, A., & Amin, E. A.-R. (2023). EFL students' perception of using AI paraphrasing tools in English language research projects. *Arab World English Journal*, 14(3), 166–181. <https://doi.org/10.24093/awej/vol14no3.11>
- Allen, T. J., & Mizumoto, A. (2026). ChatGPT over my friends: Japanese English-as-a-foreign-language learners' preferences for editing and proofreading strategies. *RELC Journal*, 57(1), 89–106. <https://doi.org/10.1177/00336882241262533>
- Amini, M., Lee, K.-F., Wang, Y., & Ravindran, L. (2026). Proposing a framework for ethical use of AI in academic writing based on a conceptual review: Implications for quality education. *Interactive Learning Environments*, 34(3), 1394–1418. <https://doi.org/10.1080/10494820.2025.2523382>
- Anani, G. E., Nyamekye, E., & Bafour-Koduah, D. (2025). Using artificial intelligence for academic writing in higher education: The perspectives of university students in Ghana. *Discover Education*, 4(1), Article 46. <https://doi.org/10.1007/s44217-025-00434-5>
- Aslan, A. (2026). Predicting EFL learners' writing proficiency through reading and writing practices: An artificial neural networks approach. *International Journal of Education, Technology and Science*, 6(2), 137–159. <https://doi.org/10.5281/zenodo.20473891>
- Azevedo, L., Mallinson, D. J., Wang, J., Robles, P., & Best, E. (2026). AI policies, equity, and morality and the implications for faculty in higher education. *Public Integrity*, 28(2), 186–201. <https://doi.org/10.1080/10999922.2024.2414957>
- Bao, T., Zhao, Y., Mao, J., & Zhang, C. (2025). Examining linguistic shifts in academic writing before and after the launch of ChatGPT: A study on preprint papers. *Scientometrics*, 130(7), 3597–3627. <https://doi.org/10.1007/s11192-025-05341-y>
- Barrett, A., & Pack, A. (2023). Not quite eye to AI: Student and teacher perspectives on the use of generative artificial intelligence in the writing process. *International Journal of Educational Technology in Higher Education*, 20(1), Article 59. <https://doi.org/10.1186/s41239-023-00427-0>
- British Educational Research Association. (2024). *Ethical guidelines for educational research* (5th ed.). <https://www.bera.ac.uk/publication/ethical-guidelines-for-educational-research-fifth-edition-2024-online>
- Boukhelkhal, O. (2025). Algerian EFL students' perspectives and practices on AI-enabled learning at the University of Medea. *Journal of Studies in Language, Culture and Society*, 8(1), 93–110.
- Cao, L., & Abdullah, A. (2025). EBA (engaged but amotivated) in AI-enhanced EFL learning: A qualitative study from a Chinese higher vocational context. *Frontiers in Psychology*, 16, Article 1643653. <https://doi.org/10.3389/fpsyg.2025.1643653>
- Chen, X., Xavier, L. G., Pires, M., & Amaro, V. (2025). Artificial Intelligence and Foreign Language Learning: ChatGPT's Limitations in Summarizing Texts. *Theory and Practice in Language Studies*, 15(11), 3585–3594. <https://doi.org/10.17507/tpls.1511.11>

- Cheng, A., Calhoun, A., & Reedy, G. (2025). Artificial intelligence-assisted academic writing: Recommendations for ethical use. *Advances in Simulation*, 10(1), Article 22. <https://doi.org/10.1186/s41077-025-00350-6>
- Cho, M.-H., Park, E. G., Lim, S., & Yu, H. (2026). Learning with AI: Student intentions for academic use and broader perspectives on AI. *Technology, Knowledge and Learning*. Advance online publication. <https://doi.org/10.1007/s10758-026-09982-7>
- Cleland, J., Driessen, E., Masters, K., Lingard, L., & Maggio, L. A. (2026). When and how to disclose AI use in academic publishing: AMEE Guide No. 192. *Medical Teacher*, 48(4), 542–553. <https://doi.org/10.1080/0142159X.2025.2607513>
- Cohen, J. (2009). *Statistical power analysis for the behavioral sciences* (2nd ed.). Psychology Press.
