

# Generative AI in Higher Education: Global Trends, Challenges, and Policy Implications

Redouan Baghit<sup>1\*</sup>, Faisal Rahman<sup>2</sup>, Ho P.H Vu<sup>3</sup>, Sreepriya R<sup>4</sup>, Imran Hassan<sup>5</sup>, EL ABBASSI Marouane<sup>6</sup>, Dr. Mohd Asad Khan<sup>7</sup>, Rajesh Shahi<sup>8</sup>

<sup>1</sup>Faculty of Letters and Human Sciences El JADIDA, Chouaib Doukkali University, El Jadida, Morocco

<sup>2</sup>Department of IT Management, Southwest Baptist University, 4431 S Fremont Ave, Springfield, MO 65804 USA

<sup>3</sup>University of Technology Malaysia (UTM), Malaysia

<sup>4</sup>Assistant Professor, Department of English, VNR Vignana Jyothi Institute of Engineering and Technology, Hyderabad, India

<sup>5</sup>Department of Electrical and Electronic Engineering, World University of Bangladesh, Dhaka, Bangladesh

<sup>6</sup>Research Scholar, LASTI Laboratory, Moulay Slimane University, Beni Mellal, Morocco

<sup>7</sup>Associate Professor, PSIT College of Higher Education, Kanpur, U.P. India

<sup>8</sup>School of Business Management, Noida International University, Greater Noida, India

## Corresponding author

Redouan Baghit

Faculty of Letters and Human Sciences El JADIDA, Chouaib Doukkali University, El Jadida, Morocco

Email: [baghit.r@ucd.ac.ma](mailto:baghit.r@ucd.ac.ma)

**Abstract:** Artificial intelligence has moved from a peripheral tool to a routine feature of student life in higher education. This paper examines how generative AI, and ChatGPT specifically, is used across global and regional student populations, and what challenges accompany that use. Drawing on secondary analysis of four public survey datasets, including a global sample of 23,218 students across 110 countries (Ravšelj et al., 2025) and three available regional datasets from Bangladesh, Colombia, and Ecuador, the study addresses three research questions concerning prevalence and institutional readiness, task level application and regional variation, and the ethical and institutional challenges tied to AI adoption. Results show that 71.4% of the global sample had used ChatGPT, primarily for brainstorming, summarizing, and research assistance, though usage intensity and institutional policy awareness varied meaningfully by region. The three regional datasets converge on a common theme even though each measure it differently: a substantial share of students report reduced verification behavior and self-perceived overreliance on AI tools, a pattern this paper terms experienced dependency to distinguish it from the abstract ethical concern measured by the global dataset. The paper argues that adoption, ethical concern, and experienced dependency function as related but distinct dimensions rather than opposite poles of a single spectrum, and that institutions should track them independently. Implications for policy, faculty development, and future research are discussed.

**Keywords:** generative AI, higher education, ChatGPT, academic integrity, institutional policy

**How to cite this article:** Baghit R, Rahman F, Vu HPH, Sreepriya R, Hassan I, Marouane EA, Khan MA, Shahi R. Generative AI in higher education: global trends, challenges, and policy implications. *Int J Drug Deliv Technol.* 2026;16(75s): 466-477. DOI: 10.25258/ijddt.16.75s.54

## INTRODUCTION

Since its public release in November 2022, ChatGPT has moved from a novelty to a fixture of student life on campuses worldwide (Dai et al., 2023). Universities that once debated whether to allow generative AI in coursework are now asking how to govern its use, whether faculty are equipped to teach alongside it, and whether existing assessment

models still measure what they were built to measure (Bearman et al., 2023). This shift happened quickly enough that policy and evidence have struggled to keep pace with practice. Surveys published in the years immediately following ChatGPT's release found adoption climbing sharply among both students and educators (Ivanov et al., 2024; von Garrel & Mayer, 2023), yet institutional guidance on acceptable use, data privacy, and assessment redesign

has often lagged behind what is actually happening in classrooms.

One term needs fixing before proceeding. Generative AI, as used throughout this paper, refers to large language model tools capable of producing novel, human like text or other content in response to a prompt, distinct from earlier rule based or predictive AI systems used in education (intelligent tutoring, learning analytics,

automated grading, and administrative AI among them). The paper's focus on ChatGPT specifically, rather than generative AI broadly, is deliberate: adoption, policy gaps, and academic integrity concerns are currently most visible and most urgent around this one tool, and it is the case this paper's four datasets were built to measure.

A second distinction matters for how the results are read. This paper treats ethical concern and experienced dependency as related but separate constructs, not synonyms. Ethical concern refers to a student's opinion about generative AI's risks in the abstract, such as agreeing that ChatGPT encourages cheating; this is what the global dataset's regulation and ethics items measure. Experienced dependency refers to a student's self report of their own reduced critical engagement or reliance on AI in their own academic behavior, such as agreeing that their critical thinking has weakened; this is what the three regional datasets measure. The two are not interchangeable, and each section below is explicit about which construct it is reporting.

The findings that follow matter to three audiences in particular. For institutional policymakers, the gap this paper documents between student AI use and student awareness of institutional policy is consistent with, though not direct proof of, a communication problem rather than only a drafting problem. For faculty, the regional and task level usage patterns indicate where assessment redesign is most urgent and where it can wait. For researchers, the cross regional comparison of overreliance measures shows that how a survey item is worded, not only what it asks, shapes what students are willing to admit, a methodological point relevant well beyond this paper's specific datasets.

