

Workplace Priorities for Entry-Level AI Competencies: A Preliminary Study

Amin Abassi-Pooya
aabassipooya@astate.edu

Kelly E. Fish
kfish@astate.edu

Paula Ruby
pruby@astate.edu

Information Systems and Business Analytics Department
Neil Griffin College of Business
Arkansas State University
Jonesboro, AR 72401 USA

Abstract

Artificial Intelligence (AI) is reshaping entry-level business work and increasing the need for competencies that combine effective AI application, critical evaluation, and responsible judgment. This study pretests a workplace-oriented instrument measuring the importance that mid-level professionals assign to AI competencies for entry-level business roles. Exploratory analysis yielded a multidimensional structure encompassing *AI Business Application*, *AI Risk Appraisal*, and *AI Governance*. Respondents placed particular emphasis on recognizing and responding to AI-related risks, while application and governance also emerged as important, complementary expectations. The measurement evidence for *AI Risk Appraisal* was less stable than for the other domains, indicating that this part of the instrument requires refinement and further testing. For employers, the framework can clarify expectations for hiring, onboarding, and responsible AI use. For information systems and business educators, it suggests that preparation for AI-enabled work should integrate practical tool use with verification, reliance decisions, ethical and privacy awareness, policy compliance, stakeholder consideration, and human accountability. The framework can guide learning objectives and performance-based assessments, while future research should validate and refine the measure across broader workplace contexts.

Keywords: AI competency, AI literacy, entry-level employment, business education, AI governance, curriculum development

Workplace Priorities for Entry-Level AI Competencies: A Preliminary Study

Amin Abbasi-Pooya, Kelly E. Fish and Paula Ruby

1. INTRODUCTION

Generative Artificial Intelligence (AI) is reshaping both the performance of business work and the capabilities organizations expect from entry-level employees. Many routine assignments through which new hires traditionally learn a profession, including preparing initial drafts, summarizing information, reconciling records, conducting basic analyses, and producing standard reports, are increasingly open to AI-enabled automation. The broader employment effects are not yet clear: The Economist (2026) characterizes early labor-market evidence as mixed, whereas Edmondson and Chamorro-Premuzic (2025) warn that removing junior roles could erode the pipeline through which organizations develop experienced professionals and future leaders. Even when entry-level positions are retained, however, their responsibilities are evolving. Drawing on more than one billion job advertisements, PwC (2026) reports that skill requirements were changing more than twice as rapidly in AI-exposed occupations as in less-exposed work and that highly exposed junior positions were much more likely to demand senior-level capabilities. Together, this evidence shifts the central question from whether AI will eliminate entry-level work to how it is reallocating routine execution and professional judgment within these roles. In this study, entry-level business roles refers broadly to early-career positions in functions such as operations, marketing, finance, and human resources rather than to a single occupation.

As routine execution becomes more automated, competencies that enable employees to use AI while preserving professional judgment become increasingly valuable. Manno (2025) observes that employers are moving beyond static knowledge toward an adaptive ability to solve problems with digital tools, verify outputs, and interpret them in context. In business practice, this capability requires users to select appropriate AI applications, formulate and refine instructions, incorporate outputs into workflows, and recognize situations in which human expertise should prevail over an automated recommendation. It also entails judging output accuracy and fitness for purpose while complying with organizational policy, privacy requirements, ethical standards, and responsibilities to stakeholders. This broader conception accords with AI-literacy research, which extends beyond tool operation to encompass understanding, application, critical evaluation, and ethical navigation (Almatrafi et al., 2024; Long & Magerko, 2020). It also reflects Chiu's (2025) distinction between foundational AI literacy and the real-world proficiency represented by AI competency. Consequently, readiness for AI-enabled work involves not only technical familiarity but also the capacity to apply AI to business tasks, appraise its risks, and recognize the governance obligations of responsible use. We therefore examine how strongly mid-level professionals value these dimensions of AI capability for entry-level business work.

