



April 10, 2026

AI in Action

2026 Conference Proceedings

AI in Action is hosted by the
Departments of Data Analytics
& Statistics and
Technical Communication

Edited by Kim Sydow Campbell and Maia Pacini

University of North Texas
Denton, TX

Table of Contents

Welcome from the Conference Team.....	1
Teaching Algorithmic Auditing as User-Centered Practice.....	2
AI as Design Partner: Co-Creating an Immersive Technical Communication Simulation.....	4
Teaching SEO in the Age of AI: A Classroom Case for Understanding Generative Engine Optimization (GEO).....	5
Leveraging AI and Automation to Alleviate Instructor Burnout.....	7
Sillion: Context-Aware AI Teaching Assistants for Personalized High-School Instruction.....	9
Human-Centered and Trustworthy AI for Fair Environmental and Health Monitoring.....	11
Theoretical Framework on Knowledge Completeness & Consistency (Jagged Intelligence).....	12
Human Expertise Enabling the Transition from AI Users to AI Contributors through a University-Led AI Data Solutions Hub.....	13
The Autonomous Strategist: Agentic AI and MCP for Real-Time Inventory Intelligence.....	15
Multi-agent orchestration for a critical healthcare financial system.....	17
Coherence: Gamified Human-AI Collaboration for Executive Digital Transformation Training and Decision-Support.....	18
Can College Instructors Tell AI from Human Writing? Evidence from 400 Classification Trials of STEM Reports.....	20
Answering IBM's Call for a Systems Theory: Agentic AI as Open Systems.....	22
Comparing LLM and human financial decision-making.....	24
Detecting and Quantifying Hallucinations in LLM-Generated References.....	25
Author Index.....	27

Welcome from the Conference Team

Welcome to the AI in Action 2026 Conference Proceedings. Held at the University of North Texas in Denton, Texas. AI in Action brought together faculty, students, researchers, and practitioners to share innovative ideas, projects, and applications of artificial intelligence across disciplines.

The conference was created as a space for exploring practical uses of AI, emerging challenges, and opportunities for interdisciplinary collaboration. We are grateful to all presenters, authors, participants, and colleagues who contributed their time, expertise, and enthusiasm to the event.

We hope these proceedings serve not only as a record of AI in Action 2026, but also as a source of ideas for future research, teaching, and collaboration.

-The AI in Action Team

Dr. Kim Sydow Campbell, Department of Technical Communication

Dr. Orhan Erdem, Department of Data Analytics and Statistics

Dr. Zeynep Orhan, Department of Data Analytics and Statistics

Teaching Algorithmic Auditing as User-Centered Practice



Ashley Rea
Department of Technical
Communication
University of North Texas
ashley.rea@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Over more than five decades, technical communicators have weathered disruptions that have affected their employability and job description – from a focus on standardization, clarity, and precision in manuals and reports during the world wars (1940-50s), growth of structured authoring tools, single-sourcing, and automation (1980-90s), print to online documentation affecting access and interactivity (1990s), global offshoring of technical writing teams (2000s), natural language processing tools for machine translation, chatbots, and grammar/style checkers (2010s) to the most recent generative AI wave for automatic text generation affecting productivity, bias, and accessibility. In response, TC researchers consistently developed new strategies for writing and editing; including recently, writing with AI (Ranade, Saravia & Johri, 2024; Ranade, 2023), teaching about AI systems (Dux Speltz et al., 2022), and writing for AI and chatbots (Hocutt, Ranade & Verhulsdonck, 2022). While AI systems are the latest disruption affecting what we write, how we write, and who we write for, AI systems offer both opportunities and challenges for technical communicators. As Selber (2024) contends, “AI requires students to know more, not less, about technical communication.” This presentation provides specific guidelines to

work alongside AI for the next generation of technical communicators in the workplace.

Researchers have discussed two primary concerns around AI: first, co-designing content or working alongside AI for content development, and second, being mindful of algorithmic bias in AI that harms populations exposed to or consuming such content (Buolamwini & Gebru, 2018; Noble, 2018; Narayanan & Kapoor, 2024). As technical communicators, we value participatory design approaches that position users as co-designers. Extending those to identify the roles and limitations of algorithms like AI can ensure successful transitions for technical communicators in their positions. Additionally, because of our longstanding focus on user-centered design, informed by design justice principles (Costanza-Chock, 2020) and the social justice turn in TC, technical communicators are uniquely well-situated to develop methods that ensure AI transparency and algorithmic fairness in TC practices. To that end, this presentation offers ‘algorithmic auditing’ as a strategy for teaching both co-design and algorithmic evaluation in TC classrooms.

Algorithmic auditing can act as both a technical and social tool to examine data, power, context, and consequences and center affected communities whose voices must be included in co-design and

co-creation practices. The presentation shares Broussard's framework for Algorithmic Auditing adapted for TC practitioners and their roles. This strategy is paired with a heuristic for evaluating AI transparency (Carey et al., 2014) and an introduction to model cards. Model cards are a genre of documentation "designed to compile information about the training and testing of AI models, as well as the features and the motivations of a given dataset or algorithmic model" (Clavell, 2023) used by practitioners. Finally, the presentation concludes with a reflection on how to apply this framework in writing and communication classrooms.