Much of the early literature on generative AI in education falls into one of two camps. One strand documents adoption and perception, asking how many students use these tools and what they think of them (Chan & Hu, 2023). A second strand focuses on risk, particularly academic integrity, and asks whether AI assisted writing can be distinguished from a student's own work (Cotton et al., 2023). Both

strands are necessary, but read separately they risk creating a misleading picture. A paper that measures only adoption can make usage look like an unqualified success story, while a paper that measures only concern can make every use of AI look like a form of cheating. Few studies attempt to hold both pictures at once, at scale, across more than one world region.

This paper attempts exactly that. It argues that adoption and concern should be treated as two related but distinct dimensions of the same phenomenon, not as opposite ends of a single scale. A region can show high adoption paired with low concern, as this paper finds for parts of Europe, or lighter adoption paired with disproportionately high task specific usage, as it finds for parts of Africa. Reading trends and challenges together, rather than as separate literatures, is necessary to make sense of these combinations.

A second contribution of this paper is methodological. Characterizing a fast moving global field convincingly requires either a large new primary survey, which was not feasible given the time and resources available for this project, or a disciplined secondary analysis of existing public data. This paper takes the second path. It draws on four publicly archived survey datasets: a global survey of higher education students spanning 110 countries (Ravšelj et al., 2025), and three independent regional surveys from Bangladesh, Colombia, and Ecuador (Acosta-Vargas, 2026; Correa, 2026; Islam et al., 2026). Analyzing these four sources together, without pooling them into a single artificial sample, allows the paper to speak with more empirical weight than a single institution study while remaining transparent about the limits of convenience sampled secondary data.

### Research Questions

Three research questions guide the analysis, moving from prevalence to application to challenge.

RQ1: What is the prevalence of generative AI adoption among higher education students, and how clearly do students perceive their institution's AI policy?

RQ2: For which academic tasks is generative AI most commonly used, and how does that usage vary across world regions?

RQ3: What ethical, technical, and institutional barriers do students associate with AI adoption, at both a global level and using three available regional datasets for illustrative depth?

The remainder of the paper proceeds as follows. The next section situates these research questions within recent literature on AI adoption, academic integrity, and institutional readiness. The Methodology section then details the four secondary datasets and how they were analyzed. The Results section reports findings organized by research question, followed by a Discussion that reads trends and challenges together and considers what they mean for institutional policy. The paper closes with limitations and directions for future research.

## LITERATURE REVIEW

### Emerging Trends in AI Adoption

Personalized learning and intelligent tutoring represent one of the oldest strands of AI application in education, predating generative AI by decades, but generative models have accelerated the pace of change considerably. Bayly-Castaneda et al. (2024), in a systematic review of AI driven personalized learning paths, found that generative tools are increasingly layered on top of earlier adaptive systems, offering conversational, on demand tutoring rather than the rule based sequencing that characterized earlier intelligent tutoring systems. Chatbots and virtual teaching assistants form a closely related trend. In a review of 67 studies, Labadze et al. (2023) found that AI chatbots most consistently benefited students in three areas: homework and study assistance, personalized learning support, and skills development, though the review also noted contradictory findings on whether chatbot use improves or dampens student motivation.

Generative AI adoption itself has grown at a pace unusual for educational technology. Ivanov et al. (2024) found that behavioral intention to adopt generative AI tools among higher education students was driven primarily by perceived usefulness and social norms, consistent with the Theory of Planned Behaviour. Strzelecki (2023) extended the Unified Theory of Acceptance and Use of Technology to ChatGPT specifically and found that performance expectancy and habit were the strongest predictors of continued use among university students, a finding that resurfaces later in this paper through one of the regional datasets, which used a related instrument. Von Garrel and Mayer (2023) documented similarly rapid uptake among German university students within a year of ChatGPT's release, while Chan and Hu (2023) found that students' overall stance toward generative AI is best described as cautiously optimistic. They value the tools for brainstorming and drafting but remain skeptical of their reliability for factual work, a distinction that recurs throughout the results reported later in this paper.

### Theoretical Framework

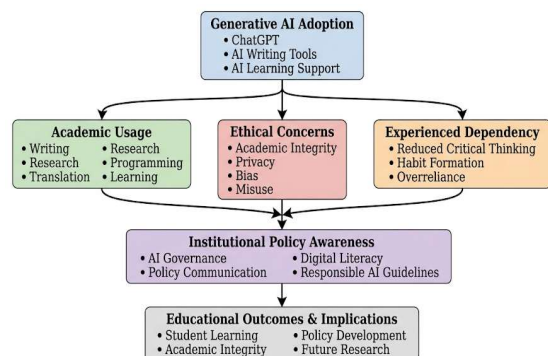
Three established technology adoption theories inform how this paper interprets its findings. The Technology Acceptance Model (TAM) proposes that perceived usefulness and perceived ease of use are the two central determinants of whether a person adopts a new technology (Davis, 1989). Applied here, TAM offers a ready explanation for the task-level usage pattern reported under RQ2: students gravitate toward brainstorming, summarizing, and research assistance, tasks where the usefulness of a conversational tool is immediately legible, and away from tasks such as calculating or professional writing, where either usefulness or ease of use may be less obvious to a student encountering the tool for the first time.

The Unified Theory of Acceptance and Use of Technology (UTAUT) extends this logic by adding social influence and facilitating conditions, organizational or institutional support that makes a technology easier or harder to use, as formal predictors of adoption alongside performance and effort expectancy (Venkatesh et al., 2003). Facilitating conditions is the construct most relevant to this paper's RQ1 findings. Institutional policy clarity is a facilitating condition in the UTAUT sense, and the finding that most students across every world region could not say whether such a policy existed is, in UTAUT terms, evidence of a facilitating condition that is absent or unobserved for the majority of the sample, independent of whatever the institution has actually decided. Strzelecki (2023) applied UTAUT directly to ChatGPT adoption among university students and found performance expectancy and habit to be the strongest predictors of continued use, a finding this paper's Ecuador dataset revisits directly, since it used items adapted from the same instrument.

