

From Allowed to Embedded: Student Perceptions of AI Use Across Engineering Courses

Dr. Mike Borowczak, University of Central Florida

Dr. Borowczak, currently an Associate Professor of Electrical and Computer Engineering at the University of Central Florida, has over two decades of academic and industry experience. He worked in the semiconductor, biomedical informatics, and storage/security sectors in early-stage and mature startups, medical/academic research centers, and large corporate entities before returning to the US public university system full-time in 2018. His current research interest are focused on automation for design, development, and assessment of resilience, robustness, and security of electronic devices and systems which currently includes topics such as: advanced/novel cryptographic logic primitives (polymorphic, homomorphic, quantum-enhanced), assessment and evaluation of assistive technologies on semiconductor design and post-manufacturing operation, development and detection methods for AI-based sabotage. He and his students have published over 100 journal and conference publications. His research has been funded (~\$8.5M since 2018) by federal, national, state, and industrial entities.

Dr. Andrea Carneal-Burrows Borowczak, University of Central Florida

Dr. Andrea C. Burrows Borowczak is the Director of the School of Teacher Education and a Professor at the University of Central Florida (UCF) in the College of Community Innovation and Education (CCIE). She received her doctorate from the University of Cincinnati. Her research focuses on STEM integration, engineering education, and partnerships. She has published 45+ journal articles, received over \$10M in grant funding, received multiple awards, and presented papers, workshops, or posters at over 250 conferences. Dr. Borowczak continues promoting STEM education and integration in traditional and non-traditional settings. She was elected and served as PCEE Division Chair for 2022-2025, and she is dedicated to advancing educational knowledge and practices.

From Allowed to Embedded: Student Perceptions of AI Use Across Engineering Courses

Abstract

As generative AI becomes embedded in university learning environments, engineering educators are moving beyond initial reactive bans toward more complex integration strategies. However, normalization does not imply effective integration. This longitudinal study, emerging from a student-initiated audit of computing and electrical engineering courses at the University of Central Florida, traces the evolution of AI policies and student perceptions across nine semesters (Spring 2023 through Fall 2025). Utilizing the authors' defined Depth-of-Integration continuum, the team employed a participatory action research design to audit 128 course experiences across 71 unique courses, mapping the trajectory of institutional adaptation from initial disruption to current practice.

Findings reveal a dramatic shift in the AI policy landscape. While the explicit bans and absence of stated policy characterized the early response (Spring 2023), these categories have substantially diminished by late 2025, with Disallowed policies dropping to 12.9% of reported courses. However, this retreat from prohibition has not yielded pedagogical clarity. The data identify the rapid rise of Allowed (Informal) policies—a category that peaked at 75% in Fall 2023 and remained prevalent through Fall 2025. These permissive-but-vague environments are associated with the lowest student clarity scores ($M=2.45/5$, $n = 47$), significantly trailing both Encouraged ($M=4.37/5$, $p < .001$) and Disallowed ($M=4.50/5$, $p < .001$) environments. The study further provides longitudinal evidence of a Shadow Curriculum: 62.5% of students in courses classified Disallowed nonetheless reported using AI for at least one task. The paper concludes that the current phase of tacit permission is associated with what we call an *ethical vacuum*—a pedagogical environment in which students lack explicit guidance on acceptable AI use, leaving ethical boundaries undefined—and advocates for a shift toward Critical Use pedagogy to support responsible AI literacy and professional formation.

Introduction

By early 2026, the integration of Generative Artificial Intelligence (GenAI) in higher education had moved past its initial disruptive phase into a period of normalization. Tools like ChatGPT, Claude, and GitHub Copilot are no longer novelties but standard utilities in the engineering student's toolkit. In the context of computing and electrical engineering education, fields governed by strict accreditation standards and professional practice expectations, the implications of this integration are profound [1]. The ability to distinguish between generated approximation and verified fact is rapidly becoming a core competency for the modern engineer [2].