Ultimately, this presentation contributes to scholarship in algorithmic accountability, leveraging technical communicators' expertise in user-centered design to provide "a networked account for a socio-technical algorithmic system, following the various stages of the system's lifecycle" (Wieringa qtd. in Bandy, 2021) and help technical communicators carve out new roles in their careers.

References

- Bandy, J. (2021). Problematic machine behavior: A systematic literature review of algorithm audits. *Proceedings of the acm on human-computer interaction*, 5(CSCW1), 1-34.
- Buolamwini, J., & Gebru, T. (2018, January). Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Conference on fairness, accountability and transparency* (pp. 77-91). PMLR.
- Carey, G., Crammond, B., & De Leeuw, E. (2015). Towards health equity: a framework for the application of proportionate universalism. *International journal for equity in health*, 14(1), 81.
- Clavell, G. G. (2024). Checklist for AI auditing. URL https://www.edpb.europa.eu/system/files/2024-06/ai-auditing_checklistfor-ai-auditing-scores_edpb-spe-programme_en.pdf. Accessed July.
- Costanza-Chock, S. (2020). *Design justice: Community-led practices to build the worlds we need*. The MIT Press.
- Hocutt, D., Ranade, N., & Verhulsdonck, G. (2022). Localizing content: The roles of technical & professional communicators and machine learning in personalized chatbot responses. *Technical Communication*, 69(4), 114-131.
- Hocutt, D., Ranade, N., & Verhulsdonck, G. (2022). Localizing content: The roles of technical & professional communicators and machine learning in personalized chatbot responses. *Technical Communication*, 69(4), 114-131.
- Narayanan, A., & Kapoor, S. (2025). AI as normal technology. *Knight First Amendment Institute*, 25.
- Noble, C. *The Politics of Information and Communication Technologies (ICTs), Artificial Intelligence (AI), and Implications for Social Work, Justice, and Advocacy*. In *AI and the Disruption of the Social* (pp. 60-71). Routledge.
- Ranade, N. (2023). AI for Editing. *The WAC Clearinghouse*. DOI: 10.37514/TWR-J.2023.2982.2.27
- Ranade, N., Saravia, M., & Johri, A. (2025). Using rhetorical strategies to design prompts: A human-in-the-loop approach to make AI useful. *AI & SOCIETY*, 40(2), 711-732.
- Speltz, E. D., Roeser, J., Bouwman, W., Chukharev, E., & Torrance, M. (2026). Writing traits: Examining the consistency of behavioral patterns in writers' composing processes. *Journal of Writing Research*.

AI as Design Partner: Co-Creating an Immersive Technical Communication Simulation



Eric Williams
Department of Technical
Communication
University of North Texas
eric.williams@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

What if AI could help faculty build ambitious learning experiences they couldn't create alone, not because they lack ideas, but because they lack time or technical skills? This presentation explores one instructor's experience using AI as a collaborative design partner to develop a four-week technical communication simulation for a technical writing course.

Students took on functional roles within a fictional game development studio facing a cascading crisis: a key employee's sudden public departure, technical setbacks, investor pressure, and community backlash. Each week introduced new complications that built on previous events; each required professional communication tailored to distinct audiences, purposes, and stakes. The simulation culminated in a final exam where students revised a flawed launch announcement—drawing on four weeks of accumulated context about the company, its stakeholders, and its challenges.

The presenter, who has no coding background, used iterative conversations with AI to brainstorm

narrative arcs, develop characters with distinct communication styles, and generate interactive artifacts including mock Slack threads, social media posts, and investor emails. The AI handled execution—HTML, styling, interactivity—while the instructor drove pedagogical and narrative decisions. The result was a cohesive, immersive experience that maintained alignment with standardized course objectives while increasing student investment.

Early student feedback highlighted the impact of this approach, with one student noting the simulation was “a first for me in the classroom environment” and describing the experience as “lots of fun and engaging.” Students reported feeling a sense of agency in the unfolding narrative, even as they practiced core professional communication skills.

Attendees will see examples from the simulation, learn about the collaborative design process between instructor and AI, and discuss how this model of human-AI partnership might support their own pedagogical experimentation, regardless of technical background.

Teaching SEO in the Age of AI: A Classroom Case for Understanding Generative Engine Optimization (GEO)



Jiaxin Zhang
Department of Technical
Communication
University of North Texas
jiaxin.zhang@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

The rapid rise of generative AI tools such as ChatGPT, Perplexity, and Gemini; and Google's AI Overviews is transforming how people access information online. Instead of navigating ranked lists of links, users increasingly encounter synthesized answers generated from multiple sources. This shift raises new questions about how digital content becomes visible and authoritative in search environments.

Traditional search engine optimization (SEO) focuses on improving website visibility through strategies such as keyword targeting, metadata, backlinks, and technical site structure. However, generative search systems introduce a different visibility dynamic: content may be selected, summarized, and cited directly within AI-generated responses (Chen et al., 2025; Manasa et al., 2025; Search Engine Land).