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Cui, Y. (2025). What influences college students using AI for academic writing? A quantitative analysis based on HISAM and TRI theory. *Computers & Education: Artificial Intelligence*, 8, Article 100391. <https://doi.org/10.1016/j.caeai.2025.100391>
- Danyaro, K. U., Abdullahi, S., Abdallah, A. S., & Chiroma, H. (2025). Hallucinations in large language models for education: Challenges and mitigation. *International Journal of Teaching, Learning and Education*, 4(6), 13–19. <https://doi.org/10.22161/ijtle.4.6.2>
- Digital Education Council. (2024). *Digital Education Council global AI student survey 2024*. <https://www.digitaleducationcouncil.com/resource-library-items/digital-education-council-global-ai-student-survey-2024>
- Fathi, J., & Rahimi, M. (2026). Utilising artificial intelligence-enhanced writing mediation to develop academic writing skills in EFL learners: A qualitative study. *Computer Assisted Language Learning*, 39(1–2), 263–302. <https://doi.org/10.1080/09588221.2024.2374772>
- George, D., & Mallery, P. (2016). *IBM SPSS Statistics 23 step by step: A simple guide and reference* (14th ed.). Routledge.
- Gonsalves, C. (2025). Addressing student non-compliance in AI use declarations: Implications for academic integrity and assessment in higher education. *Assessment & Evaluation in Higher Education*, 50(4), 592–606. <https://doi.org/10.1080/02602938.2024.2415654>
- Haireche, N. E. H., & Lamri, C. (2025). Language and strategic challenges in EFL essay writing: A case study of master’s students at the University of Sidi Bel Abbes, Algeria. *Journal Faslo El Khitab*, 14(1), 515–526.
- Hammond, K. M., Lucas, P., Hassouna, A., & Brown, S. (2023). A wolf in sheep’s clothing? Critical discourse analysis of five online automated paraphrasing sites. *Journal of University Teaching & Learning Practice*, 20(7), Article 08. <https://doi.org/10.53761/1.20.7.08>

- Jin, Y., Yan, L., Echeverria, V., Gašević, D., & Martinez-Maldonado, R. (2025). Generative AI in higher education: A global perspective of institutional adoption policies and guidelines. *Computers and Education: Artificial Intelligence*, 8, 100348. <https://doi.org/10.1016/j.caeai.2024.100348>
- Jose, B., Cherian, J., Verghis, A. M., Varghise, S. M., S, M., & Joseph, S. (2025). The cognitive paradox of AI in education: Between enhancement and erosion. *Frontiers in Psychology*, 16, Article 1550621. <https://doi.org/10.3389/fpsyg.2025.1550621>
- Kadwa, M. S. (2025). Written Assignments and the Ethical Considerations of Artificial Intelligence in Higher Education. *International Journal of Linguistics, Literature and Translation*, 8(9), 254–257. <https://doi.org/10.32996/ijllt.2025.8.9.27>
- Karani, A., & Mwancha, C. N. (2025). The role of ChatGPT in STEM education: Assessing its prospects, challenges, ethical implications, and influence on student creativity and innovation competencies. *International Journal of Education, Technology and Science*, 5(2), 116–138. <https://doi.org/10.5281/zenodo.15569957>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kim, J., Yu, S., Detrick, R., & Li, N. (2025). Exploring students' perspectives on Generative AI-assisted academic writing. *Education and Information Technologies*, 30(1), 1265–1300. <https://doi.org/10.1007/s10639-024-12878-7>
- Liang, Y., Zhang, J., Chen, J., Wu, J., Shen, H., & Liang, W. (2025). Generative large language models and academic integrity: Ethical risks, detection challenges, and governance in the age of AI. *Ethics & Behavior*. Advance online publication. <https://doi.org/10.1080/10508422.2025.2608754>
- Lien, S., Skogli, E. W., Abamosa, J. Y., Glomdal, S., & Filkuková, P. (2025). Experience and attitudes toward the use of artificial intelligence in higher education. *Uniped*, 48(1), 1–14. <https://doi.org/10.18261/uniped.48.1.1>
- Lund, B., Mannuru, N. R., Teel, Z. A., Lee, T. H., Ortega, N. J., Simmons, S., & Ward, E. (2025). Student perceptions of AI-assisted writing and academic integrity: Ethical concerns, academic misconduct, and use of generative AI in higher education. *AI in Education*, 1(1), Article 2. <https://doi.org/10.3390/aieduc1010002>
- Malik, A. R., Pratiwi, Y., Andajani, K., Numertayasa, I. W., Suharti, S., Darwis, A., & Marzuki. (2023). Exploring artificial intelligence in academic essays: Higher education students' perspectives. *International Journal of Educational Research Open*, 5, Article 100296. <https://doi.org/10.1016/j.ijedro.2023.100296>

- Martín-Moncunill, D., & Alonso Martínez, D. (2025). Students' Trust in AI and Their Verification Strategies: A Case Study at Camilo José Cela University. *Education Sciences*, 15(10), 1307. <https://doi.org/10.3390/educsci15101307>
- Masrek, M. N., Heriyanto, H., & Hussain, A. (2026). Ethical decision making and astute use of artificial intelligence. *Edelweiss Applied Science and Technology*, 10(1), 1163–1176. <https://doi.org/10.55214/2576-8484.v10i1.11868>
- Massoud, O. S., & Zhang, J. (2025). Evaluating the impact of ChatGPT on ESL students' perceptions of writing skills and academic integrity. *Sophia University Junior College Division Faculty Journal*, 46, 33–54.