However, normalization does not imply effective integration. While initial institutional concerns regarding academic integrity have been widely documented [3, 4], they have arguably been replaced by a more insidious challenge: a complex landscape of inconsistent policies and hidden curricula. In the early days of GenAI (circa 2023), the message to students was often a clear, binary “No”—or, more commonly, silence. Today, that message has fractured into a spectrum of permissions, often communicated vaguely or omitted entirely from course syllabi.

This study investigates the evolution of AI policies in computing and electrical engineering courses at a large public research university in the southeastern United States. Unlike cross-sectional snapshots that capture a single moment in time, this work utilizes a longitudinal dataset spanning Spring 2023 to Fall 2025. This nine-semester window captures the full arc of the GenAI adoption cycle, from the release of GPT-4 to the present day. We posit that the primary challenge in engineering education is no longer preventing AI use, but bridging the Clarity Gap—the phenomenon where the decline of explicit bans has been replaced by ambiguous informal allowance. This ambiguity is associated with environments where students are unsure of ethical boundaries, potentially hindering the development of the professional verification skills needed for the workforce.

Related Work

The framework for this analysis draws upon three streams of literature: the evolution of academic integrity, engineering-specific pedagogical challenges, and emerging models of human-AI collaboration.

From Prohibition to Post-Plagiarism

Early institutional responses to GenAI were characterized by restriction, driven by concerns over assessment validity and the fear of widespread contract cheating. Susnjak and McIntosh (2024) [4] demonstrated how large language models could compromise the integrity of online examinations, while broader analyses revealed similar vulnerabilities across assessment types [5]. Bearman et al. (2024) [6] note that this reactive phase focused heavily on surveillance and detection, often at the expense of learning. However, as detection tools proved unreliable and the utility of GenAI became undeniable, the field began to shift toward what Eaton (2021) [3] describes as a “post-plagiarism” era. In this paradigm, the focus shifts from policing unauthorized tools to teaching ethical attribution and hybrid writing processes.

Engineering Pedagogy and Verification

Within engineering, the stakes of AI integration are distinct. Unlike creative writing or general humanities, engineering outputs such as code, structural calculations, and circuit designs, must function correctly in the physical world. Finnie-Ansley et al. (2022) [7] highlighted early on that while AI models like OpenAI Codex could pass introductory programming exams, they often struggled with edge cases and complex system integration.

Consequently, the pedagogical focus has shifted from code generation to code verification. Recent systematic reviews, such as Nikolic et al. (2023) [5], confirm that while AI can pass many standard engineering assessments, it fundamentally changes the nature of what we must assess, moving from

solution generation to solution auditing. The authors of this study have also identified a cautious adoption pattern among students, specifically that engineering students often treat AI as a fallible assistant rather than an oracle [8, 9]. This "Trust but Verify" mindset is consistent with the broader engineering ethos but requires explicit curricular support to become embedded in practice.

Models of Integration

To operationalize these concepts, educators have turned to models like the seven approaches described by Mollick and Mollick (2023) [10]. UNESCO guidelines now emphasize the need for human-centered AI approaches that prioritize critical thinking over mere detection, advocating for curricula that explicitly teach the limitations of these tools [11]. However, Lodge et al. (2023) [12] warn that without explicit pedagogical scaffolding, students may fail to develop the critical evaluative skills necessary to use these tools safely, relying on surface-level correctness rather than deep understanding.

Methodology

Participatory Action Research Design

This study utilized a participatory action research (PAR) design, in which engineering students acted as co-researchers to audit their own curricular experiences. This approach was chosen to mitigate *desirability bias*, the tendency for respondents to provide answers that are socially acceptable rather than accurate [13], which is often present in faculty-reported surveys, where instructors may overstate the clarity of their own policies. By surveying students directly about their lived experiences in the classroom, we obtain a more direct view of how policies are received and interpreted on the ground.

Respondents were recruited via an open call distributed through course announcements, departmental mailing lists, and word of mouth among undergraduate and graduate students in computing and electrical engineering programs. Participation was voluntary, and respondents provided experiences from courses they had taken across the study window. We acknowledge upfront that this is not a random sample: students who opted into a faculty-organized AI policy audit may be systematically more engaged with AI tools and more sensitive to policy ambiguity than the broader student population. Several respondents are members of the lead author's research lab; others are not. We address the implications of this composition in the Limitations section.

