

Using Machine Learning Techniques to Develop Attitudinal Surveys on Generative AI Tool Adoption in the Engineering Classroom

Dr. David M. Feinauer P.E., Virginia Military Institute

Dr. Feinauer is a Professor of Electrical and Computer Engineering at Virginia Military Institute. His scholarly work spans a number of areas related to engineering education, including the first-year engineering experience, incorporating innovation and entrepreneurship practice in the engineering classroom, and P-12 engineering outreach. Additionally, he has research experience in the areas of automation and control theory, system identification, machine learning, and energy resilience. He holds a PhD and BS in Electrical Engineering from the University of Kentucky.

Dr. Michael Cross, Norwich University

Michael Cross, Ph.D., is Chair and Associate Professor of Electrical & Computer Engineering at Norwich University. He teaches classes in first-year engineering, circuits, electronics, and energy systems. His research interests are in battery electric vehicles, autonomous vehicles, and grid-connected energy systems. Cross received his B.S. in Physics from the Rochester Institute of Technology and his M.S. and Ph.D. in Materials Science from the University of Vermont, where he investigated semiconductor thin films and nanomaterials for sustainable energy systems. As a development engineer at IBM Microelectronics, he contributed to the development of advanced semiconductor technologies. He is also the Director of the Governor's Institutes of Vermont Engineering Institute, promoting STEM outreach and workforce development in the state.

Kundan P. Kushwaha

Ali Al Bataineh, Norwich University

Using Machine Learning Techniques to Develop Attitudinal Surveys on Generative AI Tool Adoption in the Engineering Classroom

Introduction

As faculty, administrators, and students are bombarded with polarizing opinions regarding the use of Generative AI (GenAI) tools in higher education learning environments as either a boon or a bane, it is important to explore the perceptions of these constituencies and to examine their readiness to embrace AI technologies throughout their studies. Policies are rapidly evolving and student sentiment towards both the policies and the application of AI tools to learning is hard to interpret.

In preparation for a future study on this topic, the authors developed a Large Language Model (LLM)-based topic modeling and sentiment analysis of course-level, department-level, and institution-level policies relating to the use of AI in the classroom. The LLM-enabled analysis provides a methodological demonstration of the technique, an early thematic map of AI policies or positions, and lays the groundwork for future survey instrument design.

The results of the modeling and analysis will inform the future development of a detailed survey and other instruments that explore key themes and sentiments about AI use cases and limitations. The long-term result of this effort will be an analysis of constituent survey results that can inform future pedagogical approaches to teaching with AI technologies and contribute to the growing body of literature regarding academic integrity in the age of GenAI [1] and productive uses for GenAI in the classroom [2]. The contribution towards this goal reported in this paper will be a detailed discussion of the application of LLM techniques to capture the broader institutional contexts related to Generative AI adoption for education.

This project involves students and faculty at two universities with well-established honor codes that are characterized by core values of honesty and integrity with harsh penalties for failing to live according to those values. During in-class discussions with students regarding their use of AI, the authors have found it difficult to gauge true student sentiment. Institutional AI policies combined with course-level policies that vary widely from faculty member to faculty member signal to the students the likely attitudes or beliefs related to AI use, influencing the opinions reported or shared. The broader honor code context of the institutions complicates this. Throughout their engineering education, students are called upon to conceptualize, organize, plan, create, analyze, evaluate, experiment, memorize, report, communicate, and synthesize. Various AI-tools could provide great assistance for many of these modes of learning, but the nature of most surveys could lead constituents to consider the listed use cases, amplifying their perceived value. Through the LLM topic modeling and sentiment analysis techniques presented, the authors explore the influence of course policies on prevailing attitudes and the adoption of artificial intelligence in academic practices at institutions with notably strong honor codes addressing academic integrity and personal responsibility.