Current AI-literacy frameworks offer a strong basis for specifying these capabilities, but they leave the workforce-priority question examined in this study largely unresolved. Almatrafi et al. (2024) organize AI literacy into six domains, among them using and applying AI, evaluating AI, and navigating AI ethically. The Scale for the Assessment of Non-Experts' AI Literacy assesses technical understanding, critical appraisal, and practical application in diverse nonexpert populations (Laupichler et al., 2023), while Ng et al. (2024) developed a self-report measure for secondary-school students encompassing affective, behavioral, cognitive, and ethical dimensions. Recent measures focused on generative AI further link literacy to privacy-protection and information-verification behaviors (Yan et al., 2026). Taken together, these studies establish the multidimensional character of AI literacy but provide limited evidence about which specific competencies employers emphasize when considering candidates for entry-level business roles. Instruments designed for students or general nonexperts may not adequately capture the task requirements, accountability structures, and organizational controls that govern AI use in the workplace.

This study addresses that gap by pretesting an instrument that captures the importance mid-level professionals assign to AI skills, given their likely involvement in hiring, supervising, and collaborating

with entry-level business employees. Candidate competencies are derived from prior literature, and exploratory principal components analysis and reliability analysis are used to examine the 13 retained items. A within-respondent comparison then assesses the relative importance of the resulting competency dimensions. By centering employer judgments, the study evaluates a preliminary workplace-oriented measure spanning three interrelated requirements of competent AI use: applying AI to business tasks, recognizing and responding to AI risks, and governing AI responsibly. The framework is intended to support information systems and business programs in aligning learning objectives, instructional activities, and assessment practices with the evolving demands of AI-enabled entry-level work.

2. LITERATURE REVIEW

As Artificial Intelligence becomes more deeply embedded in organizational activity, competency in its use is increasingly important for employees and future business professionals. The literature distinguishes AI literacy, the foundational knowledge needed to understand AI, from AI competency, the ability to mobilize that knowledge when addressing real-world problems. Competent use therefore goes beyond familiarity with AI tools: it requires knowing AI's capabilities and limitations, applying it to suitable tasks, critically assessing AI-generated content, and acting responsibly and ethically. Chiu (2025) characterizes AI competency as practical proficiency in using, engaging with, interacting with, developing, or managing AI systems in real-world settings. Likewise, Wu et al. (2026) argue that non-technical business professionals need AI literacy not so that they can build models as programmers do, but so that they can adopt, integrate, and use AI critically and responsibly.

Within this multidimensional perspective, the three study constructs capture distinct yet mutually reinforcing dimensions of AI competency in business contexts. *AI Business Application* represents the practical, task-focused dimension and includes identifying business problems for which AI performs well, refining prompts, incorporating AI tools into business activities, and combining AI outputs with human judgment. *AI Risk Appraisal* represents the evaluative dimension, encompassing recognition of business problems for which AI performs poorly, awareness of misleading output, and decisions about when such output should be rejected. *AI Governance* represents the responsible-use dimension and covers AI's probabilistic operation, bias recognition, independent verification, ethical and privacy risks, human accountability, organizational AI-use policies, and effects on stakeholders. Together, these dimensions correspond to Almatrafi et al.'s (2024) broad conception of AI literacy, particularly its Use and Apply, Evaluate, and Navigate Ethically constructs.

2.1 AI Business Application

AI Business Application captures the capacity to convert a general understanding of AI into productive action within business processes. This capacity begins with achieving task-AI fit: managers must specify the business need, assess the model's capabilities and limitations, and choose an appropriate tool for the situation (Robertson et al., 2024). Generative AI can assist with translation, summarization, classification, information provision, analytics, process automation, and customer service, yet the benefits of such uses depend on the task, context, and content involved (Ferraro et al., 2024; Sundberg & Holmström, 2024). Annapureddy et al. (2025) similarly identify knowledge of AI capacities, limitations, and appropriate contexts as core generative-AI competencies. Accordingly, the first item captures the ability to match business needs to AI capabilities: *Ability to identify business problems where AI performs well*.

Identifying an appropriate use is not sufficient unless AI can be incorporated into the work being performed while preserving human contribution. Retkowsky et al. (2024) found that early ChatGPT adopters progressed from experimentation to "intertwinement," in which the tool became part of routine knowledge-work activities such as information search, brainstorming, drafting, refinement, and checking. Sundberg and Holmström (2024) likewise recommend scaling from bounded applications to broader workflow integration, while Ferraro et al. (2024) show how AI can automate routine customer-service processes and transfer complex cases to employees. Chiu (2025) and Wu et al. (2026) further emphasize context-specific human-AI collaboration and human-in-the-loop oversight. These findings support two items: *Ability to integrate AI tools into business tasks* and *Ability to combine AI-generated outputs with human judgment*.