This presentation explores how SEO instruction can be updated to reflect these changes through the integration of Generative Engine Optimization (GEO) into technical and professional communication curricula. GEO emphasizes writing practices that make information legible and citable for generative AI systems, including definitional clarity, structured claims, visible evidence, and

clear signals of authority and credibility (Aggarwal et al., 2024; Chen et al., 2025; Lüttgenau et al., 2025).

The presentation introduces a classroom teaching case implemented in a junior-level content strategy course. Students participate in a comparative activity examining how the same query appears across multiple platforms, including a traditional search engine, Google AI Overviews, and generative AI tools. Working in small groups, students analyze differences in how information is defined, how sources are cited, and how authority is signaled across platforms.

By integrating SEO and GEO in the curriculum, this teaching case encourages students to critically examine how algorithmic systems shape information access and credibility online. Such instruction supports broader goals of digital literacy and algorithmic awareness in higher education (Atkins & Reilly, 2019; Muller et al., 2025).

References

- Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., & Deshpande, A. (2024). GEO: Generative Engine Optimization. arXiv preprint arXiv:2311.09735. <https://doi.org/10.48550/arXiv.2311.09735>

AI in Action 2026 sponsored by

The Departments of Technical Communication and Data Analytics & Statistics at the University of North Texas

- Atkins, A., & Reilly, C. (2019). Pedagogical strategies for integrating SEO into technical communication curricula. *Communication Design Quarterly*.
<https://doi.org/10.1145/3309578.3309585>
- Chen, M., Wang, X., Chen, K., & Koudas, N. (2025). Generative Engine Optimization: How to Dominate AI Search. *ArXiv*, abs/2509.08919.
<https://doi.org/10.48550/arxiv.2509.08919>
- Lüttgenau, F., Colic, I., & Ramirez, G. (2025). Beyond SEO: A Transformer-Based Approach for Reinventing Web Content Optimisation. *ArXiv*, abs/2507.03169.
<https://doi.org/10.48550/arxiv.2507.03169>
- Manasa, G., Raju, P., Reddy, S., Riyaz, S., & Nausheen, S. (2025). Optimizing for the Artificial Intelligence Driven Search Era: An Integrated Framework for Search Engine Optimization, Generative Engine Optimization, and Answer Engine Optimization. *International Journal of Scientific Research in Engineering and Management*.
<https://doi.org/10.55041/ijssrem53069>
- Muller, C., Sagot, S., & Domenget, J. (2025). Educational Implications of SEO Courses on Algorithm Awareness. *Education for Information*, 41, 199 - 226.
<https://doi.org/10.1177/01678329251323453>
- Search Engine Land. (n.d.). AI SEO. Retrieved January 7, 2026, from
<https://searchengineland.com/library/ai-seo>

Leveraging AI and Automation to Alleviate Instructor Burnout



Erin Friess
Technical Communication
University of North Texas
erin.friess@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Instructor burnout in higher education is all too real. Faculty report high levels of exhaustion tied not only to teaching and research expectations but also to the accumulation of administrative and emotional labor embedded in the profession (Pope-Ruark, 2022). In teaching, much of the strain faculty experience comes from forms of invisible labor, the routine but unacknowledged tasks required to keep courses running smoothly (Crain, Poster, & Cherry, 2016). Scholars have also noted that such hidden forms of work are frequently unevenly distributed and undervalued within institutions (Babcock, Recalde, Vesterlund, & Weingart, 2022). This work often takes the form of repeated clarification emails, troubleshooting software tools students are using, and mediating group conflicts. Individually these tasks seem minor, but collectively they consume substantial time and cognitive energy.

Artificial intelligence and automation tools offer an opportunity to reduce some of this invisible work. I propose to discuss three ways in which AI and automation can support instructors by reducing the routine coordination and troubleshooting work that often contributes to burnout.

Stop Asking Me focuses on shifting routine technical troubleshooting to generative AI. In many courses, students use a wide range of software

platforms. Instructors often spend significant time helping students navigate tool specific problems that are not central to the learning goals. An early semester assignment teaches students how to prompt generative AI tools to generate step by step instructions tailored to their specific software and version. This approach trains students to use AI as a first line of technical support, reducing the volume of repetitive troubleshooting questions directed at the instructor.

Stop Getting It Wrong uses AI to identify patterns of error in student work perhaps caused by holes in instruction. By analyzing assignment submissions or rubric feedback, AI tools can help instructors detect recurring problem areas and evaluate whether content delivery, assignment instructions, or rubrics need clarification. Rather than repeatedly addressing the same issues across multiple semesters, instructors can refine assignment design based on these insights.

Stop Arguing with Your Teammates focuses on group management through automation. Student team projects frequently require instructors to monitor collaboration and intervene when progress stalls. Using an automation tool such as Geekbot integrated into Microsoft Teams, instructors can implement asynchronous standup check-ins based on Scrum project management practices. These automated updates create transparency within

teams and allow instructors to quickly identify issues without manually collecting status reports.

These strategies demonstrate how AI and automation can help address the less visible aspects of teaching that contribute to instructor burnout. By reducing routine troubleshooting, repeated clarifications, and group coordination overhead, these tools allow instructors to focus more on the aspects of teaching that remain uniquely human.