Background

A systematic literature review of GenAI with respect to Higher Education [3] evaluated 41 studies and detailed the prevailing sentiment of the need for institutions to enhance students' skills with respect to digital literacy and the use of GenAI tools while also mitigating risks related to academic dishonesty [3]. In [4], a mixed methods approach of analyzing AI policy statements at 50 universities, topic modeling revealed 4 broad areas of treatment in the policies – “Integration of GenAI in Learning and Assessment, GenAI in Visual and Multimodal Media, Security and Ethical Considerations in GenAI, and GenAI in Academic Integrity.” The same study highlighted a highly positive sentiment towards GenAI overall but revealed that students had the lowest average sentiment score of all constituencies, perhaps due to the regulations and risks associated with GenAI usage. In a survey of faculty and students at a large university [5], it was found that students and faculty were adopting GenAI usage at similar rates and that both groups felt use would have more negative than positive effects on problem-solving, teamwork, self-efficacy, test anxiety, intrinsic motivation, and engagement. Among the students, it was found that the STEM males had the most positive outlook on use of GenAI for learning with non-STEM females demonstrating the least positive outlook [5].

In [6], a study at one institution reported that cultural background significantly impacted student attitudes towards the use of GenAI for learning, and that most instructors permitted use of AI under limited conditions, but they lacked formalized classroom policies. Studies such as this highlight an increased need for instructor and institutional readiness. To support faculty in developing course-specific GenAI policies, many university centers for teaching and learning have developed support resources with example considerations and an array of policies that address use of GenAI tools. The resources at Stanford [7] offer statements with varying degrees of acceptable use that could be broadly described in 4 categories related to the level of GenAI use allowed: ***Yes Always, Yes But, No But, and No Never***. Even with policies that are open without restriction, such as Yes Always, there is room for caution and restrictions related to the data shared and a requirement for usage to be cited. No Never policies are characterized as a total prohibition of using GenAI tools for graded work, Yes But policies default to allowed use of GenAI tools, with specific contexts named where usage is prohibited. No But policies tend to default to a prohibition of tools, with a list of limited circumstances where usage would be allowed. In a 2023 policy paper, the Office of Educational Technology provided recommendations for drafting AI policy statements related to teaching and learning to guide people; address equity; ensure safety, ethics and effectiveness; and promote transparency [8]. With variations in sentiment among faculty and students and cultural background and academic discipline affecting sentiments related to the use of GenAI for teaching and learning, future studies would benefit from these considerations as surveys are developed to inform the creation of more effective AI policies.

As instructors and institutions address the everchanging landscape of GenAI tools for teaching and learning with policies, it is important for those policies to address Academic Integrity,

acknowledging that Institutional cultures and honor codes will shape disclosure and self-reporting of integrity violations. Furthermore, students' attitudes will be affected by the local policy climate. As students try to navigate this complex policy environment, sensemaking will be a challenge.

To analyze GenAI policies, the authors constructed a study built on transformer-based natural language processing techniques selected for short texts. Transformer sentence embeddings (e.g., *allMiniLML6v2* and *DistilBERT*) capture contextual meaning beyond surface-level word frequency analyses, enabling comparison of policy language even when documents are brief or stylistically varied. These embeddings support both classical and neural topic modeling approaches. BERTopic analysis combined the transformer embeddings with clustering and class-based TF-IDF (Term Frequency, Inverse Document Frequency) methods to produce more coherent and interpretable topics. Sentiment analysis was then conducted using VADER (Valence Aware Dictionary and sEntiment Reasoner), a lexicon and rule-based method built on research-informed, standardized "feeling scores" that accounts for intensity modifiers, emphasis, and negation in a linguistic context.

After multiple iterations of analyzing the collected policy statements, the authors determined that the collected corpus of policies were insufficient in size for the analysis technique originally designed and described above. Rather than scratch the study, the details of the study are shared for background context, and the analysis proceeded using natural language prompts with an LLM to create initial results. The method for collecting the policies and synthesizing and categorizing them is described in the Methodology section. The results of this initial LLM-based analysis of the policies and the experience using the author-designed analysis technique described above will inform the design of a future analysis and study as well as the policy collection process.

Methodology

Institutional, departmental, and course-specific GenAI policy statements were collected by the authors. The authors work at universities with strong honor codes related to the military affiliation of the institutions. To avoid over-representing or under-representing various constituencies in the study, a variety of data collection techniques were used. At Virginia Military Institute, each department was required to develop a departmental policy on the use of GenAI that faculty could adopt for their courses. As such, one statement representative of each department was included in the corpus of policies studied. At Norwich University, the authors collected syllabi from colleagues that included AI policy statements, restricting the number of statements from a participating department so as not to bias the collection towards a limited set of disciplines. To augment the number of policies collected, the authors conducted an internet search and collected additional statements that were publicly available, selecting similar institutions with strong honor codes and military affiliations.

