

Beyond Optimism: An Initial Experience Report from a Broad-Access AI for Business and Society Course

Yang Liu
yliu8@nccu.edu
Computer Information Systems

Siobahn Day Grady
sday@nccu.edu
Information Science

Tianduo Zhang
tianduo.zhang@nccu.edu
Mass Communication

North Carolina Central University
Durham, NC, 27710, USA

Abstract

Artificial intelligence literacy courses for students outside computer science must support both operational confidence and critical judgment. This initial experience report examines a 16-week undergraduate AI for Business and Society course at a public historically Black university. The evaluation compares an anonymous first-week survey ($n = 65$) with an unlinked post-course survey ($n = 43$) from Spring 2026. The instrument measured perceived understanding, trust, governance, workforce expectations, societal impact, and open-ended concerns and hopes. Post-course respondents reported higher self-rated understanding than first-week respondents (4.27 versus 3.65); this was the only one of 11 focal items that remained statistically distinguishable after Holm correction (unadjusted $p = .0001$, adjusted $p = .001$, Cliff's delta = .42). Other differences, including stronger support for regulation and lower personal career anxiety, remain descriptive. The report contributes transferable lessons for information systems and computing educators: connect technical ideas to consequential organizational cases, distinguish confidence from competence, treat evaluation design as course infrastructure, and assess AI literacy through multidimensional and performance-based evidence.

Keywords: artificial intelligence literacy, information systems education, responsible AI, curriculum design, course evaluation, HBCU

Beyond Optimism: An Initial Experience Report from a Broad-Access AI for Business and Society Course

Yang Liu, Siobahn Day Grady, Tianduo Zhang

1. INTRODUCTION

Artificial intelligence (AI) now shapes organizational decisions in hiring, marketing, finance, cybersecurity, customer service, healthcare, and public administration. Information systems graduates and other business students therefore need more than familiarity with popular tools. They need to recognize when AI is being used, understand basic capabilities and limitations, question outputs, identify affected stakeholders, and reason about appropriate oversight. In this paper, AI refers broadly to data-driven prediction, machine learning, natural language processing, computer vision, and responsible deployment. The course was not organized as a generative-chatbot or prompt-engineering course, although students' everyday encounters with systems such as ChatGPT formed part of the surrounding context.

This broader conception complicates course evaluation. A successful course need not make students uniformly more favorable toward AI. Learning about model limitations, data quality, algorithmic bias, privacy, labor displacement, and unregulated deployment may increase confidence in some applications while reducing generalized trust in others. Responsible AI education consequently emphasizes informed judgment rather than optimism alone (Dignum, 2021; Miao et al., 2024). In this paper, differentiated judgment refers to the ability to distinguish among applications, articulate both benefits and risks, and identify conditions for responsible use. The present survey offers only indirect evidence of that capability; a revised performance-based design is proposed to measure it directly.

The course examined here, AI for Business and Society, was offered through the business school of a public historically Black university in the southeastern United States. It welcomed students from any major and required no programming, mathematics, or prior AI coursework. The design combined AI foundations, hands-on activities, organizational applications, workforce preparation, and responsible technology use. Although students came from several disciplines, information technology and cybersecurity were the largest groups. This mixture made the course particularly relevant to information systems educators seeking to broaden access without reducing AI to tool demonstrations.

The evaluation addresses three questions:

RQ1. How did perceived understanding differ between the first-week and post-course respondent groups?

RQ2. How did the groups' views of trust, governance, work, and societal impact differ?

RQ3. What design and evaluation lessons can information systems and computing educators transfer to other broad-access AI courses?

The paper is framed as an initial experience report. Spring 2026 survey data describe respondent cohorts, while the central contribution is a critical account of what the course design and evaluation made visible, what they could not establish, and how the next iteration will be strengthened.

2. RELATED WORK AND CONCEPTUAL FRAMING

Long and Magerko (2020) define AI literacy through competencies that include recognizing AI, understanding how systems learn, interpreting data, evaluating outputs, and considering ethical implications. Ng et al. (2021) organize related competencies into knowing and understanding AI, using and applying AI, evaluating and creating with AI, and addressing ethical issues. A review of AI literacy in higher and adult education found a field still taking shape and called for clearer definitions, stronger evidence about training, and better validated assessment tools (Laupichler et al., 2022). Together, these frameworks imply that frequent tool use or confidence alone cannot establish literacy.