Data Collection

Data were collected via a structured survey distributed to undergraduate and graduate students in computing and electrical engineering courses at the University of Central Florida. The dataset comprises $N = 128$ course experience reports spanning nine semesters from Spring 2023 to Fall 2025, covering 71 unique courses reported by 30 individual student respondents across four disciplinary areas: Computer Engineering ($n = 67$), Computer Science ($n = 44$), Electrical Engineering ($n = 16$), and General Engineering ($n = 1$). Of the 128 reports, 96 (75%) came from undergraduate students and 32 (25%) from graduate students. Table 1 provides the per-semester breakdown of course reports, unique courses, and student level.

Table 1: Per-semester breakdown of course reports by student level.

Semester	Reports	Unique Courses	UG	Grad
Spring 2023	8	6	8	0
Summer 2023	2	2	2	0
Fall 2023	12	9	8	4
Spring 2024	14	13	11	3
Summer 2024	3	3	3	0
Fall 2024	22	21	14	8
Spring 2025	27	23	19	8
Summer 2025	9	7	9	0
Fall 2025	31	26	22	9
Total	128	71	96	32

The survey instrument included items on:

- **Policy Classification:** Students categorized the course policy based on syllabus text and verbal instructions using predefined categories along the Depth-of-Integration continuum.
- **Instructor Clarity:** A 5-point Likert scale assessing how well the student understood the boundaries of permissible AI use. The prompt read: “How clear was the instructor’s policy regarding AI usage?” with anchors at 1 (Very Unclear) and 5 (Very Clear).
- **Task Inventory:** A multi-select checklist of how AI was actually used in each course, with options including Explanation (Concepts/Lectures), Debugging (Code/Troubleshooting), Verification (Checking math/logic), Ideation (Brainstorming), Critique (Analyzing AI errors), and None (Did not use AI).

The instrument also included items on policy consistency across assignments and a free-text space for syllabus excerpts. Sentiment items related to usefulness, fairness, ethical clarity, and confidence were also collected but are not analyzed in the present paper; these will be reported separately in subsequent work that focuses specifically on student affect.

The AI Depth-of-Integration Continuum

To analyze instructor policies, we synthesized the work of Mollick and Mollick (2023) [10] and Eaton (2021) [3] into an AI Depth-of-Integration continuum. This classification system moves beyond a simple allowed/banned binary to capture the nuance of pedagogical intent:

1. **Disallowed:** Strict prohibition. Usage results in academic misconduct.
2. **No Policy Stated:** No AI policy appears in the syllabus or course communications; students have no basis to determine whether use is permitted.
3. **Allowed (Informal):** No explicit ban is present, but no specific guidance is given. Usage is tacitly permitted through silence or vague statements.

4. **Encouraged (Optional):** Instructor suggests usage as a helpful tool (e.g., "You may use AI to summarize readings"), but it is not central to assessment.
5. **Critical Use:** The course explicitly requires students to use AI and then critique, verify, or modify the output. This category includes courses where AI proficiency is a core learning objective or where AI is mandatory for specific assignments. The AI is a subject of study, not just a tool.

Coding and Analytic Notes

Two methodological points warrant explicit explanation. First, the policy classification (assigned by the student-respondent based on syllabus text and verbal instructions) and the clarity rating (the student's own subjective assessment) are not fully independent: the absence of explicit guidance is part of how Allowed (Informal) is defined, and unclear guidance is what the clarity item measures. We therefore treat the difference between policy categories on clarity not as a discovery ("vague policies are perceived as vague"), but as a baseline against which to examine within-category variation and the more substantive comparison among Disallowed, Encouraged, and Critical Use, categories that are separately defined on policy intent rather than on perceived clarity. Second, because Likert items are ordinal, we report means with standard deviations alongside non-parametric robustness checks where applicable, and we apply Bonferroni correction to all reported pairwise tests. This transformation establishes a stricter significance threshold.