References

- Babcock, L., Recalde, M., Vesterlund, L., & Weingart, L. (2022). *The No Club: Putting a Stop to Women's Dead-End Work*. Scribner.
- Crain, M., Poster, W. R., & Cherry, M. A. (2016). *Invisible Labor: Hidden Work in the Contemporary World*. University of California Press.
- Pope-Ruark, R. (2022). *Unraveling Faculty Burnout: Pathways to Reckoning and Renewal*. Johns Hopkins University Press.

Sillion: Context-Aware AI Teaching Assistants for Personalized High-School Instruction



Kai Gu
Honors College
Texas Academy of Mathematics
and Science
kaigu@my.unt.edu

Shreyes Suribhotla
Honors College
Texas Academy of Mathematics
and Science
shreyessuribhotla@my.unt.edu

Aleksander Stevens
Honors College
Texas Academy of Mathematics
and Science
aleksanderstevens@my.unt.edu

Bryan Wang
Honors College
Texas Academy of Mathematics
and Science
bryanwang2@my.unt.edu

Weishi Shi
Department of Computer
Science and Engineering
Active Machine Learning Lab
weishishi@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Educators routinely face a fundamental tension: the quantitative data available through learning management systems – grades, attendance, completion rates – captures what students produce, but not how they think, where they struggle, or what they feel. In a class of thirty or more, the one-on-one conversations that surface these qualitative signals are difficult to assess and act on. While tools exist to address tutoring or analytics individually, none bridge both layers into a unified, instructor-facing system organized around the two forms of context that matter most: the content of the course and the evolving state of each learner. This workshop presents Sillion, an open-source AI teaching assistant platform that closes this gap by pairing guided conversational learning with automated insight extraction, giving instructors a richer, more actionable picture of student learning.

Sillion assigns each class a dedicated AI teaching assistant. Rather than relying on static prompting alone, Sillion frames personalization as a continual

learning problem: accumulating structured memory from coursework and prior interactions to adapt future tutoring turns to each learner's evolving state and instructor feedback. Students engage through a chat interface – prompted to reflect, think critically, and work through their own reasoning rather than receive answers directly. Behind each conversation, two inference processes run in tandem. First, a large language model generates context-aware responses shaped by persistent pedagogical personas – Socratic Tutor, Encouraging Mentor, Assessment Helper – allowing the AI's behavior to align with individual teaching philosophies. Second, a parallel analysis pipeline extracts structured signals from every student message: sentiment, academic concepts discussed, areas of struggle, and expressed emotions. These signals are aggregated into a teacher-facing dashboard that surfaces engagement trends, flags students who may need intervention, and provides per-student conversation histories.

The platform integrates directly into Canvas LMS, indexing course modules and materials so that AI

AI in Action 2026 sponsored by

The Departments of Technical Communication and Data Analytics & Statistics at the University of North Texas

responses are grounded in actual class content rather than generic knowledge. An AI-assisted grading workflow further extends this integration: the system generates draft scores and written feedback for Canvas assignments, which instructors can review, adjust, and sync back to the gradebook. This keeps the human instructor in the loop while reducing repetitive evaluation overhead.

We evaluate Sillion through a mixed-methods approach across pilot course deployments. Quantitative measures include student engagement metrics – message frequency, session length, and return rate – alongside instructor time-on-task for grading and flagged interventions. Qualitative data is gathered through end-of-term surveys assessing student-perceived learning support and instructor-perceived utility, with conversation logs analyzed to surface recurring struggle patterns and inform iterative platform improvements.

During this workshop, attendees will see a live walkthrough of the teacher and student experience, including class creation, enrollment, AI-guided tutoring across pedagogical modes, conversation review, and Canvas integration, followed by a discussion of design decisions, practical limitations, and the ethical considerations of deploying conversational AI in classrooms.

Human-Centered and Trustworthy AI for Fair Environmental and Health Monitoring



Masoud Rostami
Department of Data Science
University of Texas at Arlington
masoud.rostami@uta.edu

Behnaz Balmaki
Department of Environmental
Science
Texas Women's University
bbalmaki@twu.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Artificial intelligence is increasingly used to support environmental and health monitoring, yet many AI systems remain poorly aligned with human priorities when data are scarce, imbalanced, or ecologically structured. In domains such as pollen monitoring, rare signals often carry disproportionate importance for public health, environmental assessment, and food security, yet conventional models optimized for overall accuracy systematically overlook them. This creates not only technical limitations but also trust and fairness concerns in real-world decision support. This work presents a human-centered and trustworthy AI framework for fair environmental and health monitoring, demonstrated through a case study on rare pollen analysis.

Using a curated microscopic pollen image dataset with severe class imbalance, we evaluate bias-aware deep learning strategies including transfer learning, metric-based few-shot learning, and conditional generative augmentation. Trustworthiness is operationalized through equitable performance across classes, robustness under data scarcity, and expert oversight of

synthetic data generation. Macro-averaged evaluation metrics are used to reflect human and ecological priorities rather than majority-class dominance.