The GenAI pertinent aspects of the policies were extracted from their source documents, cleaned, de-identified, and boilerplate language was removed. The portion of the policy specific

to use of GenAI was labeled by the authors regarding its level of application (institution, department, or course) and it was categorized according to the perceived approach to the use of GenAI tools (Yes Always, Yes But, No But, and No Never) according to the Stanford policy development guidance [7].

As discussed in the background section, the author-designed transformer-based model and methods did not produce meaningful topic and sentiment analysis results. The corpus of text collected from 31 policy statements divided into 4 categories was insufficient for meaningful analysis. As a result, the authors modified their plans for more thorough collection of statements to inform future work, and they turned to a modern large language model (LLM) tool to analyze the policies. The policy excerpts were analyzed with the CoPilot LLM (ver. 2.20260108.46.0) with enterprise data restrictions enabled. The CoPilot GPT was prompted to provide a high-level synthesis of the policies, to identify categories or levels of policy restriction, to sort the policies into those categories, and to analyze trends or theme variations among the policies based on the level at which the policy was developed and applied (institution, course, or department). The author-generated categorization of the policies was not provided to the GPT. The output was reviewed by the authors for validity and summarized in the results section. Implications of the findings are discussed in the Conclusion section.

Results

Thirty-one policies across multiple Senior Military Colleges (SMCs) and US Service Academies (SAs) were analyzed, including policies from multiple disciplines. The policies were manually categorized into three levels: institutional policy, departmental policy, or course policy. The policies were also categorized into restriction levels using the Stanford methodology: Yes Always, Yes But, No But, and No Never [7]. The first and final levels allow for either the use of AI tools or none at all, respectively. An example of the Yes But category would be allowing most AI tools to be used, with proper citation, while restricting the use of tools for a limited set of tasks or contexts. An example of the No But category would be allowing the use of spelling/grammar checkers or another restrictive list of pre-approved tools after the draft work has been completed by the student. A summary of the author-labeled classifications of the 31 GenAI policies used in this study is shown below in Figure 1.

The all-encompassing nature of the institutional policies suggests that they might be best considered separately from the more course delivery focused policies. Removing the institutional policies resulted in the plots shown in Figure 2.

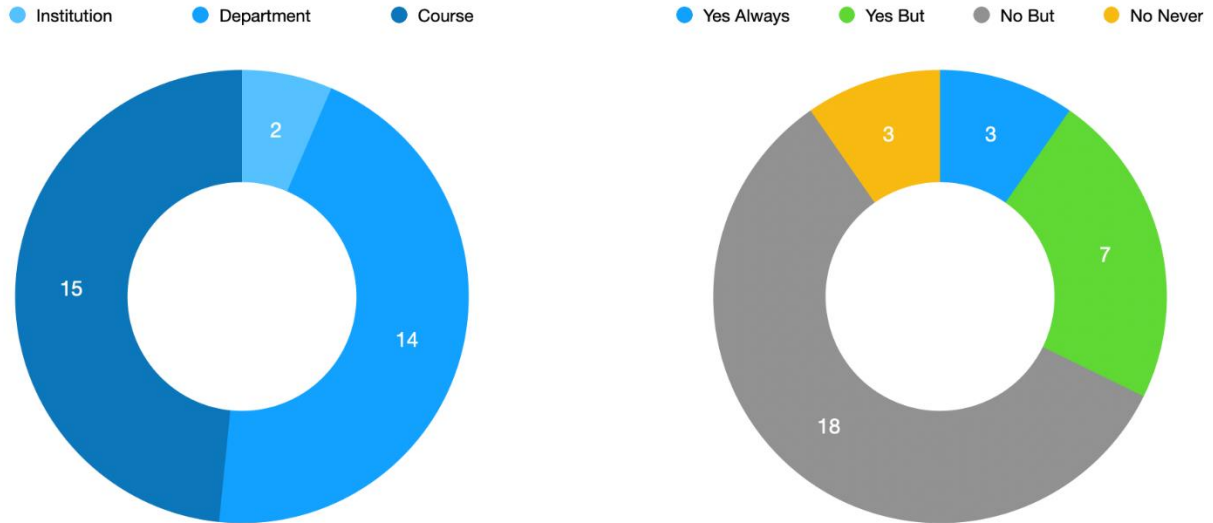


Figure 1: Summary of analyzed statements showing (left) the policy level and (right) the restriction level of the policy. The total count was $N = 31$.