Diffusion of Innovation theory (Rogers, 2003) offers a complementary lens for the regional variation reported under RQ1 and RQ2. Rogers describes adoption as an S-shaped process that moves through distinct populations, innovators, early adopters, early majority, late majority, and laggards, at a pace shaped by an innovation's relative advantage, compatibility with existing practice, and observability. Regions further along this curve would be expected to show not only higher adoption but also more differentiated, task-specific usage patterns, since later adopters within a diffusing population typically adapt an innovation to a wider range of specific uses rather than a single generic one. This offers one lens, among others, for interpreting why African respondents in this paper's sample combine mid-range overall adoption with unusually intensive, broadly

distributed task-level usage; a diffusion of innovation reading would treat this as a signature of a population still in an active diffusion phase rather than one that has already reached a plateau, though this study's cross-sectional design cannot confirm that trajectory directly.

None of these three theories was built for generative AI specifically, and none is tested directly here. The datasets weren't designed to formally estimate TAM, UTAUT, or Diffusion of Innovation — they're used as interpretive lenses instead, brought in to make sense of what turned up in the Results and picked back up in the Discussion. Figure 1 lays out a framework developed by the authors linking the study's central constructs: generative AI adoption, ethical concern, experienced dependency, institutional policy awareness, and regional context. The dashed arrows mark relationships consistent with the evidence, not causal pathways that have actually been tested.



**Figure 1: Conceptual Framework of Generative AI Adoption, Ethical Concerns, Dependency, Institutional Policy, and Educational Outcomes in Higher Education**

Note. Developed by the authors based on the synthesis of four secondary datasets and relevant theoretical perspectives, including the Technology Acceptance Model (TAM), Unified Theory of Acceptance and Use of Technology (UTAUT), and Diffusion of Innovation Theory. The framework illustrates conceptual relationships identified in the literature and the present analysis and should not be interpreted as an empirically validated causal model.

### Key Challenges Documented in Prior Research

Academic integrity is the most heavily researched challenge associated with generative AI in higher education. Cotton et al. (2023) were among the first to argue that ChatGPT complicates rather than simply amplifies existing plagiarism concerns, since AI generated text is not a direct copy of any single source and therefore evades traditional

plagiarism detection. Their call for redesigned assessment, rather than detection based enforcement alone, has since become a common recommendation in the wider literature.

A second, more recent challenge concerns the cognitive costs of AI reliance. Stadler et al. (2024) found experimentally that students who used large language models during scientific inquiry tasks expended less mental effort but also produced shallower reasoning, a pattern the authors describe as a trade of cognitive ease for cognitive depth. This finding is directly relevant to the self reported overreliance and reduced verification behavior documented later in this paper's regional datasets.

Algorithmic bias and equity form a third challenge. Baker and Hawn (2022), in a widely cited review, distinguish between bias that is known but unmeasured and bias whose existence is not yet known at all, arguing that most educational algorithms have been audited for only a narrow set of protected characteristics. The OECD (2023) has echoed this concern at a policy level, noting that the fairness of algorithmic systems in education remains poorly understood even in well resourced national systems, let alone in the lower income contexts represented by two of this paper's regional datasets.

Faculty and institutional readiness represent a fourth challenge, one closely tied to the institutional policy findings reported later in this paper. EDUCAUSE (2024) found that most higher education institutions in its readiness assessment had not yet formalized generative AI policy at the time of data collection, and UNESCO's AI competency framework for teachers (Miao & Cukurova, 2024) was developed specifically to address the resulting gap in educator preparedness. Industry survey data reinforces this picture from the administrative side: Ellucian (2024) found that concerns among higher education administrators about AI's effect on data privacy and academic integrity rose year over year even as adoption accelerated, suggesting that institutional anxiety has not been resolved simply by growing familiarity with the technology.

## METHODOLOGY

### Research Design

This study employs a secondary data analysis design using a cross regional comparative approach. Characterizing a fast moving global field with the breadth this paper requires, within the time and resources available, was not achievable through new primary data collection. Instead, the study synthesizes four existing, publicly archived survey datasets, applying original research questions and

analysis to data that other researchers collected for related but distinct purposes.

The design combines one large scale global dataset, used as the evidentiary base for prevalence, usage pattern, and baseline ethical concern findings, with three smaller regional datasets used to add contextual depth on the challenges dimension of the study. These three regional datasets are not evenly distributed geographically: one covers South Asia (Bangladesh) and two independently cover different parts of Latin America (Colombia and Ecuador). This imbalance reflects what could be located, obtained, and independently verified as genuine secondary data within the project's timeframe, rather than a deliberate attempt at balanced global coverage, and it is treated explicitly as a limitation in the closing section of this paper rather than glossed over. These three datasets are not framed as theoretically selected case studies; they are the regional data that were available and verifiable, used here for descriptive illustration of challenge-related themes rather than as a basis for regional generalization. No dataset is statistically pooled with another. Each is analyzed independently, using its own response scale and population, and findings are compared descriptively and thematically across sources.

### Data Sources

Table 1 summarizes the four datasets. The global dataset (Ravšelj et al., 2025) was collected by the Faculty of Public Administration at the University of Ljubljana in partnership with more than 200 international academic collaborators between October 2023 and February 2024. The questionnaire, distributed in seven languages, covered sociodemographics, usage patterns, perceived capabilities, regulation and ethics, satisfaction, study outcomes, skills development, labor market perceptions, and emotional responses. After cleaning, the dataset comprised 23,218 valid responses from students studying across 110 countries, recruited through classroom promotion and university communication channels using convenience sampling. It is archived on Mendeley Data and underpins a peer reviewed companion publication in PLOS ONE.