Effective application also requires users to improve their instructions when initial output does not meet the intended objective. Robertson et al. (2024) describe prompt engineering as an iterative process involving context, structure, evaluation, and refinement, including increasing clarity, removing ambiguity, modifying phrasing or context, and specifying the desired response format. Sundberg and Holmström (2024) similarly argue that task-specific detail can increase output accuracy and value. Knoth et al. (2024) provide empirical support by finding that higher-quality prompt engineering predicted higher-quality large-language-model output, and Annapureddy et al. (2025) identify prompting as a distinct competency for tailoring outputs to specific objectives. Thus, the final item captures iterative improvement of instructions: *Ability to refine AI prompts to improve output quality*.

Although *AI Business Application* is broad and its components are distinct, the items form a reasonable sequence of AI-enabled business work: identifying a suitable problem and selecting an appropriate AI tool, integrating the tool into the business task, refining prompts to improve output, and combining the output with human judgment. The construct therefore captures a connected application process rather than asserting that the component skills are interchangeable or that every employee must master each one to the same degree; expected proficiency may vary across entry-level roles.

2.2 AI Risk Appraisal

AI Risk Appraisal denotes the capacity to identify circumstances in which AI may be unreliable, evaluate the resulting risk, and respond appropriately. AI-literacy scholarship emphasizes that competence requires an understanding of AI's strengths and limitations rather than the mere ability to operate a tool (Laupichler et al., 2023; Long & Magerko, 2020; Ng et al., 2024). In the same vein, Annapureddy et al. (2025) identify knowledge of generative AI's capabilities, limitations, and appropriate contexts as a core competency. In business settings, complex, ambiguous, sensitive, or empathy-dependent problems may require human involvement (Ferraro et al., 2024), and Wu et al. (2026) specifically identify ambiguity, ethical dilemmas, and tasks requiring genuine understanding or empathy as areas in which AI struggles. These findings justify the item *Ability to identify business problems where AI performs poorly*.

Users must also recognize that fluent and convincing AI output can be inaccurate, unsupported, or fabricated. Ji et al. (2023) document hallucination in natural-language generation, while Walters and Wilder (2023) show that ChatGPT can produce fabricated and erroneous citations. Shin et al. (2025) further show that users vary in how they understand and process generative-AI misinformation. Consistent with this evidence, Annapureddy et al. (2025) state that large-language-model output can combine factual material with false or fabricated claims and therefore identify output assessment as a core competency. To capture recognition of this risk, we use the item *Awareness that AI can produce misleading results*.

Recognition must be followed by an appropriate reliance decision. Human-AI decision-making research defines appropriate reliance as accepting accurate and useful AI advice while rejecting incorrect or unsuitable advice, rather than trusting or distrusting AI indiscriminately (Lee & See, 2004; Schemmer et al., 2023). This judgment is important because users may over-rely on AI even when its recommendations conflict with contextual evidence or their own judgment (Klingbeil et al., 2024). Wu et al. (2026) likewise specify that critical evaluation includes deciding whether an AI output should be accepted, revised, or rejected. Therefore, the third item captures the user's reliance decision: *Ability to decide when AI outputs should be rejected*.

2.3 AI Governance

Mäntymäki et al. (2022) define AI governance as the organizational rules, practices, processes, and tools that align AI use with strategic aims, legal obligations, and ethical principles. For individual users, governance begins with understanding how AI generates information and applying appropriate controls to assess that information. Annapureddy et al. (2025) characterize generative AI as a statistical model, while Wu et al. (2026) explain that such systems rely on probabilistic methods to generate human-like responses without possessing genuine understanding or intent. Chiu (2025) identifies critical evaluation of output accuracy, bias, and relevance as part of AI literacy, while Ashok et al. (2022) treat fairness and bias as governance concerns. Finally, output assessment requires

checking claims against trusted or independent sources (Annapureddy et al., 2025; Metzger, 2007; Wineburg & McGrew, 2019; Yan et al., 2026). These findings support three items: *Understanding that AI relies on probabilities, not true reasoning*; *Ability to recognize bias in AI-generated information*; and *Ability to verify AI output accuracy using alternative sources*.