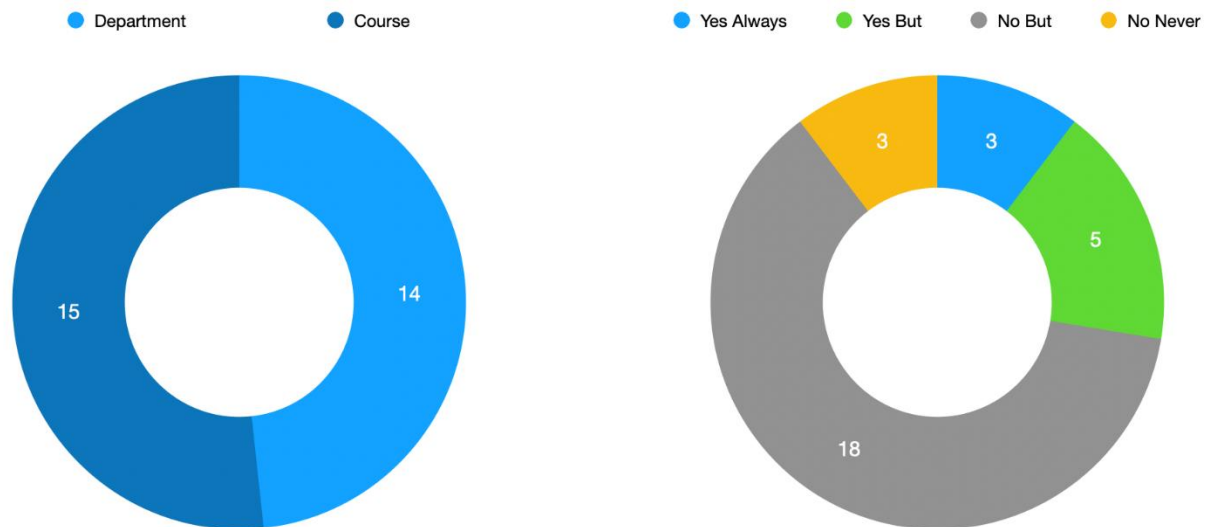


Figure 2: Summary of analyzed statements with institutional policies removed, showing (left) the policy level and (right) the restriction level of the policy. The total count was $N = 29$.

LLM analysis as described in the methodology section identified key themes among the policies, generating the labels provided in the list below. The parenthetical quantity represents the number of policies that addressed the theme, as identified by the LLM.

- Disclosure & Citation expectations (17)
- Academic integrity / plagiarism / Honor Code (17)
- Allowed study-support uses (10)
- Prohibited “submission-replacing” uses (16)
- Privacy and data security (3)
- Scope limits by task type (3)

The identified themes capture the breadth and complexity of the issue and the difference in approaches to providing guidance with many policies defining when AI tools can be used with others defining contexts or scopes where they may not be used. The bulk of the policies addressed academic integrity and disclosure themes, which aligns with the source policies provided being predominantly course-level or department-level policies. Allowable uses tended to focus on themes of building understanding or debugging, while disallowed uses tended to focus on authorship. The institutional level policies included the broader theme of privacy / data security. It is also interesting to consider that the LLM did not connect the academic integrity or disclosure themes to 14 of 31 policies, perhaps providing some insight into challenges with ambiguity and misinterpretation that instructors and students face when navigating expectations of the modern learning environment.

The prompt resulted in the LLM identifying six groupings or “stances” represented by the policies, along with a 7th category for institutional-level policies. This is in contrast to the four categories suggested by Stanford and manually labeled by the authors. The stances and the number of policies mapped to each stance are provided below. The stance names are exactly as generated by the LLM, and the associated sub-bullet is the descriptive, clarifying explanation of the stance, verbatim, as provided by the LLM.