Results

The longitudinal data reveal three distinct phases in the adaptation to AI: the Reactive Phase (2023), the Permissive Phase (2024), and the Stratification Phase (Late 2025).

Phase 1: The Reactive Phase (2023)

As illustrated in Figure 1, the earliest semester in the study (Spring 2023) was characterized not by widespread prohibition but by a policy vacuum: 50% of course reports indicated no stated AI policy, while Disallowed and Allowed (Informal) each accounted for 25%. By Fall 2023, the landscape had already shifted dramatically toward tacit permission, with Allowed (Informal) policies dominating at 75% of reports ($n = 9/12$) and Disallowed dropping to 16.7%. This rapid transition suggests that, rather than a prolonged ban period, many faculty quickly moved from silence to informal acceptance during the first year of widely available GenAI tools.

Phase 2: The Rise of Ambiguity (2024–Early 2025)

Through 2024 and into early 2025, the Allowed (Informal) category remained the single most prevalent policy mode, while Disallowed policies persisted but declined from 35.7% in Spring 2024 to 14.8% by Spring 2025. Notably, the Encouraged and Critical Use categories began to emerge in Fall 2024 (22.7% and 13.6%, respectively), signaling the first meaningful movement toward structured integration.

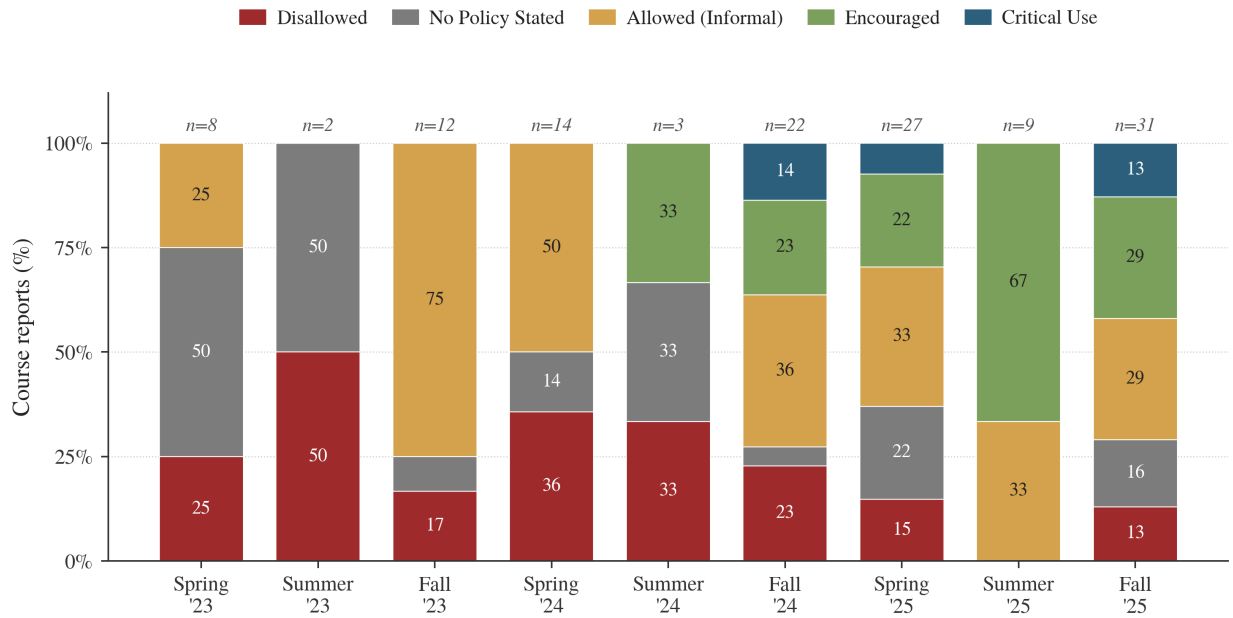


Figure 1: AI policy stratifies from a Spring 2023 vacuum to a five-category landscape by Fall 2025. Stacked bars show the percentage of course reports in each policy category by semester; sample sizes are indicated above each bar.

This period also saw the persistence of No Policy Stated reports, representing courses where students could find no evidence of an AI policy in syllabi or verbal instructions. These reports, while declining as a proportion, continued to appear through Fall 2025 (16.1%).