Results show that while standard pretrained models achieve high overall accuracy, they frequently underperform on rare pollen types. In contrast, metric-based approaches significantly improve minority performance without sacrificing generalization, and biologically validated generative augmentation further enhances rare-signal detection. These findings demonstrate that effective environmental AI systems require deliberate human guidance in model design, evaluation, and validation.

Beyond environmental and allergy monitoring, this framework directly extends to honey quality assessment and food security, where rare pollen signatures are critical for authenticity verification and safety evaluation. By treating class imbalance as a meaningful signal rather than noise, this work illustrates how human-centered and trustworthy AI can support reliable decision-making in environmental, health, and food systems.

Theoretical Framework on Knowledge Completeness & Consistency (Jagged Intelligence)



Thuan Luong Nguyen
Department of Data Analytics
and Statistics
University of North Texas
thuan.nguyen@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

The rapid ascent of generative artificial intelligence (AI) and large language models (LLMs) has revealed a critical anomaly: systems demonstrate expert-level skills in complex domains yet fail at elementary tasks. This paper presents a comprehensive theoretical framework to explain the phenomenon of jagged intelligence, exemplified by the Wharton Paradox, the coexistence of graduate-level capabilities with elementary arithmetic failures within the same system.

Employing constructive theoretical synthesis, we introduce 19 foundational concepts, including intellectual capability level, knowledge completeness, and knowledge consistency. We identify key mechanisms including epistemic decoupling (advanced patterns without foundational support), knowledge density gradient (peaks in data-rich domains, valleys in data deserts), and probabilistic drift (context-dependent activation without semantic anchoring).

We propose two novel theories: theory I (the completeness of knowledge theory) posits that current AI training paradigms, characterized by stochastic knowledge aggregation, inevitably produce extreme knowledge gaps and epistemic decoupling in AI knowledge due to the absence of the prerequisite enforcement mechanisms inherent in human systematic knowledge scaffolding. Theory II (the AI knowledge consistency test theory) operationalizes these insights into a falsifiable testing protocol to distinguish artificial from organic intelligence. Empirical validation is provided through analysis of the Wharton Paradox and corroborated by authoritative assessments from industry leadership regarding the end of the age of scaling.

We conclude that knowledge completeness is a structural prerequisite for trustworthy AI, arguing that the path to artificial general intelligence (AGI) requires a fundamental paradigm shift from scaling capability to enforcing knowledge consistency.

Human Expertise Enabling the Transition from AI Users to AI Contributors through a University-Led AI Data Solutions Hub



Zeynep Orhan
Department of Data Analytics
and Statistics
University of North Texas
zeynep.orhan@unt.edu

Mehmet Orhan
Department of Data Analytics
and Statistics
University of North Texas
mehmet.orhan@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

The proposed University-led AI Data Solutions Hub establishes a structured environment where human expertise supports AI systems through supervised evaluation, testing, training data creation, and governance-oriented feedback. The hub organizes interdisciplinary contributors within the university to provide services such as annotation, validation, adversarial testing, bias analysis, rubric-based scoring, and documented system assessment. The model connects educational activities, research capacity, and industry collaboration while creating paid experiential opportunities for participants. The project positions the university as an active contributor to AI ecosystems rather than a passive consumer of technology.

Universities hold structural strengths that enable this transition. Academic institutions bring multidisciplinary expertise, continuous participant renewal, guided learning environments, and institutional oversight. Students, faculty, and staff represent a wide range of disciplinary perspectives and professional experiences that allow AI systems to be evaluated under realistic social, linguistic, and domain-specific conditions. Structured onboarding,

supervision, and quality assurance transform individual participation into coordinated human expertise. The university setting supports training, mentoring, and governance practices that commercial contractor platforms often struggle to sustain consistently.

Market demand demonstrates strong momentum for human-in-the-loop workflows. AI companies rely heavily on human feedback to improve model performance, evaluate risks, and adapt to rapidly changing tasks. Companies such as Scale AI and Mercor illustrate the growth of this sector through rapid expansion, significant investment rounds, and large-scale operational deployment. Public job postings reveal demand across technical and nontechnical fields, including healthcare, law, education, finance, and creative disciplines. Competitive compensation structures and referral incentives signal persistent shortages of qualified contributors and confirm the economic viability of organized human expertise.

The proposed hub generates mutual benefits for universities, industry partners, and society. Participants gain applied experience, professional skills, and income opportunities aligned with emerging labor markets. Industry partners receive access to coordinated human expertise supported

AI in Action 2026 sponsored by

The Departments of Technical Communication and Data Analytics & Statistics at the University of North Texas

by institutional quality standards and ethical oversight. Society benefits from improved accountability, reduced bias, and more transparent evaluation processes for AI systems. Ethical responsibility forms a central pillar of the model. Governance structures guide data handling, fairness assessment, evaluation criteria, and documentation practices to support responsible innovation.

The University-led AI Data Solutions Hub demonstrates how academic institutions can enable a transition from AI users to AI contributors. The model integrates education, research, workforce development, and revenue generation while reinforcing human expertise as a critical component of trustworthy AI systems.