- McDonald, N., Johri, A., Ali, A., & Collier, A. H. (2025). Generative artificial intelligence in higher education: Evidence from an analysis of institutional policies and guidelines. *Computers in Human Behavior: Artificial Humans*, 3, 100121. <https://doi.org/10.1016/j.chbah.2025.100121>
- Meyer, J. G., Urbanowicz, R. J., Martin, P. C. N., O'Connor, K., Li, R., Peng, P.-C., Bright, T. J., Tatonetti, N., Won, K. J., Gonzalez-Hernandez, G., & Moore, J. H. (2023). ChatGPT and large language models in academia: Opportunities and challenges. *BioData Mining*, 16(1), 20. <https://doi.org/10.1186/s13040-023-00339-9>
- Mo, Z., & Crosthwaite, P. (2025). Exploring the affordances of generative AI large language models for stance and engagement in academic writing. *Journal of English for Academic Purposes*, 75, 101499. <https://doi.org/10.1016/j.jeap.2025.101499>
- Mustaffa, N. E., Lai, K. E., Preece, C. N., & Wong, F. Y. (2025). A bibliometric review of large language model hallucination. *International Journal of Research and Innovation in Social Science*, 9(9), 5025–5037. <https://doi.org/10.47772/IJRISS.2025.909000409>
- Nelson, A. S., Santamaría, P. V., Javens, J. S., & Ricaurte, M. (2025). Students' perceptions of generative artificial intelligence (GenAI) use in academic writing in English as a foreign language. *Education Sciences*, 15(5), Article 611. <https://doi.org/10.3390/educsci15050611>
- Nguyen, K. V. (2025). The Use of Generative AI Tools in Higher Education: Ethical and Pedagogical Principles. *Journal of Academic Ethics*, 23(3), 1435–1455. <https://doi.org/10.1007/s10805-025-09607-1>
- Nik Kamarudin, N. N. A., Shaari, S. N. M., Alias, A. Z., Abdul Raman, S., & Che Hasan, H. (2025). Exploring accounting students' perceptions and ethical concerns regarding the use of AI (LLMs) in coursework. *International Journal of Modern Education*, 7(27), 234–254. <https://doi.org/10.35631/IJMOE.727016>
- Noddings, N. (2013). *Caring: A relational approach to ethics and moral education* (2nd ed.). University of California Press.
- Pecorari, D. (2010). *Academic writing and plagiarism: A linguistic analysis*. Continuum.

- Pérez-Pérez, I., González-Afonso, M. C., Plasencia-Carballo, Z., & Pérez-Jorge, D. (2026). Transparency mechanisms for generative AI use in higher education assessment: A systematic scoping review (2022–2026). *Computers*, 15(2), Article 111. <https://doi.org/10.3390/computers15020111>
- Plata, S., De Guzman, M. A., & 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), 743–758. <https://doi.org/10.24191/ajue.v19i4.24697>
- Resnik, D. B., & Hosseini, M. (2026). Disclosing artificial intelligence use in scientific research and publication: When should disclosure be mandatory, optional, or unnecessary? *Accountability in Research*, 33(2), Article 2481949. <https://doi.org/10.1080/08989621.2025.2481949>
- Rest, J. R. (1986). *Moral development: Advances in research and theory*. Praeger.