The dominance of informal allowance is associated with substantial confusion. As shown in Figure 2, Allowed (Informal) policies received the lowest clarity ratings from students ($M=2.45$, $SD = 1.14$, $n = 47$). A one-way ANOVA revealed significant differences in clarity across policy types, $F(3, 103) = 38.84$, $p < .001$. Post-hoc Welch's t -tests confirmed that Allowed (Informal) clarity scores were significantly lower than Disallowed ($t = -9.25$, $p < .001$, $d = -2.15$), Encouraged ($t = -8.54$, $p < .001$, $d = -1.96$), and Critical Use ($t = -6.80$, $p < .001$, $d = -2.09$). All three of these comparisons remain significant under Bonferroni correction at $\alpha_{adj} = 0.0083$. In contrast, there were no significant differences among Disallowed ($M=4.50$), Encouraged ($M=4.37$), and Critical Use ($M=4.44$) categories (all $p > .05$). A non-parametric Kruskal–Wallis test produced consistent results, $H(3) = 58.7$, $p < .001$. These large effect sizes indicate that the clarity deficit associated with informal policies is not merely statistically significant but substantively meaningful. Importantly, these patterns were consistent across student levels. Undergraduate students reported clarity means of 2.51 ($n = 37$) for Allowed (Informal) versus 4.30 ($n = 20$) for Encouraged, while graduate students reported 2.20 ($n = 10$) versus 4.57 ($n = 7$) for the same categories. The consistency of the Clarity Gap across both populations suggests that ambiguity is a structural feature of informal policies rather than a function of student experience level.

The headline pattern is U-shaped: clarity is achievable in either prohibition or structured permission; it is not achieved in the ambiguous middle. Disallowed, Encouraged, and Critical Use categories are statistically indistinguishable on clarity, whereas Allowed (Informal) sits roughly two points lower