The Autonomous Strategist: Agentic AI and MCP for Real- Time Inventory Intelligence



Deepeka Gurunathan
Department of Data Analytics &
Statistics
University Of North Texas
deepekagurunathan@my.unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

The Problem

Traditional dealership inventory management relies on reactive decision-making, where managers manually cross-reference regional market trends with local stock levels. This process is labor-intensive, prone to human error, and often results in inventory gaps—missed sales opportunities due to insufficient stock of high-demand vehicles. Industry estimates suggest these gaps can result in 15-20% revenue loss due to stockouts and misaligned inventory allocation.

Innovation

This project introduces the Intelligent Inventory Balancer, an autonomous strategist that utilizes Agentic AI orchestrated through n8n and the Model Context Protocol (MCP). By adopting MCP, the system replaces rigid, hard-coded pipelines with a modular universal adapter architecture. MCP enables the AI Agent to dynamically discover and invoke specific tools—such as querying a PostgreSQL inventory database or analyzing regional vehicle sales trends—only when the task context requires it.

Technical Methodology

The system is built on an n8n backbone, utilizing the MCP Server Trigger to expose local data sources as executable tools for a large language model (LLM) like Gemini 2.5 Flash.

1. **Local Context (PostgreSQL):** A dedicated MCP tool allows the agent to pull real-time stock counts for specific make and models (e.g., Toyota Camry, Honda Civic) in the Dallas, Texas area.
2. **Regional Intelligence:** The agent queries a market-analysis tool to extract used car sales volume by make, model, and segment (e.g., sedans, SUVs, trucks) to identify high-demand vehicles in the regional market.
3. **Autonomous Execution:** Instead of following a predefined path, the AI Agent uses the MCP Client to determine which data points are missing. If local stock is low while regional demand is high, the agent autonomously generates a strategic restock recommendation and delivers it via Gmail.

Impact and Future Trends

This research demonstrates that MCP significantly reduces context overload by ensuring the AI only processes the most pertinent data points. By decoupling the model's core intelligence from its

AI in Action 2026 sponsored by

The Departments of Technical Communication and Data Analytics & Statistics at the University of North Texas

knowledge sources, the system becomes highly scalable; new data sources can be plugged in via MCP without rewriting the core workflow.

This project showcases a future where human-AI collaboration is defined by autonomous agents that not only display data but also actively perceive, reason, and act to solve complex logistical challenges. The system enables faster decision-making, reduces cognitive load on inventory managers, and transforms raw data into strategic business intelligence.

Future work will explore multi-region optimization and integration with predictive demand forecasting models.

Key Contributions

Agentic AI Architecture: Demonstrates autonomous decision-making through tool-aware LLM agents.

MCP Integration: Showcases the Model Context Protocol as a scalable solution for connecting AI agents to enterprise data sources.

Real-World Application: Addresses a concrete business problem (inventory gaps) with measurable impact potential.

Modular Design: Enables rapid adaptation to new data sources and business contexts through the MCP framework.

Multi-agent orchestration for a critical healthcare financial system



Sonali Sabnam
Optum | Department of Data
Analytics & Statistics (Alum)
University of North Texas
sonalisabnam89@gmail.com

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Healthcare is one of the most complex businesses in today's world which involves humans and their well-being. The tiniest errors can have severe impacts on business and human life.

What looks like a sweet (often seamless) relationship between a provider and member on the outside is much more complicated inside the operations and management processes.

I had the privilege to work on one of such business processes which is tedious, cumbersome, and goes through multiple windows for the simplest update.

Current Process

The financial tagging system goes through multiple checks from multiple teams to identify any changes to the financial files and make the updates at the correct rule and rule structure.

Solution

Automate this process with a multi-agent orchestration system using LangGraph and HITL (human in the loop).

Details

Technology stack: Python, LangGraph, Fast API

Architecture

Data Analyst Agent: Identifies the new and updated group policy numbers and generates a report. It also identifies the details like medical combination or behavioral combination.

Natural Language Chat: Users can ask questions about the Group Policies and details. The chat response is generated via a RAG (retrieval-augmented generation) system.

Rule Generator Agent: takes the input from the data analyst agent output and generates the rules. It asks the user for confirmation using the HITL functionality. The user has three options for the proposed changes/additions: "accept," "reject," and "modify."

Conclusion

The automation improves the efficiency of the process and ensures transparency with HITL.

Coherence: Gamified Human-AI Collaboration for Executive Digital Transformation Training and Decision-Support



Scott J. Warren
Department of Learning
Technologies
University of North Texas
scott.warren@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

As AI adoption accelerates, organizations need workforce training and teaching approaches that prepare leaders to collaborate effectively with AI while making high-stakes decisions before, during, and after digital transformation. This 20-minute talk shares Coherence, a functioning (and still-in-development) gamified web application designed to train leaders to make more coherent decisions that support AI, automation, and other IT adoption at scale. The tool is meant not to treat “digital transformation” but as a set of learned principles and clear practices for implementation to support realistic adoption at scale in organizations.