- Rosalina, U., Shahat, S. A. A., & Sulaeman, N. F. (2026). Investigating students' views on the role of generative AI in academic writing. *Journal of General Education and Humanities*, 5(2), 2545–2558. <https://doi.org/10.58421/gehu.v5i2.1208>
- Rowland, D. R. (2023). Two frameworks to guide discussions around levels of acceptable use of generative AI in student academic research and writing. *Journal of Academic Language & Learning*, 17(1), T31–T69. <https://journal.aall.org.au/index.php/jall/article/view/915>
- Saad, S., & Zolkifli, A. N. F. (2026). ESL Undergraduates' Attitudes and Practices in AI-Assisted Academic Writing: A Mixed-Methods Survey. *International Journal of Research and Innovation in Social Science*, 10(4), 1898–1911. <https://doi.org/10.47772/IJRISS.2026.100400142>
- Sabbaghan, S., & Eaton, S. E. (2025). Navigating the Ethical Frontier: Graduate Students' Experiences with Generative AI-Mediated Scholarship. *International Journal of Artificial Intelligence in Education*, 35(4), 1860–1886. <https://doi.org/10.1007/s40593-024-00454-6>
- Smit, M., Wagner, R. F., & Bond-Barnard, T. J. (2025). Ambiguous regulations for dealing with AI in higher education can lead to moral hazards among students. *Project Leadership and Society*, 6, 100187. <https://doi.org/10.1016/j.plas.2025.100187>
- Spirgi, L., Seufert, S., Delcker, J., & Heil, J. (2024). Student perspectives on ethical academic writing with ChatGPT: An empirical study in higher education. In *Proceedings of the 16th International Conference on Computer Supported Education* (pp. 179–186). SCITEPRESS. <https://doi.org/10.5220/0012555700003693>
- Stephenson, R., & Armstrong, C. (2026). *Student generative AI survey 2026* (HEPI Report No. 199). Higher Education Policy Institute. <https://www.hepi.ac.uk/wp-content/uploads/2026/03/HEPI-Report-199-Gen-AI-Survey-2026.pdf>
- Subaveerapandiyan, A., Kalbande, D., & Ahmad, N. (2025). Perceptions of effectiveness and ethical use of AI tools in academic writing: A study among PhD scholars in India. *Information Development*, 41(3), 728–746. <https://doi.org/10.1177/02666669251314840>

- Tekir, S. (2026). Generative AI use in EFL writing: Associations with originality, critical reasoning, and metacognitive engagement in a Turkish higher education context. *Computer Assisted Language Learning*. Advance online publication. <https://doi.org/10.1080/09588221.2026.2617399>
- Tong, S. T., DeTone, A., Frederick, A., & Odebiyi, S. (2025). What are we telling our students about AI? An exploratory analysis of university instructors' generative AI syllabi policies. *Communication Education*, 74(3), 261–282. <https://doi.org/10.1080/03634523.2025.2477479>
- Triana, M. D., Naibaho, S. A., Sihite, S. H. B., & Hartati, R. (2025). The role of ChatGPT in enhancing paraphrasing skills to minimize plagiarism in student assignments. *Journal of Citizen Research and Development*, 2(1), 565–572. <https://doi.org/10.57235/jcrd.v2i1.4787>
- United Nations. (2019). *The Sustainable Development Goals report 2019*. United Nations.
- Utami, S. P. T., Andayani, A., Winarni, R., & Sumarwati, S. (2023). Utilization of artificial intelligence technology in an academic writing class: How do Indonesian students perceive it? *Contemporary Educational Technology*, 15(4), Article ep450. <https://doi.org/10.30935/cedtech/13419>
- Vendrell, M., & Johnston, S.-K. (2026). Scaffolding critical thinking with generative AI: Design principles for integrating large language models in higher education. *Computers & Education: Artificial Intelligence*, 10, Article 100572. <https://doi.org/10.1016/j.caeai.2026.100572>
- Virvou, M., Tsihrintzis, G. A., & Tsihrintzi, E.-A. (2025). Hallucinations in generative AI: Epistemic risks for learners in educational applications. In *2025 16th International Conference on Information, Intelligence, Systems & Applications (IISA)* (pp. 1–8). IEEE. <https://doi.org/10.1109/IISA66859.2025.11311276>
- Wang, J. (2025). EDCEW-LLM: Error detection and correction in English writing: A large language model-based approach. *Alexandria Engineering Journal*, 129, 1153–1164. <https://doi.org/10.1016/j.aej.2025.08.005>
- Xu, Y., Guo, J., Zhou, D., & Li, M. (2026). Chatbot or cheatbot? Moral judgments of AI usage in academic writing. *Education and Information Technologies*. Advance online publication. <https://doi.org/10.1007/s10639-026-13947-9>
- Yao, G. (2026). AI ethical awareness in L2 AI-assisted academic writing: Development and initial validation of the AEAL2-AWS. *Ethics & Behavior*. Advance online publication. <https://doi.org/10.1080/10508422.2026.2634324>
- Zhu, W., Huang, L., Zhou, X., Li, X., Shi, G., Ying, J., & Wang, C. (2025). Could AI ethical anxiety, perceived ethical risks, and ethical awareness about AI influence university students' use of generative AI products? An ethical perspective. *International Journal of Human–Computer Interaction*, 41(1), 742–764. <https://doi.org/10.1080/10447318.2024.2323277>