Higher education institutions increasingly treat AI literacy as an across-the-curriculum responsibility rather than a computing-major specialty. Southworth et al. (2023) connect such efforts with workforce

readiness in multiple disciplines. Candon et al. (2025) describe a course for nontechnical learners seeking to become safe and informed AI users without advanced mathematics or programming. Biswas et al. (2025) and Xu et al. (2026) report university-wide one- and three-credit designs that combine technical foundations with disinformation, employment, and ethics. Bhattacharya et al. (2024) describe design-based AI education at a minority-serving institution using disciplinary modules, surveys, focus groups, and iterative revision. These general-education and broad-access models are not confined to HBCUs. The present course complements them through a 16-week business-school design centered on organizational cases and responsible deployment; the HBCU setting is context, not the sole basis of contribution.

Constructive alignment provides a useful design principle: intended learning outcomes, learning activities, and assessment evidence should support the same capabilities (Biggs, 1996). Course learning was assessed through quizzes and examinations, spreadsheet and data reports, a guided notebook model, written case analyses, a reflective essay, and a final presentation. These artifacts required students to explain concepts, interpret evidence, analyze affected stakeholders, and justify safeguards. The research study, however, relied on an anonymous perception survey that could not be linked to grades or artifacts. Course assessment and study evidence were related in topic but not integrated at the data level. Table 1 makes that alignment, and the remaining gaps, explicit.

3. COURSE DESIGN AND TEACHING CONTEXT

The course ran for 16 weeks in Fall 2025 and Spring 2026; the evaluation reported in this paper is limited to Spring 2026. Weeks 2-5 introduced AI, its history and subfields, data reasoning, correlation and causation, and linear regression through small Excel and Python exercises. Weeks 6-11 addressed supervised and unsupervised learning, natural language processing, computer vision, and neural networks through recommendation, fraud detection, spam filtering, sentiment analysis, facial recognition, and autonomous-vehicle examples. A guided Jupyter Notebook project asked students to implement and run a simple model. Later modules examined biased hiring, fairness and accountability, AI careers, skill development, and changing work. The semester concluded with a reflective essay and final presentation. No single textbook was prescribed; instructor-curated readings, cases, demonstrations, datasets, software documentation, and guided notebooks were used across both offerings.

3.1 Access Without Removing Technical Work

Removing prerequisites made the course accessible, but it did not eliminate technical work. The progression began with descriptive statistics and regression in familiar spreadsheets, moved to small Python exercises, and then used a guided notebook for the model project. This sequence provided several entry points and kept syntax from becoming the sole gatekeeper. The trade-off was selective depth: the course emphasized how data, objectives, and evaluation shape AI systems rather than the mathematical coverage of a computer science machine-learning course.

3.2 Consequential Cases as Bridges

Familiar applications linked technical concepts to organizational consequences. Recommendation systems illustrated prediction and personalization; facial recognition made error distribution and surveillance concrete; automated hiring connected classification with organizational accountability; and autonomous vehicles connected object detection with safety. The same scenarios appeared in the survey, improving curricular relevance but creating a measurement risk. A changed response could reflect familiarity with the example, a different social judgment, or a different respondent group. The scenarios therefore describe the response profile; they are not treated as module-level learning outcomes.

3.3 Ethics Across the Sequence

The original sequence placed its most explicit treatment of bias, fairness, and accountability in Week 12. In retrospect, that organization can make ethics appear to be a final qualification added after technical material. A stronger design would introduce stakeholder impact and data limitations with the first predictive example and revisit them in recommendation, vision, hiring, and workforce modules.

Responsible deployment would then operate as a recurring analytic lens rather than a standalone ethics week.

Course Goal	Representative Experiences	Evaluation Evidence
Explain AI concepts and limitations	Foundations; data reasoning; guided notebook model	Self-rated understanding; future scenario performance task
Evaluate applications in context	Recommendation, facial recognition, and autonomous-vehicle cases	Application judgments; societal-impact items
Reason about responsible deployment	Hiring-bias case; privacy, fairness, and accountability discussions	Trust, regulation, hiring, inequality, and open responses
Connect AI with business and work	Fraud and recommendation examples; careers module; final presentation	Business-change, employment, career-anxiety, and tool-use items

Table 1. Alignment among course goals, learning experiences, and evaluation evidence. Survey evidence reflects perceptions and judgments, not objective mastery.