- A. Zero-use ban (total prohibition) (4)
 - Most stringent; any use is misconduct
- B. Ban with narrow allowances (9)
 - Functionally a ban on GenAI for content/coding, but allows mechanical support (proofing, citations)
- C. Permission required / opt-in by instructor (9)
 - Local control, assignment-by-assignment; places burden on students to ask/verify
- D. Limited study-support use (2)
 - AI can teach (explanations, debugging guidance), but cannot author
- E. Allowed with disclosure; learning aid only (4)
 - AI use is allowed/encouraged as a learning tool, but the final work must represent student understanding; disclosure/citation required
- F. Fully encouraged/integrated (with reflection) (1)
 - Use is permitted at all stages (brainstorm → polish) with reflection/disclosure and no wholesale copy-paste
- U. University meta-guidelines (2)

A mapping of the categories by level of restriction that identifies the level of the policy and frequency of occurrence is shown in Figure 3.

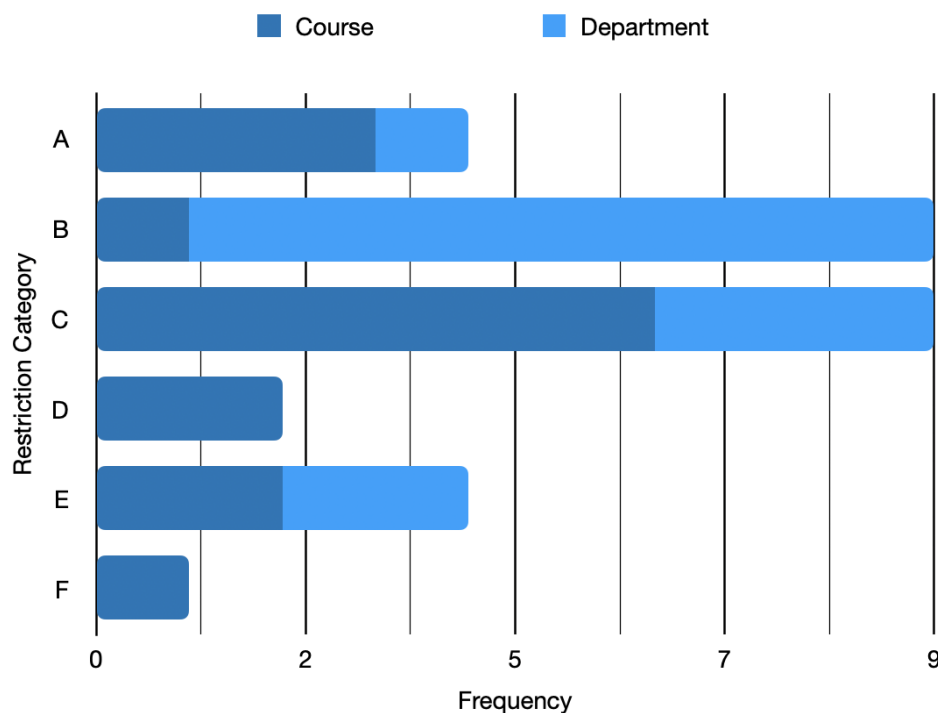


Figure 3: Number of policies categorized into levels A→F, with color indicating policy-level.

Upon analyzing the categorization, it is interesting to note that the two endpoints comprised of a total ban and fully encouraged use were preserved, but the categories or stances related to limited or restricted use were expanded from a paradigm that included either “default no, with limited pre-approved exceptions” or “default yes, with limited defined prohibitions or exceptions” to a categorization that looked at the restrictions more contextually (e.g. study support use, usage with disclosure).

The first difference worth exploring is at the extreme limit, with the LLM identifying four total prohibitions, while the authors labeled three statements in that category. An excerpt of the additional statement that the LLM labeled as a total prohibition is below.

“In general, I do NOT want you using Chat GPT or other AI generated work. If AI or other computer-generated writing method is used, then you must note it at the top of your submission...”