Formal governance also requires employees to recognize ethical and privacy risks and comply with organizational controls. Janssen et al. (2020) conclude that trustworthy AI requires data stewardship, information-sharing controls, risk-based governance, and compliance with legal and ethical requirements. Fjeld et al. (2020) find broad agreement that responsible AI should protect privacy and provide affected people with control over their data. Reviews of AI ethics identify fairness, accountability, autonomy, privacy, transparency, and societal impact as recurring principles and risks (Ashok et al., 2022; Jobin et al., 2019), while Birkstedt et al. (2023) and Papagiannidis et al. (2025) emphasize governance processes such as oversight, auditing, and risk or impact assessment. This literature supports the items *Awareness of potential ethical risks associated with using AI*; *Understanding data privacy related to AI tools*; and *Understanding the importance of following organizational AI use policies*.

AI governance also locates accountability with people and requires attention to those affected by AI adoption. Fjeld et al. (2020) identify accountability, human control, and professional responsibility as common responsible-AI themes, including the expectation that consequential decisions remain subject to human review. Birkstedt et al. (2023) likewise place oversight and accountability within AI-governance processes. Beyond the organization itself, Clarke (2019) argues that AI risk assessment should consider stakeholder interests, and Birkstedt et al. (2023) distinguish internal stakeholders from external groups such as clients, regulators, civil society, and affected individuals. Miller (2022) further classifies AI project stakeholders and highlights potential harm to passive stakeholders. These studies justify the final two items: *Awareness that responsibility for AI-assisted decisions lies with humans* and *Awareness that AI adoption can affect employees, customers, and other stakeholders*.

3. THE STUDY

We pretested the survey with an online panel of 110 respondents who identified as mid-level professionals. The instrument contained 15 substantive AI-competency items rated on a five-point importance scale. The items were intentionally tool-neutral and referred to generative AI broadly rather than to specific branded systems. Although *AI Business Application* includes competencies involving tool use, other items address misleading output, privacy, ethics, organizational policy, human accountability, and stakeholder effects; tying the full instrument to particular products would therefore narrow its intended scope and reduce its durability as tools change. We screened each respondent's pattern across all 15 substantive items and operationalized straight-lining as selecting the same scale response for every item. Twenty-six cases met this rule and were excluded, leaving 84 responses for analysis.

All 84 retained respondents provided complete numeric responses for Items 1 through 15, and the same 84 cases were used in every analysis. We first conducted exploratory factor analysis on all 15 items. After examining the rotated loading pattern, we removed Item 1 because of insufficient loading separation and low communality (below .4) and Item 8 because of substantial cross-loadings. We then repeated the factor, reliability, and linear mixed model analyses using the 13 retained items.

After Items 1 and 8 were removed, the 13-item analysis produced an overall Kaiser-Meyer-Olkin statistic of .768, and Bartlett's test remained significant, $\chi^2(78) = 363.228$, $p < .001$. The Kaiser criterion indicated three components (eigenvalues = 4.180, 2.155, and 1.352). The fixed three-component solution explained 59.14% of total variance. After rotation, *AI Governance* accounted for 24.39%, *AI Business Application* for 20.06%, and *AI Risk Appraisal* for 14.69%. Absolute primary loadings ranged from .534 to .800.

In the final 13-item solution, *AI Business Application* comprised Items 2, 5, 6, and 7; *AI Risk Appraisal* comprised Items 3, 4, and 10; and *AI Governance* comprised Items 9, 11, 12, 13, 14, and 15. Although Item 4 exhibited a secondary loading, it was retained because its primary loading (.587) and communality (.522) were acceptable and because the item provided important conceptual coverage of

AI Risk Appraisal. Cronbach's alpha was .813 for all 13 remaining items, .817 for *AI Governance*, .802 for *AI Business Application*, and .525 for *AI Risk Appraisal*.