4. EVALUATION DESIGN

4.1 Survey Instrument

The 20-question Student Perceptions of Artificial Intelligence instrument covered trust and regulation, business and work, societal impact, personal experience, general attitudes, tool use, and academic background. Eleven focal statements used a five-point agreement scale. Other questions addressed prior AI training, frequency and types of tool use, three application scenarios, and overall emotional orientation. Two open prompts asked about students' greatest concern and greatest hope. Appendix A provides an item-level map and response formats.

Concepts and scenarios were adapted in part from a national study of public attitudes toward AI (Rainie et al., 2022) and supplemented with course-specific questions about regulation, hiring, career obsolescence, inequality, tool use, and perceived understanding. The adapted instrument measures perceptions rather than objective knowledge and does not inherit psychometric validation from the source. Candidate multi-item composites had inadequate pre-survey reliability, including Cronbach's alpha values of .31 for concern and .58 for optimism, so items were analyzed separately. Q14—"I have a good understanding of what AI is and how it works"—is also double-barreled and is treated only as a broad self-assessment, not a validated knowledge measure.

4.2 Participants and Administration

The Spring 2026 first-week survey, administered January 13-22 before substantive instruction, received 65 responses. The post-course survey, administered April 23-May 2, received 43 responses. Item-valid sample sizes ranged from 63 to 65 and from 39 to 43, respectively. Respondents were predominantly information technology and cybersecurity majors, with academic classifications distributed from first year through senior year. Approximately 30% of first-week and 28% of post-course respondents reported previous AI-focused coursework or training.

A Fall 2025 pilot used a retrospective end-of-term then-test rather than a true first-week baseline. Because that administration is not comparable with the Spring design and is vulnerable to recall, anchoring, and response-shift bias (Howard & Dailey, 1979), Fall responses were excluded from all results reported here. The failed administration is retained only as a design lesson: baseline timing must be specified before a course begins and cannot be reconstructed after instruction.

4.3 Analysis and Interpretation Boundaries

Likert responses were coded from 1 to 5. We report item-valid means, agreement percentages, two-sided Mann-Whitney U tests with Holm correction across 11 items, and Cliff's delta. The survey did not use anonymous linkage codes. Some students may have responded at both time points, but their records

cannot be matched. Consequently, the tests' independence assumption may be violated, and the analysis cannot estimate within-student change. Inferential statistics are retained as an exploratory check on descriptive differences, not as confirmatory evidence of course effects. A single coder used a keyword-assisted scheme for Spring open responses and checked proposed themes against the full text; quotations are illustrative rather than comparative evidence.

5. RESULTS

5.1 Perceived Understanding

Table 2 reports the Spring 2026 first-week and post-course distributions. Post-course respondents reported higher self-rated understanding than first-week respondents: the mean for 'I have a good understanding of what AI is and how it works' was 4.27 post-course and 3.65 first week. Agreement was 82.9% and 66.2%, respectively. This was the only focal item that remained statistically distinguishable after Holm correction (unadjusted $p = .0001$, adjusted $p = .001$, Cliff's delta = .42). Because Q14 is double-barreled and the surveys were unlinked, the result concerns broad perceived understanding among respondent groups. It does not demonstrate technical mastery or individual learning.

Item	Construct	First Week Mean (SD)	Post Mean (SD)	% Agree First-Post	Delta	p (Holm)
Q1	Trust AI fairness	3.09 (1.16)	3.02 (1.35)	44.6-39.5	-.03	.796 (1.00)
Q2	Support regulation	3.14 (1.26)	3.35 (1.40)	36.9-55.8	.10	.356 (1.00)
Q3	Privacy concern	4.22 (0.96)	3.79 (1.28)	86.2-67.4	-.18	.099 (.888)
Q4	AI changes business	4.57 (0.73)	4.12 (1.43)	93.8-76.2	-.08	.396 (1.00)
Q5	Net job creation	3.09 (1.13)	3.00 (1.23)	32.3-38.1	-.03	.767 (1.00)
Q6	Career obsolescence worry	3.32 (1.09)	3.00 (1.25)	50.8-38.1	-.15	.184 (1.00)
Q7	AI hiring good idea	2.48 (1.22)	2.62 (1.34)	25.0-38.1	.05	.660 (1.00)
Q8	Widens inequality	3.57 (0.95)	3.44 (1.10)	44.6-56.1	.00	.973 (1.00)
Q9	Solves societal problems	3.48 (1.03)	3.80 (1.01)	52.3-73.2	.20	.065 (.652)
Q14	Understand AI	3.65 (0.96)	4.27 (1.14)	66.2-82.9	.42	.0001 (.001)
Q16	More good than harm	3.19 (1.03)	3.28 (1.12)	39.7-46.2	.06	.596 (1.00)