The authors viewed and coded this statement as a “No But” policy, interpreting the restriction to be default no use, with no allowed use for authorship, but some allowed uses when explained. However, it is also easy to see that the language is direct and emphatic about no use of AI generated work, explaining the LLM categorization and highlighting how ambiguity in policy statements presents many challenges for learners and instructors.

The two statements that the authors labeled as “Yes Always” that were not categorized this way by the LLM were:

“Cadets are welcome to utilize Generative AI, but assistance must be documented...In other words, if you use AI tools in your work, you must cite it... Do Not: Upload sensitive, proprietary, CUI or classified data into an AI tool. Think carefully about uploading any data that you did not generate or that is not a publicly available data set.”

“In this course, use of generative AI tools is permitted for all assignments, given that the student is using the AI tool to contribute to their learning of the course content, but not to replace original work. You must properly acknowledge when and how AI tools have been used in your submissions.”

The authors categorized both of the statements above as policies falling into the “Yes Always” category, while the LLM used the directives related to acknowledgement and disclosure to categorize the policies into a new category based on required disclosure.

In analyzing how policy level relates to the restrictions, 8/14 departmental policies were categorized as “ban with narrow allowances,” with 8/9 policies labeled in the category originating as departmental policies. In looking at the course level policies, 6/15 course policies were labeled as “permission required / opt-in by instructor”, with 6/9 policies in that category originating at the instructor / course level. No departmental policies fell within the “fully encouraged” or “limited study-support use” categories, while course-level policies spanned all 6 categories, indicating a wider variability of stance and interpretation among course-level policies as opposed to those implemented at a departmental level.

The predominate operating norm is to proceed with caution, aware of restrictions, with 18 policies falling within the “ban with narrow allowances” or “permission required / instructor opt-in” categories. The creation of the separate category for the institutional policies was an interesting development. Upon further exploration, the nature of those policies tended to be broader, encompassing all aspects of the institution, and the references to safety, privacy, and standards were differentiators that likely led to the LLM removing those policies from categorization by restriction.

As the authors examined the policies, one near-consensus norm existed, requiring disclosure of GenAI usage. Among policies crafted to provide guiderails regarding the context in which GenAI could be used, authorship versus assistance was the key distinguishing factor. Many of the course and departmental policies did not address privacy of information or standards for how attribution could / should be made, suggesting opportunity for university level policies to emphasize standards with respect to these concerns. With the uniformity of the requirement for attribution, the LLM level policy category should likely be removed / flattened, reducing the classifications to 5. It is interesting to consider the merits of expanding the suggested development of policy guidance and categorization into additional categories beyond the four suggested by Stanford. Is there value in taking the “Yes But” and No But” groupings and adding

resolution by moving to categories of “Ban with narrow allowances,” “Instructor Permission Required / Opt-In,” and “Limited Study Support / Learning Aid Use.”

Lastly, it is worth noting that the policies submitted to the LLM were not labeled by the associated discipline of the department or course. However, the LLM synthesis results suggested that policies from disciplines where languages are studied tended to trend to be the strictest. This was particularly interesting to the authors as they had identified the same trend from examining the collection, but the authors felt there was not enough of a sampling of policy statements to draw that conclusion. However, the discipline-specific tendencies inferred by the LLM based on the contexts and examples within the policies, without being explicitly provided disciplinary labeled data, was intriguing.

Conclusions

Manual analysis of the syllabi from SMCs and SAs indicated that 90% of the courses surveyed allowed some level of AI use, compared with 87% for the LLM analysis. Moving forward, a larger dataset including syllabi from multiple disciplines could enable deeper analysis using the proposed modeling techniques.

The analysis validated the challenges that students and instructors face with sensemaking in a rapid changing environment related to the use of GenAI in higher education. The key themes of disclosure, academic integrity, plagiarism, allowable use contexts, prohibited uses, and data privacy and security were readily identified among many of the policies, aligning with key guidance from many sources on best practices for course policy development.

The analysis also highlighted opportunities for future exploration and development in this domain. As students try to navigate this complex environment, alignment of policies and language detailing the permissible contexts and restricted uses would be helpful. Given the number of policies categorized as “Yes But” and “No But” or “ban with narrow allowances” and “permission required”, there is great opportunity for instructors to adopt standardized language from provided policy templates that could improve coherence and understanding for the parties involved. In the environment of the institutions where the policies were collected, most policies fell into a stance where the default position was no use of generative AI for graded work with limited exceptions for certain contexts or tools (“No But” or “Ban with Narrow Allowances” + “Permission Required”).