Pretest Item	Statement	Factor 1	Factor 2	Factor 3
1	Understanding that AI relies on probabilities, not true reasoning.	-.028	.296	.427
2	Ability to identify business problems where AI performs well.	.803	.280	.004
3	Ability to identify business problems where AI performs poorly.	.288	.538	.115
4	Awareness that AI can produce misleading results.	.051	.582	.427
5	Ability to refine AI prompts to improve output quality.	.731	.059	.047
6	Ability to integrate AI tools into business tasks.	.807	-.222	.211
7	Ability to combine AI-generated outputs with human judgment.	.799	.095	.120
8	Ability to recognize bias in AI-generated information.	.218	.434	.496
9	Ability to verify AI output accuracy using alternative sources.	.173	.364	.629
10	Ability to decide when AI outputs should be rejected.	-.087	.776	.107
11	Awareness of potential ethical risks associated with using AI.	.050	-.307	.769
12	Understanding data privacy related to AI tools.	.070	.069	.750
13	Awareness that responsibility for AI-assisted decisions lies with humans.	.045	.381	.649
14	Understanding the importance of following organizational AI use policies.	.146	.296	.656
15	Awareness that AI adoption can affect employees, customers, and other stakeholders.	.105	.143	.715

Table 1. Kaiser-Normalized Varimax-Rotated Component Loadings for All 15 Pretest Items
Note. N = 84. Values are loadings from a fixed three-component principal components analysis with Kaiser-normalized Varimax rotation; absolute loadings of .40 or higher are bold. Factor 1 = *AI Business Application*; Factor 2 = *AI Risk Appraisal*; Factor 3 = *AI Governance*.

To compare factor ratings within respondents, we averaged each participant's retained item responses within *AI Business Application*, *AI Risk Appraisal*, and *AI Governance*. We estimated a random-intercept linear mixed model with factor as a fixed effect and respondent as the grouping factor, thereby accounting for the three repeated factor scores contributed by each participant (Cnaan et al., 1997; Muhammad, 2023). *AI Business Application* served as the mixed-model reference category.

Construct	Number of Items	Mean	Standard Deviation
1 (<i>AI Business Application</i>)	4	3.976	.736
2 (<i>AI Risk Appraisal</i>)	3	4.286	.606
3 (<i>AI Governance</i>)	6	4.123	.635

Table 2. Descriptive Statistics for the Three Retained AI Competency Factors

As shown in Table 2, *AI Risk Appraisal* had the highest observed mean ($M = 4.286$, $SD = .606$), followed by *AI Governance* ($M = 4.123$, $SD = .635$) and *AI Business Application* ($M = 3.976$, $SD = .736$). The random-intercept linear mixed model yielded a significant omnibus fixed effect of factor, Wald $\chi^2(2) = 13.128$, $p = .001$. With *AI Business Application* as the reference category, *AI Risk Appraisal* was .310 points higher ($SE = .085$, $z = 3.622$, $p < .001$, 95% CI [.142, .477]), whereas *AI Governance* was .147 points higher ($SE = .085$, $z = 1.718$, $p = .086$, 95% CI [-.021, .314]).

Because *AI Business Application* was the reference category, the reported coefficients did not directly estimate the contrast between *AI Risk Appraisal* and *AI Governance*. Their descriptive ordering should therefore not be interpreted as a statistically tested difference.

4. DISCUSSION OF RESULTS

The initial 15-item analysis was used only as an item-screening step, and Items 1 and 8 were removed before substantive interpretation. The final 13-item analysis produced substantively recognizable *AI Business Application*, *AI Risk Appraisal*, and *AI Governance* groups. Its three-component solution explained 59.14% of the variance, the Kaiser criterion recovered three

components, and absolute primary loadings ranged from .534 to .800. These results provide preliminary support for the intended 13-item structure.

Because the same 84 cases were used for item screening and estimation of the final component solution, the structure was not cross-validated and may reflect sample-specific patterns. The findings should therefore be interpreted as preliminary rather than as evidence of a finalized measurement model.

Important limitations remain. Cronbach's alpha was .813 for all 13 remaining items. At the construct level, alpha was .817 for *AI Governance*, .802 for *AI Business Application*, and .525 for *AI Risk Appraisal*. Internal consistency was therefore good for *AI Governance* and *AI Business Application* but poor for the three-item *AI Risk Appraisal* factor. The low coefficient may partly reflect the factor's short length and heterogeneous content, but it means that *AI Risk Appraisal* should not yet be treated as a well-established scale. Although retaining Item 4 was conceptually justified, its secondary loading and the low factor reliability indicate that the risk-appraisal items should be revised or expanded and re-evaluated in the main study.