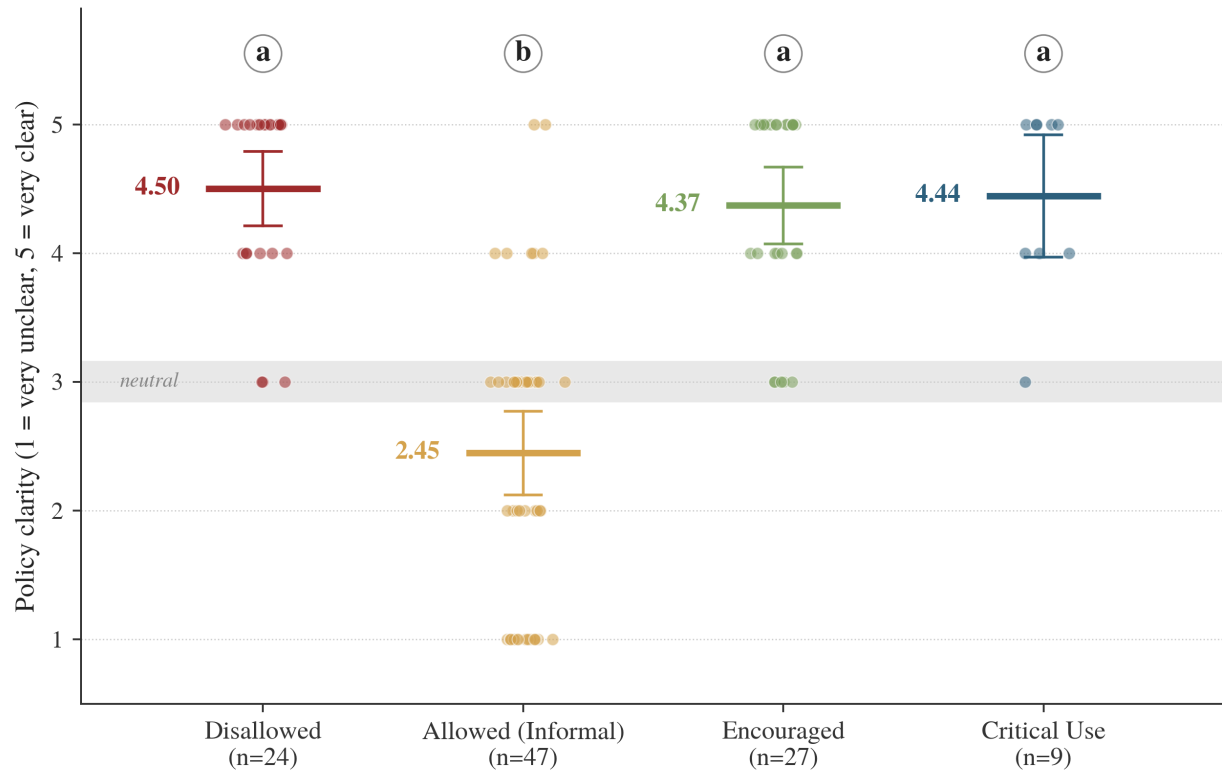


Figure 2: Allowed (Informal) policies are perceived as significantly less clear than every other policy type. Each point is one course report; thick horizontal bars show category means and the vertical whiskers show 95% confidence intervals. Categories sharing a letter (*a* or *b*) do not differ significantly under Welch's *t*-tests with Bonferroni correction ($\alpha_{adj} = .0083$); Group *b* is significantly lower than Group *a* ($p < .001$, $|d| > 1.9$ for all pairwise comparisons). No Policy Stated courses are omitted because respondents could not rate what was absent.

on a five-point scale. This framing addresses the natural concern that the Allowed (Informal) result is a definitional artifact: if low clarity were merely an automatic consequence of being classified as informally permissive, we would not expect Disallowed and Critical Use, which differ from each other on essentially every other dimension, to converge on the same high clarity scores.

Phase 3: Stratification (Late 2025)

A positive trend in the most recent semesters is the stratification of the policy landscape. By Fall 2025, no single policy category dominated: Allowed (Informal) and Encouraged each accounted for 29.0%, followed by No Policy Stated (16.1%), Disallowed (12.9%), and Critical Use (12.9%). The emergence of Critical Use policies, courses that explicitly require students to use AI and then critique or verify its output, represents a qualitative shift toward pedagogically grounded integration. In these courses, students reported high levels of ethical clarity and confidence, consistent with the "Trust but Verify" pedagogy advocated in previous work [8, 9].

Longitudinal Persistence of the Shadow Curriculum

A critical finding of this audit is the persistence of the Shadow Curriculum, the unauthorized or unaddressed use of AI tools. Even in courses classified Disallowed, 62.5% of student respondents ($n = 15/24$) reported using AI for at least one task during the term; only 37.5% reported abstaining entirely. Table 2 summarizes task-level AI use by policy category.

Table 2: Reported AI use by task and policy category (% of reports in each category).

Task	Disallowed ($n = 24$)	No Policy ($n = 21$)	Allowed ($n = 47$)	Encouraged ($n = 27$)	Critical Use ($n = 9$)
Explanation (Concepts/Lectures)	45.8	57.1	72.3	70.4	55.6
Ideation (Brainstorming)	33.3	47.6	25.5	59.3	55.6
Debugging (Code/Troubleshooting)	29.2	38.1	36.2	63.0	88.9
Verification (Math/Logic)	16.7	28.6	27.7	48.1	66.7
Critique (Analyzing AI errors)	16.7	14.3	8.5	25.9	66.7
Used AI for ≥ 1 task	62.5	66.7	78.7	88.9	100.0
Abstained entirely	37.5	33.3	21.3	11.1	0.0

The pattern within Disallowed courses warrants particular attention. The most common reported uses were Explanation (45.8%), Ideation (33.3%), and Debugging (29.2%), tasks that students often distinguished from "cheating" (e.g., generating assignment solutions), viewing them instead as study aids. However, because these behaviors were banned or unaddressed, they occurred in secrecy, preventing faculty from guiding students toward effective verification strategies. Critique, the highest-order use, in which students analyze AI errors, was rarely reported outside of Critical Use courses (66.7%), where it was central to the assignment design.