The ambiguity in AI use policy language observed by the authors and the challenge of clearly identifying and categorizing policies present sensemaking challenges for students and instructors. Informed by the experience described in this paper, the authors intend to develop future attitudinal surveys for both students and instructors to better inform and guide the development of policies. Future surveys could focus on Likert scale attitudinal items to understand students’ anxiety levels, clarifying communication behavior, and the impact on assignment time and effort based on example AI policy and policy language. Based on the

literature regarding common uses for AI in the higher education classroom and in writing, the authors could develop a series of descriptive scenarios of how AI was used on an assignment and use the survey to have respondents indicate whether that usage would be allowed for each of a series of policy statements. Additionally, for each scenario, the survey could explore whether the respondent feels that use case is “right”. Based on examples from this effort, confusion around acknowledgement / attribution, and bans on AI for authorship or content creation while allowing use for self-improvement, learning, and content editing were key areas to explore. Example scenarios could involve descriptions of students using AI for “learning / study” and explore whether the students feel each use adheres to various policy statements, and whether the example acknowledgement / attribution statements adhere to the policies. This would allow further explanation and understanding of policy language that is more confusing or ambiguous, help identify policies that are less ambiguous and suggest how they could be used as a model, and help explore whether the perceived ambiguities or gray areas with the policy statements stem from the language use or one’s own belief in the validity of or right to enforce a given policy.

If one considered Yes, Always and No, Never as the traffic green light and red light, respectively, one could consider the most common policy stances—Yes, But and No, But—as flashing yellow lights and flashing red lights, respectively. Imagine navigating our world where the predominant traffic control at each intersection was either flashing yellow or flashing red lights. As students, instructors, policymakers, and scholars work to develop operating norms related to GenAI, there is great need and opportunity for better signaling and increased clarity with respect to how one can safely proceed.

References

- [1] D.R.E. Cotton, P.A. Cotton, and J.R. Shipway, “Chatting and cheating: Ensuring academic integrity in the era of ChatGPT,” *Innovations in Education and Teaching International*, vol. 61, no.2, 2023. Available: <https://doi.org/10.1080/14703297.2023.2190148>
- [2] E. Kasneci *et al.*, “ChatGPT for good? On opportunities and challenges of large language models for education,” *Learning and Individual Differences*, vol. 103, pp. 102274, 2023. Available: <https://doi.org/10.1016/j.lindif.2023.102274>
- [3] K. Bittle and O. El-Gayar, “Generative AI and Academic Integrity in Higher Education: A Systematic Review and Research Agenda,” *Information*, vol. 16, no. 4, pp. 296, 2025. Available: <https://doi.org/10.3390/info16040296>
- [4] Y. An, J.H. Yu, and S. James, “Investigating the higher education institutions’ guidelines and policies regarding the use of generative AI in teaching, learning, research, and administration.” *Int J Educ Technol High Educ*, vol. 22, no. 10, 2025. Available: <https://doi.org/10.1186/s41239-025-00507-3>

[5] J. Kim *et al.* “Examining Faculty and Student Perceptions of Generative AI in University Courses.” *Innov High Educ*, vol. 50, pp. 1281–1313, 2025. Available: <https://doi.org/10.1007/s10755-024-09774-w>

[6] R. Sah, R., *et al.* “Generative AI in higher education: student and faculty perspectives on use, ethics, and impact.” *Issues in Information Systems*, vol. 20, no. 2, pp. 372–386, 2025. Available: https://doi.org/10.48009/2_iis_129

[7] Stanford University, “Creating your course policy on AI.” Accessed January 21, 2026. Available: <https://teachingcommons.stanford.edu/teaching-guides/artificial-intelligence-teaching-guide/creating-your-course-policy-ai>

[8] U.S. Department of Education, “Artificial Intelligence and the Future of Teaching and Learning.” 2023. Available. <https://www.ed.gov/sites/ed/files/documents/ai-report/ai-report.pdf>