Coherence is built as a structured learning and decision-support environment. In the app, users move through scenario-based “decision missions” that simulate the real constraints leaders face (conflicting stakeholder goals, capability gaps, governance issues, workforce impacts, and uncertainty about emerging technologies).

The AI layer is deliberately not a generic chatbot. Instead, it is used as a guided collaborator that helps users

1. clarify transformation intent,
2. surface assumptions and risks,
3. compare strategic options,

4. anticipate second-order effects on people and processes, and
5. produce actionable outputs that can be reviewed, revised, and used for instruction or organizational discussion.

The platform is intended to support workforce training, higher education teaching, K-12 technology adoption, and ethical thinking about whether we should integrate these tools using business, practice, and feasibility principles. In instructional contexts, Coherence supports graduate-level learning in digital transformation, technology leadership, and decision-making by turning abstract frameworks into interactive practice. In professional contexts, it serves as a scalable training tool for leadership development. The goal is not just “knowing about AI,” but making better decisions about whether to adopt AI or other advanced technologies, if the organization is ready, and how to think about proper integration to ensure measurable return-on-investment with minimal unintended consequences.

The talk will highlight

6. the instructional design logic behind the gamified structure,

AI in Action 2026 sponsored by

The Departments of Technical Communication and Data Analytics & Statistics at the University of North Texas

7. how the platform operationalizes decision quality (i.e., coherence, defensibility, and alignment across stakeholders and capabilities), and
8. how Coherense supports multiple transformation lenses including AI/automation, data/analytics, and emerging IT domains such as quantum as a strategic horizon topic.

Attendees will see how the tool generates structured outputs for reflection, coaching, and organizational use rather than producing fragile one-off chat responses.

Finally, I will describe an ongoing Delphi study with corporate executives examining their learning experience and usability/UX perceptions while using Coherense. Special attention is given to how leaders and operations actors in organizations should reason about AI, automation, quantum, and other emerging IT under guided human-AI collaboration. The session closes with practical takeaways on designing AI-enhanced learning systems that improve real decision-making, not just engagement.

Can College Instructors Tell AI from Human Writing? Evidence from 400 Classification Trials of STEM Reports



Mason Pellegrini
Technical Communication
University of North Texas
mason.pellegrini@unt.edu

David Rowe
Medical Decision Modeling
david@rowe-stats.com

Paul Hunter
Department of English and
Modern Languages
Longwood University
hunterpt@longwood.edu

Adrianna Deptula
Department of English
Case Western Reserve
University
axd1309@case.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

As generative AI becomes more common in higher education, instructors are increasingly expected to determine whether student writing is human-authored or AI-generated. This expectation has major implications for assessment and academic integrity processes, yet there is still limited empirical evidence about how accurately instructors can make these judgments. Many studies investigate whether people can distinguish AI and human text, but nearly all of these studies randomly recruit participants from Prolific and do not attempt to recreate the educational context (Chein et al., 2024; Gunser et al., 2022; Köbis & Mossink; 2021 Porter & Machery, 2024). This presentation reports findings from a study designed to address that gap.

The study examined whether experienced college writing instructors could distinguish between student-written and AI-generated excerpts from STEM reports without any contextual information about the writer, class, or assignment. Forty instructors completed a total of 400 classification

trials. In each trial, participants read a 300–400 word excerpt and judged whether it had been written by a student or generated by AI, while also indicating their level of confidence. Student-written samples were drawn from the Michigan Corpus of Upper-Level Student Papers, and AI samples were generated to mirror those reports in discipline, title, and major features.

Overall, instructors classified texts correctly 71% of the time. They were more successful at identifying AI-generated text (79%) than student-written text (66%), and confidence was positively related to accuracy. However, no significant relationships were found between accuracy and demographic variables such as years of teaching experience, institution type, or self-reported familiarity with AI. Most importantly, the findings reveal a persistent false-positive problem: instructors sometimes labeled authentic student writing as AI-generated, even when they indicated a high level of certainty in their judgments.

These results matter beyond writing studies. Across higher education, many faculty are now

being asked to monitor AI use in student work, often without reliable tools or clear evidence about what human judgment can and cannot do. This presentation argues that while instructors may detect some patterns, policing AI use through text-based suspicion remains a fraught and limited practice. This is especially true considering the harsh penalties for students who may be wrongly accused of unauthorized AI use. Therefore, we urge colleges and universities to place less emphasis on catching AI use through subjective detection and more emphasis on AI literacy, transparent expectations, and assessment design that encourages accountable, context-aware writing. In this way, the project speaks directly to ongoing conversations about ethical and effective human-AI collaboration in teaching and learning.