**Table 1: Overview of Data Sources**

| Dataset                       | Source                | N      | Scope                      | Sampling method                             |
|-------------------------------|-----------------------|--------|----------------------------|---------------------------------------------|
| Global ChatGPT Student Survey | Ravšelj et al. (2025) | 23,218 | 110 countries; 7 languages | Convenience (classroom/university channels) |

|                                          |                      |       |                                             |                                  |
|------------------------------------------|----------------------|-------|---------------------------------------------|----------------------------------|
| AI Dependency and Academic Vulnerability | Islam et al. (2026)  | 1,104 | 30 universities, Bangladesh                 | Convenience                      |
| Study Habits and AI Use                  | Correa (2026)        | 357   | 1 public university, Manizales, Colombia    | Proportional quota (14 programs) |
| Higher Education ChatGPT Usage Dataset   | Acosta-Vargas (2026) | 540   | Universidad de Las Américas, Quito, Ecuador | Convenience                      |

Note. N reflects raw, item-level responses used for analysis in this paper. All datasets are cross-sectional and were accessed as publicly archived files on Mendeley Data.

The Bangladesh dataset (Islam et al., 2026) is a cross sectional survey of 1,104 students across 30 Bangladeshi universities, using a nine item Likert battery measuring AI dependency, critical thinking erosion, verification behavior, and submission of unreviewed AI generated work. The Colombia dataset (Correa, 2026) surveyed 357 undergraduates at a single Colombian public university using proportional quota sampling across 14 degree programs, and includes a direct item asking whether academic AI use constitutes fraud. The Ecuador dataset (Acosta-Vargas, 2026) surveyed 540 students using items adapted from the Unified Theory of Acceptance and Use of Technology, including direct measures of habitual use and self perceived dependence.

### Data Preparation and Analysis

Because the four instruments differ in language, response scale, and item wording, they were not merged into a single pooled sample. Each dataset was cleaned and decoded independently. Numeric response codes in the global dataset were mapped to their questionnaire labels using the order in which options appear in the original instrument, a standard survey coding convention. Countries in the global dataset were classified into five world regions using the `pycountry_convert` library, supplemented

by manual classification for a small number of edge cases, achieving 99.7% coverage.

Analysis was conducted in Python using pandas, SciPy, and NumPy. Descriptive statistics (frequencies and percentages) were computed for all four datasets. For the global dataset, chi square tests of independence and Cramér's V were used to assess regional associations, since chi square alone becomes statistically significant for almost any difference at this sample size, while Cramér's V provides a scale independent measure of effect size. Regional datasets were analyzed descriptively given their smaller, non probability samples.

Cross dataset comparisons on shared underlying constructs, most notably self reported overreliance on AI, are reported descriptively and should be read as suggestive triangulation rather than as evidence of a controlled, statistically equivalent effect, given differences in population, instrument wording, and response scale length across the four sources.

### Ethical Considerations

All four datasets are secondary, de identified, and publicly shared for research reuse under their original collectors' informed consent and ethics protocols. No new human subjects data collection occurred in this study, and no new institutional ethics approval was required.

## RESULTS

### RQ1: Prevalence and Institutional Readiness

Among the 23,218 students in the global sample, 71.4% reported having used ChatGPT ( $n = 22,414$  valid responses). Among users, extent of use skewed toward occasional rather than habitual engagement: 32.0% reported occasional use, 28.8% moderate use, 21.2% rare use, 13.7% considerable use, and only 4.4% extensive use. The free version of the tool dominated, used by 88.6% of respondents, while only 3.3% relied exclusively on the paid subscription tier.

Adoption varied by world region, as shown in Table 2; these regional percentages come from one consistent instrument and are useful for descriptive comparison, but the convenience sampling within each region means they should not be read as precise, nationally representative estimates for any single country. South America showed the highest reported adoption at 83.4%, followed by Europe at 74.1%, while Africa, Asia, and North America clustered closely together between 65.0% and 66.6%. This regional association was statistically significant but modest in size, Cramér's  $V = 0.138$ ,  $\chi^2(4) = 420.4$ ,  $p$

$< .001$ , meaning region explains only a small share of the variation in whether a student has used ChatGPT.

Institutional policy awareness told a different story. Across every region, the modal response to whether a student's institution had a formal ChatGPT policy was "I don't know," ranging from 46.7% in Africa to 61.4% in Europe. This suggests that unclear policy communication is a global pattern rather than a regional one, consistent with EDUCAUSE's (2024) finding that most institutions had not yet formalized generative AI guidance at the time these data were collected. Where students did report a confirmed answer, Asia showed the highest rate of confirmed policy presence (32.3% "Yes"), while Africa showed the highest rate of confirmed policy absence (30.8% "No"). The regional association was again statistically significant but small, Cramér's  $V = 0.114$ ,  $\chi^2(8) = 374.9$ ,  $p < .001$ .

**Table 2: ChatGPT Adoption and Institutional Policy Awareness by World Region**

| Region        | Used ChatGPT (%) | Policy Yes (%) | Policy No (%) | Policy Don't Know (%) |
|---------------|------------------|----------------|---------------|-----------------------|
| Africa        | 66.6             | 22.5           | 30.8          | 46.7                  |
| Asia          | 65.0             | 32.3           | 20.9          | 46.8                  |
| Europe        | 74.1             | 21.9           | 16.7          | 61.4                  |
| North America | 65.8             | 25.0           | 19.6          | 55.4                  |
| South America | 83.4             | 23.8           | 17.3          | 58.8                  |

Note. Data from Ravšelj et al. (2025) global dataset. Percentages for policy awareness are row percentages restricted to respondents who answered the ChatGPT usage item. Cramér's  $V = 0.138$  (usage by region) and 0.114 (policy awareness by region), both  $p < .001$ .