Discussion

The Clarity Gap and Student Anxiety

The data suggest that computing and electrical engineering education has successfully moved past widespread prohibition but may currently be stalled in a Clarity Gap. The prevalence of Allowed (Informal) policies is associated with a potential failure of pedagogical communication. When faculty do not set explicit boundaries, students report significantly higher ambiguity regarding academic integrity, consistent with Moorhouse et al.'s (2023) [14] findings on assessment guidelines. A policy of tacit permission may leave students vulnerable to inconsistent enforcement and could deny them the opportunity to learn professional norms for AI use.

Implications for Engineering Accreditation

These findings may have implications for engineering accreditation, particularly regarding outcomes related to ethical and professional responsibility. If accreditation bodies move toward expecting graduates to demonstrate competency in AI verification, an evolution that recent disciplinary discussions suggest is plausible, programs that ignore AI, or permit it only tacitly, may not be in a strong position to evidence such outcomes. The Shadow Curriculum data are consistent with

the concern that students are acquiring AI tools without ethical guidance: in courses classified Disallowed, the majority of respondents reported using AI anyway, but for tasks they distinguished from cheating. The Critical Use model offers one potential pathway to map AI competencies directly to outcome assessment, though the evidence here is suggestive rather than definitive given the small Critical Use subsample ($n = 9$).

Toward a "Trust but Verify" Pedagogy

The evidence presented here suggests that to address the Shadow Curriculum, faculty should consider moving beyond informal allowance toward structured integration. The high clarity scores associated with Critical Use courses indicate that students may benefit from explicit boundaries and structured engagement with AI tools. We propose that engineering educators consider a Verification First approach, where assignments are designed assuming AI availability and the assessment metric emphasizes the rigorous validation of AI-generated outputs rather than their production. This approach would shift the cognitive load from syntax (which AI handles well) to semantics and system logic (which require human expertise), aligning with the professional competencies that accreditation bodies have historically sought to develop.

Limitations

This study has several limitations that should be considered when interpreting the findings. First, the data are self-reported by students and may be subject to recall bias. Second, the study is situated within a single institution and mostly within computing and electrical engineering disciplines (Computer Engineering, Computer Science, Electrical Engineering, and General Engineering); the findings may not generalize to other engineering fields or to institutions with different AI policy cultures. Third, although 30 individual students contributed 128 course experience reports, the dataset is not a random sample. Students who opted into the audit may hold systematically different views than the broader population, likely skewing toward higher AI engagement and greater sensitivity to policy ambiguity, and several respondents are members of the lead author's research lab. The direction of bias matters: if our sample over-represents AI-engaged students, then the Shadow Curriculum prevalence reported here is likely an upper bound for the broader student population, while the Clarity Gap reported here may be more pronounced in our sample than in a random one. Fourth, sample sizes for individual semesters, particularly the early semesters and summer terms, are small, and proportional analyses should be interpreted with caution. Fifth, the Critical Use category, while showing promising clarity scores, has a small sample size ($n = 9$ after merging task-specific and embedded policy variants) and warrants further investigation as more courses adopt this approach. Finally, as noted in the Methodology, the Allowed (Informal) classification and the clarity rating share definitional overlap; the more interpretable comparison is therefore among Disallowed, Encouraged, and Critical Use, where policies differ on intent rather than on perceived clarity.

Future work should expand this audit to other engineering disciplines and institutions, incorporate faculty perspectives to compare stated policy intent with student perception, and analyze the correlation between specific AI policies and student performance outcomes.

Conclusion

This audit reveals an emerging trajectory: explicit bans have largely receded, but they have often been replaced by ambiguity rather than clarity. The data suggest that to support the next generation of engineers, programs should consider replacing informal policies with explicit guidance. Silence, the evidence indicates, is associated with the lowest clarity scores in the study; meanwhile, 62.5% of students in courses that disallowed AI reported using it anyway, predominantly for study-aid tasks they did not consider to be cheating. The Critical Use model, while still emergent, offers a promising framework: courses that require students to generate, verify, and critique AI outputs are associated with the highest levels of student clarity and ethical confidence. By formalizing a "Trust but Verify" approach, engineering programs have an opportunity to transform AI from a threat to academic integrity into a catalyst for deeper professional formation.