The random-intercept linear mixed model identified an overall difference among the three factor means. *AI Risk Appraisal* was .310 points higher than *AI Business Application* and this contrast was statistically significant ($p < .001$). *AI Governance* was .147 points higher than *AI Business Application*, but that contrast did not reach the .05 level ($p = .086$). Because the reported model used *AI Business Application* as the reference category, it did not directly compare *AI Risk Appraisal* with *AI Governance*. Thus, the supported pairwise conclusion is that respondents rated risk-appraisal items more highly than business-application items; the descriptive ordering of all three competency domains requires confirmation after the *AI Risk Appraisal* items are strengthened.

5. CONCLUSIONS, IMPLICATIONS, AND FUTURE RESEARCH

This pretest provides preliminary evidence for a workplace-oriented framework of entry-level AI competency. After item screening, the retained items formed interpretable domains of *AI Business Application*, *AI Risk Appraisal*, and *AI Governance*, consistent with scholarship that distinguishes appropriate application, critical evaluation, and responsible use (Almatrafi et al., 2024; Chiu, 2025). The structure is promising rather than definitive: *AI Risk Appraisal* showed weaker internal consistency than the other domains. The instrument should therefore be treated as a foundation for refinement and further testing rather than as a validated scale.

For employers, the framework clarifies what competent AI use can involve in entry-level business roles. Job descriptions, hiring exercises, onboarding, and performance discussions can assess whether candidates and employees can match AI to appropriate tasks, refine and integrate AI outputs, recognize misleading results, make sound reliance decisions, and operate within policy, privacy, ethical, stakeholder, and accountability requirements. The significant contrast between *AI Risk Appraisal* and *AI Business Application* suggests that employers value judgment about AI failures in addition to operational fluency. Because the risk-appraisal measure requires refinement and the model did not directly compare *AI Risk Appraisal* with *AI Governance*, the observed domain ordering should guide further inquiry rather than serve as a fixed ranking.

For information systems and business educators, the implications are to integrate *AI Business Application*, *AI Risk Appraisal*, and *AI Governance* rather than teach tool use in isolation. Learning activities can require students to assess task-AI fit, refine prompts, incorporate AI into a workflow, verify material claims against authoritative sources, document accept, revise, reject, or escalate decisions, and analyze privacy, policy, ethical, stakeholder, and accountability issues. Assessment should make students' reasoning visible through annotated outputs, verification records, workflow maps, and decision rationales instead of evaluating only the polish of an AI-assisted product. This integration of application and judgment more closely reflects workplace expectations and can be evaluated through authentic, performance-based tasks (Miao et al., 2024).

Future revisions should use expert review to refine *AI Risk Appraisal* item wording, expand coverage of risk recognition, evaluation, and reliance decisions, and pilot the revised items before structural testing. The probabilistic-behavior item was removed from the retained solution; if reconsidered, it

should be evaluated as a foundational AI-literacy or risk-appraisal indicator rather than assigned a priori to *AI Governance*.

Future research should revise and expand *AI Risk Appraisal*, reassess its weaker items, and test whether the items removed during screening can be reworded to reduce cross-loadings. The planned final study will target approximately 500 respondents across industries, organizational levels, and geographic settings, allowing separate exploratory and confirmatory subsamples and comparison of plausible factor structures. A 7-point importance scale should be tested against the current response format to determine whether greater response differentiation improves reliability and factor separation. Finally, importance ratings should be paired with performance-based and longitudinal evidence to determine whether the competencies predict workplace performance and whether education or training interventions develop them over time. Because generative AI tools, capabilities, and workplace practices are evolving rapidly, items tied to particular products or current system features may quickly become outdated. Future versions of the instrument should therefore retain tool-neutral wording that emphasizes transferable competencies, behaviors, and professional judgment rather than proficiency with named platforms. The item pool should also be versioned and reviewed periodically, with additions, revisions, and retirements documented so that changes can be tracked. When substantive changes are made, the revised instrument should be piloted and its content coverage, factor structure, reliability, and comparability with earlier versions reassessed before scores are interpreted across time or workplace settings.

6. ACKNOWLEDGEMENTS

Generative AI tools supported drafting, language refinement, and manuscript formatting. The authors verified and remain responsible for all content.