### RQ2: Task-Level Application and Regional Variation

Across the full sample, students reported using ChatGPT most often (often or always) for brainstorming (29.0%), summarizing (27.4%), and research assistance (25.3%). Creative writing (11.3%), professional writing (12.5%), and calculating help (13.1%) were the least common uses. This ordering is broadly consistent with Chan and Hu's (2023) finding that students value generative AI most for idea generation and synthesis tasks rather than final output generation.

Breaking task usage down by region, shown in Table 3 and subject to the same descriptive-comparison caveat noted above, reveals a pattern not

visible in the overall ranking: intensity of use and breadth of adoption are separate dimensions. Africa reported the heaviest task specific usage across nearly every category, despite a mid range overall adoption rate of 66.6%, while North America reported the lightest usage across the board despite majority adoption. For academic writing specifically, 30.7% of African respondents reported often or always using ChatGPT, compared with 15.1% of North American respondents, a gap of more than two to one. This distinction between whether a student has adopted a tool and how intensively they use it once adopted is easy to lose in adoption statistics reported at a single, aggregate level, and appears to be an important nuance for future studies of generative AI diffusion to track separately. This regional breakdown does not control for discipline, age, or year of study, so an alternative explanation is that regional samples in this dataset differ in composition on one of these variables (for instance, a heavier STEM concentration in one region's respondents) rather than that region itself drives the usage difference; the global dataset used here does not support the subgroup analysis needed to rule this out, and future work with matched samples should test whether the pattern holds within discipline.

**Table 3: Task-Level Usage Frequency by World Region (% Reporting Often or Always)**

| Task                | Africa | Asia | Europe | N. America | S. America | M    |
|---------------------|--------|------|--------|------------|------------|------|
| Brainstorming       | 34.3   | 29.6 | 27.2   | 26.2       | 27.6       | 29.0 |
| Summarizing         | 29.4   | 27.3 | 26.0   | 22.3       | 30.1       | 27.0 |
| Research assistance | 34.1   | 21.7 | 23.0   | 18.4       | 28.4       | 25.1 |
| Academic writing    | 30.7   | 19.3 | 20.0   | 15.1       | 22.1       | 21.4 |
| Study assistance    | 33.7   | 21.2 | 17.9   | 12.0       | 19.5       | 20.9 |
| Personal assistance | 28.2   | 17.1 | 16.1   | 19.5       | 22.2       | 20.6 |
| Proofreading        | 21.7   | 17.4 | 15.4   | 19.5       | 25.8       | 20.0 |
| Translating         | 28.2   | 20.2 | 18.0   | 12.0       | 17.0       | 19.1 |
| Coding assistance   | 23.3   | 16.5 | 19.5   | 10.3       | 20.0       | 17.9 |
| Calculating help    | 18.3   | 13.9 | 12.0   | 8.1        | 11.0       | 12.7 |

|                      |      |      |      |     |      |      |
|----------------------|------|------|------|-----|------|------|
| Professional writing | 15.9 | 11.8 | 12.0 | 9.0 | 11.8 | 12.1 |
| Creative writing     | 15.2 | 10.9 | 11.0 | 7.9 | 9.8  | 11.0 |

Note. Data from Ravšelj et al. (2025) global dataset. M = unweighted mean of the five regional percentages, reported for cross-regional comparability rather than as a population estimate; it differs slightly from the sample-weighted overall percentage reported in text because regions contribute unequal numbers of respondents to the sample.

### **RQ3: Challenges (Ethical, Institutional, and Regional)**

At the global level, the most widely shared ethical concern was that ChatGPT encourages cheating (44.9% agree or strongly agree), closely followed by concern about misleading or inaccurate information (43.9%) and encouragement of plagiarism (43.5%). These findings align closely with Cotton et al.'s (2023) argument that generative AI complicates rather than simply extends existing academic integrity concerns. Students showed more trust in institutional than governmental oversight: 51.9% agreed ChatGPT should be subject to university or faculty ethical guidelines, compared with 38.2% who supported government regulation.

Regional breakdown of these same items, shown in Table 4 and again offered as descriptive comparison rather than a controlled cross-national test, revealed only small effect sizes overall (Cramér's V ranging from 0.055 to 0.098 across items), but one pattern stood out. Africa consistently scored highest on concerns framed around social and humanistic costs, including reduced human interaction (41.7%), increased social isolation (37.5%), and the risk that AI could replace formal education (34.9%), while Europe scored lowest on nearly every concern item despite reporting the second highest adoption rate in the sample. North America, by contrast, reported the highest concern specifically about misinformation (51.9%) and the strongest support for external, non institutional regulation. Read alongside the adoption findings from RQ1 and RQ2, this suggests three distinct regional postures toward generative AI: high use paired with low concern in Europe, moderate use paired with high social concern and heavy task specific reliance in Africa, and lighter use paired with high regulatory concern in North America.