Table 2. Spring 2026 first-week and post-course responses (unlinked). Delta is Cliff's delta. Mann-Whitney p values are followed by Holm-adjusted values. Item-valid n = 63-65 first week and 39-43 post-course. SD: Standard Deviation.

5.2 A Descriptive Profile of Judgment

No attitude item survived multiplicity correction. Several descriptive differences are nevertheless useful for course revision. Agreement that AI could help solve important social problems was 52.3% in the first-week group and 73.2% in the post-course group (unadjusted $p = .065$, adjusted $p = .652$, Cliff's delta = .20). Support for government regulation was 36.9% and 55.8%, while trust that AI systems can make fair and unbiased decisions was 44.6% and 39.5%. Concern that AI could make one's own career obsolete was 50.8% and 38.1%. The broad judgment that AI will do more good than harm changed little, at 39.7% and 46.2% agreement.

This profile is compatible with more differentiated judgment: respondents could express optimism about socially useful applications while remaining wary of unfair decisions, job loss, and weak oversight. The design does not establish that interpretation. Unknown overlap, attrition, and differences in respondent composition can produce the same pattern. For this reason, the descriptive profile is used to identify questions for future assessment rather than to claim attitudinal change.

5.3 What Students Hoped for and Worried About

Spring open responses provided contextual examples but were coded by one researcher and are not treated as comparative evidence. Post-course respondents commonly connected practical benefits with social risks. One student wanted AI 'to do small jobs like organizing data.' Another warned that AI 'could replace jobs and be used in unfair or biased ways, especially if it's not properly regulated.' These quotations illustrate the kind of conditional judgment the course sought to encourage; they are not representative and cannot establish change.

6. LESSONS FOR IS AND COMPUTING EDUCATORS

6.1 Teach Through Consequential Organizational Cases

Hiring, surveillance, recommendation, fraud, misinformation, healthcare, and environmental costs connect technical choices to institutional responsibility. Prompts should ask students to identify a system's target, data limitations, affected stakeholders, and conditions for acceptable use. Cases are particularly useful in information systems courses because they join technical feasibility with governance, organizational process, and accountability.

6.2 Assess Confidence and Competence Separately

Tool use and self-rated understanding are valuable but incomplete outcomes. Students may become more confident without being able to explain a failure mode or justify a deployment decision. Conversely, students may become more cautious as their knowledge improves. Course evaluation should therefore retain a small set of perception items while adding scenario-based tasks that assess technical explanation, evaluation of evidence, stakeholder analysis, and justification of safeguards.

6.3 Treat Evaluation Design as Course Infrastructure

A true first-week baseline, an anonymous linkage code, explicit semester identifiers, and completion checks are inexpensive when planned in advance. They are costly or impossible to reconstruct later. The excluded Fall pilot demonstrates that having a survey instrument is not enough; administration timing and data architecture determine what questions the evidence can answer. Graded artifacts likewise become research evidence only when consent, de-identification, alignment, and linkage are designed before the course begins.

6.4 Keep AI Literacy Multidimensional

Separate questions about understanding, trust, regulation, work, social impact, and application conditions revealed tensions that a single favorability score would have concealed. Open responses then showed how students connected benefits with harms. This multidimensionality should be preserved, but future studies should designate a small set of primary outcomes in advance and treat the remainder as exploratory.

6.5 Integrate Responsible AI from the First Technical Example

Responsible AI should not appear only after students have learned the technical material. Data provenance, error distribution, stakeholder impact, and accountability can be introduced alongside regression and classification and revisited as applications become more complex. This design also makes ethical reasoning assessable throughout the semester rather than only in a final reflection.

7. REVISED EVALUATION BLUEPRINT

The next iteration will collect the baseline during the first class meeting and use a short anonymous code generated by each student. The perception battery will be reduced to a prespecified core covering understanding, fairness, regulation, workforce expectations, and social impact. The double-barreled understanding item will be split into 'I can explain what AI is' and 'I can explain at a high level how an AI or machine-learning system learns from data.' Candidate multi-item constructs will undergo cognitive interviewing, item-total analysis, and factor testing; a reliability coefficient of at least .70 will be required before a composite is used for exploratory group comparison.