7. REFERENCES

- Almatrafi, O., Johri, A., & Lee, H. (2024). A systematic review of AI literacy conceptualization, constructs, and implementation and assessment efforts (2019–2023). *Computers and Education Open*, 6, Article 100173. <https://doi.org/10.1016/j.caeo.2024.100173>
- Annapureddy, R., Fornaroli, A., & Gatica-Perez, D. (2025). Generative AI literacy: Twelve defining competencies. *Digital Government: Research and Practice*, 6(1), Article 13. <https://doi.org/10.1145/3685680>
- Ashok, M., Madan, R., Joha, A., & Sivarajah, U. (2022). Ethical framework for artificial intelligence and digital technologies. *International Journal of Information Management*, 62, Article 102433. <https://doi.org/10.1016/j.ijinfomgt.2021.102433>
- Birkstedt, T., Minkkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI governance: Themes, knowledge gaps and future agendas. *Internet Research*, 33(7), 133–167. <https://doi.org/10.1108/INTR-01-2022-0042>
- Chiu, T. K. F. (2025). AI literacy and competency: Definitions, frameworks, development and future research directions. *Interactive Learning Environments*, 33(5), 3225–3229. <https://doi.org/10.1080/10494820.2025.2514372>
- Clarke, R. (2019). Principles and business processes for responsible AI. *Computer Law & Security Review*, 35(4), 410–422. <https://doi.org/10.1016/j.clsr.2019.04.007>
- Cnaan, A., Laird, N. M., & Slasor, P. (1997). Using the general linear mixed model to analyse unbalanced repeated measures and longitudinal data. *Statistics in Medicine*, 16(20), 2349–2380. [https://doi.org/10.1002/\(SICI\)1097-0258\(19971030\)16:20%3C2349::AID-SIM667%3E3.0.CO;2-E](https://doi.org/10.1002/(SICI)1097-0258(19971030)16:20%3C2349::AID-SIM667%3E3.0.CO;2-E)
- Edmondson, A. C., & Chamorro-Premuzic, T. (2025, September 16). The perils of using AI to replace entry-level jobs. *Harvard Business Review*. <https://hbr.org/2025/09/the-perils-of-using-ai-to-replace-entry-level-jobs>

- Ferraro, C., Demsar, V., Sands, S., Restrepo, M., & Campbell, C. (2024). The paradoxes of generative AI-enabled customer service: A guide for managers. *Business Horizons*, 67(5), 549–559. <https://doi.org/10.1016/j.bushor.2024.04.013>
- Fjeld, J., Achten, N., Hilligoss, H., Nagy, A. C., & Srikumar, M. (2020). *Principled artificial intelligence: Mapping consensus in ethical and rights-based approaches to principles for AI*. Berkman Klein Center for Internet & Society, Harvard University. <https://dash.harvard.edu/entities/publication/739ea327-7c95-43d9-9fdc-a702bf420e40>
- Janssen, M., Brous, P., Estevez, E., Barbosa, L. S., & Janowski, T. (2020). Data governance: Organizing data for trustworthy artificial intelligence. *Government Information Quarterly*, 37(3), Article 101493. <https://doi.org/10.1016/j.giq.2020.101493>
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Chen, D., Dai, W., Chan, H. S., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248, 1–38. <https://doi.org/10.1145/3571730>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Klingbeil, A., Grützner, C., & Schreck, P. (2024). Trust and reliance on AI: An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, Article 108352. <https://doi.org/10.1016/j.chb.2024.108352>
- Knoth, N., Tolzin, A., Janson, A., & Leimeister, J. M. (2024). AI literacy and its implications for prompt engineering strategies. *Computers and Education: Artificial Intelligence*, 6, Article 100225. <https://doi.org/10.1016/j.caeai.2024.100225>
- Laupichler, M. C., Aster, A., Haverkamp, N., & Raupach, T. (2023). Development of the “Scale for the Assessment of Non-Experts’ AI Literacy” (SNAIL): An exploratory factor analysis. *Computers in Human Behavior Reports*, 12, Article 100338. <https://doi.org/10.1016/j.chbr.2023.100338>
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- 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>
- Manno, B. V. (2025, October 30). A new AI career ladder. *Stanford Social Innovation Review*. <https://doi.org/10.48558/68hp-1740>
- Mäntymäki, M., Minkinen, M., Birkstedt, T., & Viljanen, M. (2022). Defining organizational AI governance. *AI and Ethics*, 2, 603–609. <https://doi.org/10.1007/s43681-022-00143-x>
- Metzger, M. J. (2007). Making sense of credibility on the Web: Models for evaluating online information and recommendations for future research. *Journal of the American Society for Information Science and Technology*, 58(13), 2078–2091. <https://doi.org/10.1002/asi.20672>
- Miao, F., Shiohira, K., & Lao, N. (2024). *AI competency framework for students*. UNESCO. <https://doi.org/10.54675/JKJB9835>
- Miller, G. J. (2022). Stakeholder roles in artificial intelligence projects. *Project Leadership and Society*, 3, Article 100068. <https://doi.org/10.1016/j.plas.2022.100068>
- Muhammad, L. N. (2023). Guidelines for repeated measures statistical analysis approaches with basic science research considerations. *Journal of Clinical Investigation*, 133(11), e171058. <https://doi.org/10.1172/JCI171058>