**Table 4: Ethical Concerns About ChatGPT by World Region (% Agree or Strongly Agree)**

| ChatGPT might..              | Africa | Asia | Europe | N. America | S. America | M    |
|------------------------------|--------|------|--------|------------|------------|------|
| Encourage cheating           | 38.1   | 49.2 | 46.6   | 48.7       | 41.2       | 44.8 |
| Mislead with inaccurate info | 28.7   | 44.7 | 48.0   | 51.9       | 44.4       | 43.5 |
| Encourage plagiarism         | 38.6   | 45.9 | 45.3   | 47.3       | 40.0       | 43.4 |
| Hinder learning              | 43.1   | 45.3 | 40.7   | 42.6       | 35.2       | 41.4 |
| Threaten study ethics        | 32.2   | 42.6 | 34.4   | 41.9       | 35.9       | 37.4 |
| Reduce human interaction     | 41.7   | 39.9 | 35.2   | 33.5       | 26.2       | 35.3 |
| Encourage unethical behavior | 29.6   | 35.7 | 29.0   | 44.1       | 37.0       | 35.1 |
| Increase social isolation    | 37.5   | 31.9 | 26.2   | 29.1       | 20.0       | 28.9 |
| Replace formal education     | 34.9   | 28.3 | 24.9   | 21.7       | 20.6       | 26.1 |
| Invade privacy               | 21.1   | 28.9 | 23.6   | 25.2       | 17.8       | 23.3 |

Note. Data from Ravšelj et al. (2025) global dataset,  $n = 14,450$  to  $14,502$  valid responses across items.  $M$  = unweighted mean across the five regions. Cramér's  $V$  ranged from 0.055 to 0.098 across items (all  $p < .001$ ).

The three regional datasets add a layer of depth the global instrument's ethical concern items cannot provide on their own, because they measure

experienced dependency, students' self-reports about their own behavior and cognitive habits, rather than opinions about the technology in the abstract. In Bangladesh, 57.9% of students agreed their critical thinking had weakened due to AI reliance, and 59.8% agreed they had submitted academic work without fully understanding AI generated content, findings consistent with Stadler et al.'s (2024) experimental result that language model use reduces mental effort at the cost of reasoning depth. In Colombia, 37.3% of students perceived themselves as AI dependent, though only 10.1% considered academic AI use to constitute fraud, a considerably smaller share. In Ecuador, 44.6% agreed ChatGPT use had become a habit and 46.9% agreed it was essential for completing assignments, yet only 5.6% were willing to describe themselves as "addicted" to the tool. These indicators are summarized in Table 5.

Read together, these three regional findings suggest that reported experienced dependency is highest when an instrument asks about cognitive effects a student can observe directly, such as reduced critical thinking or difficulty recalling information, and lowest when an instrument asks students to accept a stigmatized label, such as fraud or addiction, for the same underlying behavior. Because the three instruments differ in wording, scale length, and population, this pattern should be read as a suggestive cross context observation rather than a controlled comparison; it nonetheless recurs consistently enough across three independent samples to merit attention in future instrument design.

**Table 5: Regional Overreliance and Dependency Indicators (Three Regional Datasets)**

| Dataset (region) | Indicator                                             | % Endorsing | N     |
|------------------|-------------------------------------------------------|-------------|-------|
| Bangladesh       | Critical thinking has weakened due to AI reliance     | 57.9        | 1,104 |
| Bangladesh       | Submitted work without fully understanding AI content | 59.8        | 1,104 |
| Colombia         | Perceives self as AI dependent                        | 37.3        | 357   |
| Colombia         | Believes academic AI use constitutes fraud            | 10.1        | 357   |
| Ecuador          | ChatGPT use has become a habit                        | 44.6        | 540   |
| Ecuador          | ChatGPT is essential for completing                   | 46.9        | 540   |

|         |                                         |     |     |
|---------|-----------------------------------------|-----|-----|
|         | assignments                             |     |     |
| Ecuador | Self-describes as "addicted" to ChatGPT | 5.6 | 540 |

Note. Indicators come from three independent instruments (Correa, 2026; Islam et al., 2026; Acosta-Vargas, 2026) with different response scales and wording. Percentages are not directly comparable across datasets and are presented for descriptive, thematic comparison only, not as a controlled measure of the same construct.

## DISCUSSION

The findings reported here support this paper's central argument that adoption, ethical concern, and experienced dependency function as related but distinct dimensions rather than opposite poles of a single continuum. Restricting attention first to ethical concern, as measured by the global dataset: Europe's combination of high adoption and low concern, Africa's combination of moderate adoption, heavy task specific use, and elevated social concern, and North America's combination of lighter use and high regulatory concern would each be flattened into a misleading single number if trends and challenges were reported separately, as much of the existing literature tends to do (Chan & Hu, 2023; Ivanov et al., 2024). Treating adoption rate as a proxy for institutional risk, or treating concern level as a proxy for actual usage, would misrepresent all three regions in this dataset.

The theoretical frameworks introduced earlier help explain why these regional patterns take this particular shape rather than some other one. Read through UTAUT (Venkatesh et al., 2003), the finding that most students across every region cannot say whether their institution has an AI policy points to facilitating conditions, the institutional and organizational support structures UTAUT identifies as a distinct predictor of technology use, that are functionally invisible to most of the population meant to benefit from them, regardless of what institutions have actually decided. Read through diffusion of innovation theory (Rogers, 2003), Africa's combination of mid-range adoption with disproportionately intensive, broadly distributed task usage is consistent with a population still moving through an active diffusion phase, where early and later adopters coexist and adapt the innovation to a widening range of specific uses, rather than a population that has already settled into a narrower, more habitual pattern of use, as Europe's combination of high adoption and low concern arguably suggests.

These theoretical readings are offered as plausible interpretations consistent with the data, not as formally tested models, since this paper's datasets were not constructed to estimate TAM, UTAUT, or diffusion of innovation parameters directly.