Perceptions will be paired with parallel performance tasks. At entry, students might examine an AI-supported hiring scenario and identify the system target, likely data limitations, affected stakeholders, and one condition for responsible use. At the end of the course, a structurally similar scenario in healthcare or credit would be scored with the same rubric. The rubric would separately score technical explanation, evaluation of evidence, identification of harms, and justification of safeguards. Descriptive distributions and changes in rubric dimensions would then be reported alongside perceptions. Using a parallel but unfamiliar case reduces recognition of classroom examples as an alternative explanation.

Open responses will remain because fixed-choice items cannot show how students connect benefits and harms. Future qualitative analysis will use a documented codebook, two independent coders, and agreement checks before wave differences are interpreted. A small set of interviews could then probe why a student accepts an application under some conditions but rejects it under others. Linked surveys, performance tasks, and qualitative follow-up would convert the present cohort-level description into a stronger account of learning and judgment.

8. LIMITATIONS AND ETHICAL STATEMENT

The Spring surveys were anonymous and unlinked, with 65 first-week and 43 post-course responses. Attrition, item nonresponse, or differences in respondent composition could produce apparent differences. Possible repeat respondents also weaken the independence assumption of the exploratory tests. Without a comparison group, the study cannot separate instruction from external events or support causal claims. The adapted instrument had weak candidate-composite reliability, contained a double-barreled self-assessment item, and included no objective learning measure. Course grades and artifacts were not included in the approved de-identified research dataset and therefore cannot be reported or linked to survey responses. Subgroup comparisons are especially fragile because of small, shifting samples.

Qualitative responses were unevenly available and coded by one researcher without inter-rater reliability. The excluded Fall retrospective baseline is not equivalent to a first-week measure and is not analyzed. Finally, one convenience sample at one HBCU, composed largely of information technology and cybersecurity majors, cannot represent other institutions, disciplines, or broad-access courses. The institutional context broadens the evidence base but does not by itself establish that findings are distinctive to HBCUs.

The surveys were part of routine course assessment. Participation was voluntary and anonymous, and no identifying information is reported. The authors' Institutional Review Board determined that secondary analysis of the existing, fully de-identified data was Not Human Research under 45 CFR 46.102(e). Automatically logged IP-address and coarse-location fields were removed before analysis and are absent from shared materials.

9. CONCLUSION

In this initial Spring 2026 experience report, post-course respondents described greater understanding of AI, while views of governance, work, and social impact remained mixed. Only the self-rating difference was statistically robust, and the unlinked design prevents a claim of individual learning or technical mastery. The more durable contribution is methodological and pedagogical. Broad-access AI education benefits from a technical progression through familiar tools, consequential organizational cases, recurring responsible-AI analysis, multidimensional outcomes, true baselines, anonymous linkage, and

direct performance measures. For information systems educators, these practices help move AI literacy beyond tool familiarity toward reasoned decisions about when and how AI should be used.

10. ACKNOWLEDGEMENTS

The authors acknowledge support from the Duke Power Endowed Professorship at the School of Business, North Carolina Central University; the IBM-HBCU Quantum Center Faculty Fellowship; and the Institute for Artificial Intelligence and Emerging Research (IAIER) Seed Grant.