Acknowledgments

The authors gratefully acknowledge the engineering undergraduate and graduate students at the University of Central Florida who participated in and led the data collection for this audit. Their willingness to document their own classroom experiences, classify course policies, and share candid reflections on AI use across nine semesters made this longitudinal analysis possible.

The authors also acknowledge the use of Google Copilot and Anthropic Claude to assist in the refinement of the data transformation and analysis pipeline, along with restructuring, formatting, and clarity edits to the main body of this work. All findings, interpretations, and conclusions are the responsibility of the authors.

References

- [1] ABET Engineering Accreditation Commission, *Criteria for Accrediting Engineering Programs, 2025–2026*. ABET, Baltimore, MD, 2024.
- [2] J. Qadir, "Engineering education in the era of ChatGPT: Promise and pitfalls of generative AI for education," in *2023 IEEE Global Engineering Education Conference (EDUCON)*, pp. 1–9, IEEE, 2023.
- [3] S. E. Eaton, *Plagiarism in Higher Education: Tackling Tough Topics in Academic Integrity*. Santa Barbara, CA: Libraries Unlimited, 2021.
- [4] T. Susnjak and T. R. McIntosh, "ChatGPT: The end of online exam integrity?," *Education Sciences*, vol. 14, no. 6, p. 656, 2024.
- [5] S. Nikolic, S. Daniel, R. Haque, M. Belkina, G. M. Hassan, S. Grundy, S. Lyden, P. Neal, and C. Sandison, "ChatGPT versus engineering education assessment: A multidisciplinary and multi-institutional benchmarking and analysis of this generative artificial intelligence tool to investigate assessment integrity," *European Journal of Engineering Education*, vol. 48, no. 4, pp. 559–614, 2023.
- [6] M. Bearman, J. Tai, P. Dawson, D. Boud, and R. Ajjawi, "Developing evaluative judgement

for a time of generative artificial intelligence,” *Assessment & Evaluation in Higher Education*, vol. 49, no. 6, pp. 893–905, 2024.

- [7] J. Finnie-Ansley, P. Denny, B. A. Becker, A. Luxton-Reilly, and J. Prather, “The robots are coming: Exploring the implications of OpenAI codex on introductory programming,” in *Proceedings of the 24th Australasian Computing Education Conference, ACE '22*, (New York, NY, USA), pp. 10–19, Association for Computing Machinery, 2022.
- [8] M. Borowczak and A. C. Borowczak, “Integrating generative AI into microelectronics education: Implications for learning and pedagogical practice,” in *Proceedings of the Great Lakes Symposium on VLSI 2025 (GLSVLSI '25)*, (New York, NY, USA), ACM, 2025.
- [9] M. Borowczak and A. C. Borowczak, “Trust but verify: Engineering students’ adoption of generative ai and pedagogical implications across disciplines,” in *2026 ASEE Southeastern Section Conference*, 2026.
- [10] E. R. Mollick and L. Mollick, “Assigning AI: Seven approaches for students, with prompts,” *SSRN Electronic Journal*, 2023. Working paper, The Wharton School. Not peer-reviewed.
- [11] UNESCO, “Guidance for generative AI in education and research.” <https://unesdoc.unesco.org/ark:/48223/pf0000386693>, 2023.
- [12] J. M. Lodge, P. de Barba, and J. Broadbent, “Learning with generative artificial intelligence within a network of co-regulation,” *Journal of University Teaching and Learning Practice*, vol. 20, no. 7, pp. 1–10, 2023.
- [13] I. Krumpal, “Determinants of social desirability bias in sensitive surveys: A literature review,” *Quality & Quantity*, vol. 47, no. 4, pp. 2025–2047, 2013.
- [14] B. L. Moorhouse, M. A. Yeo, and Y. Wan, “Generative AI tools and assessment: Guidelines of the world’s top-ranking universities,” *Computers and Education Open*, vol. 5, p. 100151, 2023.