- Ng, D. T. K., Wu, W., Leung, J. K. L., Chiu, T. K. F., & Chu, S. K. W. (2024). Design and validation of the AI literacy questionnaire: The affective, behavioural, cognitive and ethical approach. *British Journal of Educational Technology*, 55(3), 1082–1104. <https://doi.org/10.1111/bjet.13411>
- Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *The Journal of Strategic Information Systems*, 34(2), Article 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
- PwC. (2026, June 15). *2026 global AI jobs barometer: Two futures for jobs in an AI era*. <https://www.pwc.com/gx/en/services/ai/ai-jobs-barometer.html>
- Retkowsky, J., Hafermalz, E., & Huysman, M. (2024). Managing a ChatGPT-empowered workforce: Understanding its affordances and side effects. *Business Horizons*, 67(5), 511–523. <https://doi.org/10.1016/j.bushor.2024.04.009>
- Robertson, J., Ferreira, C., Botha, E., & Oosthuizen, K. (2024). Game changers: A generative AI prompt protocol to enhance human-AI knowledge co-construction. *Business Horizons*, 67(5), 499–510. <https://doi.org/10.1016/j.bushor.2024.04.008>
- Schemmer, M., Kühn, N., Benz, C., Bartos, A., & Satzger, G. (2023). Appropriate reliance on AI advice: Conceptualization and the effect of explanations. In *Proceedings of the 28th International Conference on Intelligent User Interfaces* (pp. 410–422). Association for Computing Machinery. <https://doi.org/10.1145/3581641.3584066>
- Shin, D., Koerber, A., & Lim, J. S. (2025). Impact of misinformation from generative AI on user information processing: How people understand misinformation from generative AI. *New Media & Society*, 27(7), 4017–4047. <https://doi.org/10.1177/14614448241234040>
- Sundberg, L., & Holmström, J. (2024). Innovating by prompting: How to facilitate innovation in the age of generative AI. *Business Horizons*, 67(5), 561–570. <https://doi.org/10.1016/j.bushor.2024.04.014>
- The Economist. (2026, January 29). *How big a threat is AI to entry-level jobs?* <https://www.economist.com/business/2026/01/29/how-big-a-threat-is-ai-to-entry-level-jobs>
- Walters, W. H., & Wilder, E. I. (2023). Fabrication and errors in the bibliographic citations generated by ChatGPT. *Scientific Reports*, 13, Article 14045. <https://doi.org/10.1038/s41598-023-41032-5>
- Wineburg, S., & McGrew, S. (2019). Lateral reading and the nature of expertise: Reading less and learning more when evaluating digital information. *Teachers College Record*, 121(11), 1–40. <https://doi.org/10.1177/016146811912101102>
- Wu, L., Yang, W., & Hu, Y. (2026). *Proposing an AI literacy framework for business education* [Preprint]. SSRN. <https://ssrn.com/abstract=6096280>
- Yan, W., Liu, Y.-L., Mamaeva, V., Dong, F., Tao, G., Li, R., & Yang, H. (2026). Generative AI literacy: Scale development and its influence on privacy protection behaviors and information verification behaviors. *Telecommunications Policy*, 50(2), Article 103117. <https://doi.org/10.1016/j.telpol.2025.103117>