The distinction between institutional policy absence and policy invisibility is a second point worth emphasizing, though it should be read as an inference from student-side data rather than a demonstrated fact, since this study includes no faculty, administrator, or institutional-document data of its own. Across every region, "I don't know" was the most common response to whether a student's institution had a ChatGPT policy, outnumbering both confirmed presence and confirmed absence combined in most regions. One explanation is that institutions have simply not written policy yet; EDUCAUSE's (2024) readiness assessment and Miao and Cukurova's (2024) competency framework document that guidance and self-assessment tools existed during this period, but neither is direct evidence that institutions had adopted formal, written AI policy, so this explanation cannot be ruled out from the sources cited here. A second, not mutually exclusive explanation is that policy exists but has not reached students, a communication gap rather than a governance gap. The data in this paper cannot distinguish between these two explanations, since no institutional policy documents were examined directly; the hypothesis that this is chiefly a communication problem is plausible and, if confirmed by future research pairing student surveys with institutional policy audits, would imply that institutions see limited return from further policy development alone, with the more binding constraint being how policy is surfaced to students at the point of use, for instance embedded in assignment instructions or learning management systems rather than published only in standalone documents students rarely consult.

The cross regional comparison of experienced dependency indicators offers a third contribution, this time drawing on the three regional datasets rather than the global one. Bangladesh, Colombia, and Ecuador used three different instruments to ask, in effect, a related underlying question: whether students feel they have become dependent on AI in ways they themselves recognize. The share of students answering affirmatively varied considerably, from roughly one in ten on Colombia's fraud item to nearly six in ten on Bangladesh's critical thinking items, and this paper argues that much of this variation reflects framing rather than genuine underlying differences in student experience across these three national contexts. Stigmatized labels such

as fraud and addiction appear to suppress honest self report relative to behavior neutral items that ask about specific cognitive experiences, such as difficulty recalling information or submitting work without understanding it. This carries a direct methodological implication: researchers designing future instruments to measure AI dependency in student populations should be cautious about relying on a single stigmatized item as their primary outcome measure, since it likely understates the phenomenon relative to behavior specific alternatives.

These findings also speak to the algorithmic bias and equity literature (Baker & Hawn, 2022; OECD, 2023). None of the four datasets analyzed here contain direct measures of technical or infrastructure access, such as device ownership, connectivity quality, or data cost, despite these being plausible drivers of the regional adoption differences reported under RQ1 and RQ2. The proxy variables available, including urban or rural residence, self rated economic status, and institutional funding type, point toward access related explanations without confirming them directly. This is a genuine evidentiary gap in the current literature on generative AI in the Global South specifically, not simply a limitation of this paper's dataset selection, and it should be a priority for future data collection efforts.

Finally, the administrative side of this picture, documented by Ellucian's (2024) industry survey showing rising administrator concern about privacy and integrity even as adoption accelerates, suggests that the gap this paper identifies between student behavior and institutional policy clarity has a mirror image on the institutional side. Administrators appear increasingly aware of the risks associated with generative AI, but that awareness has not yet translated into policy that students can perceive. Closing this loop, from administrator concern to written policy to policy that students can actually name when asked, may be the more urgent institutional priority than debating whether to permit generative AI use at all, a debate that these data suggest has already been settled by students regardless of institutional position.

## LIMITATIONS

Several limitations qualify these findings. All four samples are non probability convenience samples; results describe the populations actually sampled, not a statistically representative global population of higher education students. Instrument heterogeneity across the four datasets means cross context comparisons, particularly the overreliance synthesis reported in the Results section, are descriptive and thematic rather than causal or

statistically pooled, and should be read with that caveat throughout.

Regional coverage is also uneven. Latin America is represented by two independent datasets (Colombia and Ecuador), while South Asia is represented by one (Bangladesh), and no dataset from East Asia, North Africa, or the Middle East met this paper's verification standard for regional depth within the available timeframe. These three datasets were selected on availability and verifiability, not on theoretical or representative grounds, and regional challenge claims in this paper should be read as descriptive evidence from the specific contexts sampled, not as representative of their broader regions or as a basis for regional generalization. Similarly, the regional differences in global-dataset task usage reported under RQ2 are not adjusted for discipline, age, or year of study, and could partly reflect differences in the composition of regional sub-samples rather than region-level effects alone. Finally, the paper's interpretation of low institutional-policy awareness as a communication gap rather than a governance gap is a hypothesis consistent with the student-side data, not a conclusion supported by direct evidence of institutional policy content, since no faculty, administrator, or institutional-document data were collected or analyzed.

Technical and infrastructure barriers, a component of RQ3 as originally framed, were addressed only through indirect proxy variables rather than direct measures of device access, connectivity, or cost, since no dataset available to this study measured these constructs directly. Finally, all data are self reported and cross sectional, collected at points between 2023 and 2026 depending on the dataset, during a period when both the underlying technology and institutional responses to it were changing rapidly. Findings should be understood as a snapshot rather than a stable characterization of a settled phenomenon.

## CONCLUSION

Generative AI has moved faster into the daily practice of higher education than institutional policy has moved to meet it. This paper's analysis of four independent survey datasets, spanning a global sample of more than 23,000 students and three available regional datasets, suggests that this gap is less a matter of institutions declining to act and more a matter of policy failing to reach the students it is meant to govern. Across every world region examined, most students could not say with confidence whether their institution had a generative AI policy at all.

The paper's central claim, that adoption, ethical concern, and experienced dependency operate as related but distinct dimensions rather than a single spectrum running from acceptance to rejection, has practical implications for how institutions monitor AI use going forward. A region or institution with high adoption is not necessarily one with low ethical concern, and a region with modest adoption is not necessarily one with light usage once students do adopt the tool. Both dimensions need to be tracked, and tracked separately, for institutional policy to respond to the situation students are actually in rather than an assumed one.