11. REFERENCES

- Bhattacharya, S., Czejdo, B., Zulli, R. A., & Smith, A. A. (2024). Enhancing AI education at an MSI: A design-based research approach. *Proceedings of the AAAI Symposium Series*, 3(1), 467-472. <https://doi.org/10.1609/aaais.v3i1.31258>
- Biggs, J. (1996). Enhancing teaching through constructive alignment. *Higher Education*, 32(3), 347-364. <https://doi.org/10.1007/BF00138871>
- Biswas, J., Fussell, D., Stone, P., Patterson, K., Procko, K., Sabatini, L., & Xu, Z. (2025). The essentials of AI for life and society: An AI literacy course for the university community. *Proceedings of the AAAI Conference on Artificial Intelligence*, 39(28), 28973-28978. <https://doi.org/10.1609/aaai.v39i28.35166>
- Candon, K., Georgiou, N. C., Ramnauth, R., Cheung, J., Fincke, E. C., & Scassellati, B. (2025). Artificial intelligence for future presidents: Teaching AI literacy to everyone. *Proceedings of the AAAI Conference on Artificial Intelligence*, 39(28), 28988-28995. <https://doi.org/10.1609/aaai.v39i28.35168>
- Dignum, V. (2021). The role and challenges of education for responsible AI. *London Review of Education*, 19(1), 1-11. <https://doi.org/10.14324/LRE.19.1.01>
- Howard, G. S., & Dailey, P. R. (1979). Response-shift bias: A source of contamination of self-report measures. *Journal of Applied Psychology*, 64(2), 144-150. <https://doi.org/10.1037/0021-9010.64.2.144>
- Laupichler, M. C., Aster, A., Schirch, J., & Raupach, T. (2022). Artificial intelligence literacy in higher and adult education: A scoping literature review. *Computers and Education: Artificial Intelligence*, 3, 100101. <https://doi.org/10.1016/j.caeai.2022.100101>
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1-16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- Miao, F., Shiohira, K., & Lao, N. (2024). *AI competency framework for students*. UNESCO. <https://doi.org/10.54675/JKJB9835>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- Rainie, L., Funk, C., Anderson, M., & Tyson, A. (2022). *How Americans think about artificial intelligence*. Pew Research Center. <https://www.pewresearch.org/internet/2022/03/17/how-americans-think-about-artificial-intelligence/>
- Southworth, J., Migliaccio, K. W., Glover, J., Glover, J., Reed, D. M., McCarty, C., Brendemuhl, J., & Thomas, A. C. (2023). Developing a model for AI across the curriculum: Transforming the higher education landscape via innovation in AI literacy. *Computers and Education: Artificial Intelligence*, 4, 100127. <https://doi.org/10.1016/j.caeai.2023.100127>

Xu, Z., Procko, K., Munje, M., Patterson, K., Sabatini, L., Biswas, J., & Stone, P. (2026). The essentials of AI for life and society: A full-scale AI literacy course accessible to all. *Proceedings of the AAAI Conference on Artificial Intelligence*, 40(47), 40730-40736.
<https://doi.org/10.1609/aaai.v40i47.41521>

APPENDIX A

Student Perceptions of Artificial Intelligence Instrument Map

Table A1 provides an item-level map of the 20-question instrument. The 11 focal agreement items are stated in full; the remaining questions are identified by content and response format. Q13 contained three application-scenario subprompts. The instrument was designed for course assessment rather than as a validated unidimensional scale.

Table A1. Item-level content and response format.

Item	Content	Format / role
Q1	I trust AI systems to make fair and unbiased decisions.	5-point agreement; focal
Q2	The government should regulate AI to protect the public.	5-point agreement; focal
Q3	I am concerned about the privacy risks associated with AI.	5-point agreement; focal
Q4	AI will significantly change how businesses operate.	5-point agreement; focal
Q5	AI will create more jobs than it eliminates.	5-point agreement; focal
Q6	I am concerned that AI could make my career obsolete.	5-point agreement; focal
Q7	Using AI in hiring decisions is a good idea.	5-point agreement; focal
Q8	AI is likely to widen social and economic inequality.	5-point agreement; focal
Q9	AI can help solve important social problems.	5-point agreement; focal
Q10	Previous AI-focused coursework or training.	Background; categorical
Q11	Frequency of AI-tool use.	Use frequency; categorical
Q12	Types of AI tools used.	Select all that apply
Q13	Judgments about three AI application scenarios.	Three scenario subprompts
Q14	I have a good understanding of what AI is and how it works.	5-point agreement; focal
Q15	Overall emotional orientation toward AI.	Categorical orientation
Q16	Overall, AI will do more good than harm.	5-point agreement; focal
Q17	What is your greatest concern about AI?	Open response
Q18	What is your greatest hope for AI?	Open response
Q19	Academic major or field of study.	Background; categorical
Q20	Academic classification.	Background; categorical

The wording of Q14 combines two claims—knowing what AI is and knowing how it works—and is therefore interpreted only as a broad self-assessment. The next instrument will split these claims, retain a smaller set of prespecified perception outcomes, and develop short multi-item constructs through cognitive interviewing, item analysis, and factor testing. Reliability of at least .70 will be required before a composite is used for exploratory group comparison; items will otherwise remain separate.