This paper's secondary data could not support several developments that future research should prioritize. The gap in direct infrastructure measurement is worth singling out: none of the four datasets analyzed here asked students directly about device ownership, internet access quality, or the affordability of data and connectivity, despite these being plausible drivers of the regional adoption and usage differences reported under RQ1 and RQ2. Given how much this paper has had to rely on indirect proxies such as urban or rural residence and self-rated economic status, closing this specific gap should be treated as a priority rather than one item among several. Beyond infrastructure, future studies should include:

- longitudinal designs capable of distinguishing genuine behavioral change in AI use from the framing effects this paper identified between stigmatized and behavior neutral survey items;
- matched student and faculty datasets within the same institutions, so that student-reported policy awareness can be checked directly against institutional practice rather than inferred;
- direct measures of digital infrastructure, including device access, connectivity quality, and affordability, rather than the indirect proxies this paper relied on; and
- institution-level policy analysis, examining actual written AI policy documents and their communication channels, to test directly whether the gap this paper identifies is one of communication, governance, or both.

As generative AI capabilities continue to change from year to year, the empirical picture presented here should be understood as a snapshot of a specific period rather than a stable baseline, and periodic replication using consistent instruments across regions would considerably strengthen the evidence base this paper has drawn on.

## REFERENCES

- Acosta-Vargas, P. (2026). Higher education ChatGPT usage dataset [Data set]. Mendeley Data, V1. <https://doi.org/10.17632/d6fnzyvjcc.1>
- Baker, R. S., & Hawn, A. (2022). Algorithmic bias in education. *International Journal of Artificial Intelligence in Education*, 32(4), 1052–1092. <https://doi.org/10.1007/s40593-021-00285-9>
- Bayly-Castaneda, K., Ramirez-Montoya, M. S., & Morita-Alexander, A. (2024). Crafting personalized learning paths with AI for lifelong learning: A systematic literature review. *Frontiers in Education*, 9, Article 1424386. <https://doi.org/10.3389/feduc.2024.1424386>
- Bearman, M., Ryan, J., & Ajjawi, R. (2023). Discourses of artificial intelligence in higher education: A critical literature review. *Higher Education*, 86(4), 369–398.
- Chan, C. K. Y., & Hu, W. (2023). Students' voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20, Article 43.
- Correa, C. (2026). Study habits and artificial intelligence use among university students: A proportionally stratified survey dataset [Data set]. Mendeley Data, V1. <https://doi.org/10.17632/mcwb2ppdsw.1>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*. Advance online publication. <https://doi.org/10.1080/14703297.2023.2190148>
- Dai, Y., Liu, A., & Lim, C. P. (2023). Reconceptualizing ChatGPT and generative AI as a student-driven innovation in higher education. *Procedia CIRP*, 119, 84–90.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- EDUCAUSE. (2024). Higher education generative AI readiness assessment. <https://library.educause.edu/resources/2024/4/higher-education-generative-ai-readiness-assessment>
- Ellucian. (2024). AI in higher education: Understanding the present and shaping the future (2nd annual AI in higher education survey report). [https://lp.ellucian.com/rs/085-MHT-312/images/Ellucian\\_2024\\_AI\\_Report.pdf](https://lp.ellucian.com/rs/085-MHT-312/images/Ellucian_2024_AI_Report.pdf)

- Islam, M. J., Sharmin, S., & Tuhin, H. H. (2026). AI dependency and academic vulnerability in higher education: A survey dataset from Bangladesh [Data set]. Mendeley Data, V1. <https://doi.org/10.17632/6s6nvmpgbb.1>
- Ivanov, S., Soliman, M., Tuomi, A., Alkathiri, N. A., & Al-Alawi, A. N. (2024). Drivers of generative AI adoption in higher education through the lens of the Theory of Planned Behaviour. *Technology in Society*, 77, Article 102521.
- Labadze, L., Grigolia, M., & Machaidze, L. (2023). Role of AI chatbots in education: Systematic literature review. *International Journal of Educational Technology in Higher Education*, 20, Article 56. <https://doi.org/10.1186/s41239-023-00426-1>
- Miao, F., & Cukurova, M. (2024). AI competency framework for teachers. UNESCO. <https://doi.org/10.54675/ZJTE2084>
- OECD. (2023). Algorithmic bias: The state of the situation and policy recommendations. In *OECD digital education outlook 2023*. OECD Publishing.
- Ravšelj, D., Keržič, D., Tomaževič, N., Umek, L., Brezovar, N., Iahad, N. A., Abdulla, A. A., Akopyan, A., Aldana Segura, M. W., AlHumaid, J., Allam, M. F., Alló, M., Andoh, R. P. K., Andronic, O., Arthur, Y. D., Aydin, F., Badran, A., Balbontín-Alvarado, R., Ben Saad, H., ... Aristovnik, A. (2025). Higher education students' perceptions of ChatGPT: A global study of early reactions. *PLOS ONE*, 20(2), Article e0315011. <https://doi.org/10.1371/journal.pone.0315011>
- Rogers, E. M. (2003). *Diffusion of innovations* (5th ed.). Free Press.
- Stadler, M., Bannert, M., & Sailer, M. (2024). Cognitive ease at a cost: LLMs reduce mental effort but compromise depth in student scientific inquiry. *Computers in Human Behavior*, 160, Article 108386.
- Strzelecki, A. (2023). Students' acceptance of ChatGPT in higher education: An extended Unified Theory of Acceptance and Use of Technology. *Innovative Higher Education*. Advance online publication. <https://doi.org/10.1007/s10755-023-09686-1>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- von Garrel, J., & Mayer, J. (2023). Artificial intelligence in studies: Use of ChatGPT and AI-based tools among students in Germany. *Humanities and Social Sciences Communications*, 10, Article 799.