References

- Chein, J. M., Martinez, S. A., & Barone, A. R. (2024). Human intelligence can safeguard against artificial intelligence: Individual differences in the discernment of human from AI texts. *Scientific Reports*, 14(1), 25989. <https://doi.org/10.1038/s41598-024-76218-y>.
- Gunser, V. E., Gottschling, S., Brucker, B., Richter, S., Çakir, D. C., & Gerjets, P. (2022). The pure poet: How good is the subjective credibility and stylistic quality of literary short texts written with an artificial intelligence tool as compared to texts written by human authors? *Proceedings of the First Workshop on Intelligent and Interactive Writing Assistants (In2Writing 2022)*, 60-61. <https://doi.org/10.18653/v1/2022.in2writing-1.8>.
- Köbis, N., & Mossink, L. D. (2021). Artificial intelligence versus Maya Angelou: Experimental evidence that people cannot differentiate AI-generated from human-written poetry. *Computers in Human Behavior*, 114, 106553. <https://doi.org/10.1016/j.chb.2020.106553>.
- Porter, B., & Machery, E. (2024). AI-generated poetry is indistinguishable from human-written poetry and is rated more favorably. *Scientific Reports*, 14(1), 26133. <https://doi.org/10.1038/s41598-024-76900-1>.

Answering IBM's Call for a Systems Theory: Agentic AI as Open Systems



Stephen Penn
Department of Data Analytics
and Statistics
University of North Texas
stephen.penn@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

IBM research recently argued that “the current development of Agentic AI requires a more holistic, systems theoretic perspective in order to fully understand their capabilities and mitigate any emergent risks” (Miehling et al., 2025). Current AI development has shown the need for prompt engineering, context engineering, and intent engineering as solutions for better control over agentic AI. This approach lacks a holistic view that misses agentic AI’s true capabilities and risks. This presentation proposes that Katz and Kahn's open systems theory provides the required broad view requested by IBM. This foundational view offers a framework to design and improve agentic AI systems through systematic feedback.

Agentic AI systems can be conceptualized as open systems whose adaptive capacity depends on how effectively their feedback mechanisms are designed and used. This presentation maps the components of agentic AI to the ten characteristics in Katz and Kahn’s open systems, which are energetic inputs, throughput, energetic outputs, cycle of events, negative entropy, feedback mechanisms, dynamic homeostasis, differentiation, integration and coordination, and equifinality. In this framework, the system’s feedback processes are organized into three layers that align with prompt, context, and

intent engineering and correspond to triple-loop learning.

The organizational learning loop, as defined by the concept of triple loop learning, fits the feedback loop existing in current agentic AI architecture. Single-loop learning operates at the prompt layer, fixing task-level errors within existing processes. Double-loop learning operates at the context layer, redesigning workflows, tools, and orchestration when processes prove inadequate. Triple-loop learning operates at the Intent layer, revising organizational goals, constraints, and success criteria when strategic priorities shift. Together, these feedback mechanisms enable negative entropy, where the system accumulates organizational knowledge over time while maintaining dynamic homeostasis, preserving stable orchestration and roles as the system adapts.

The core proposition is that increasing the quality and explicitness of the design of the feedback mechanisms across all three layers and revising them based on structured feedback from event logs and process analysis, will lead to alignment of agentic outputs with organizational expectations. This makes agentic AI systems self-sustaining open systems rather than static pipelines.

The presentation includes a detailed mapping table showing how each of Katz and Kahn's ten characteristics corresponds to specific agentic AI components, a conceptual architecture diagram illustrating the three feedback loops, and boundary conditions specifying when this theoretical framework applies.

References

Miehling, E., Ramamurthy, K. N., Varshney, K. R., Riemer, M., Bouneffouf, D., Richards, J. T., ... & Geyer, W. (2025). Agentic AI needs a systems theory. arXiv preprint arXiv:2503.00237.

Comparing LLM and human financial decision-making



Orhan Erdem
Department of Data Analytics &
Statistics
University of North Texas
orhan.erdem@unt.edu

Ragavi Pobbathi Ashok
Department of Data Analytics &
Statistics
University of North Texas
ragavipobbathiashok@my.unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

In this paper, we investigate how large language models (LLMs) make financial decisions by comparing their responses systematically with those of human participants from around the world. We administered a set of widely used financial decision-making questions to several prominent LLMs, including GPT-4, GPT-4o, GPT-5, Gemini 2.0 Flash, and DeepSeek R1. Each model was evaluated at multiple temperature settings, resulting in 21 model-temperature combinations overall. We then benchmarked these outputs against human responses from a dataset spanning 53 countries.

Our analysis yields three main findings. First, in cross-country comparisons, LLM responses tend to cluster closely together and form a distinct group that is separate from all national samples. In this sense, they do not appear to be WEIRD (Western, Educated, Industrialized, Rich, and Democratic), contrary to suggestions made in earlier studies in other settings. Second, LLMs generally display a risk-neutral decision profile, tending to favor options consistent with expected-value reasoning

in lottery-style tasks. Third, in questions involving trade-offs between present and future rewards, LLMs often produce responses that are internally consistent and economically coherent. Together, these findings deepen our understanding of how LLMs mirror human financial decision-making and point to the possible influence of cultural patterns and training data on their outputs.

Detecting and Quantifying Hallucinations in LLM-Generated References



Ankur Bhattacharjee
Department of Data Analytics &
Statistics
University of North Texas
ankurbhattacharjee@my.unt.edu

Orhan Erdem
Department of Data Analytics
& Statistics
University of North Texas
orhan.erdem@unt.edu

Copyright © 2026 by the author(s). Permission granted to AI in Action to publish.

Abstract

Large language models (LLMs), a type of AI, are developed by training on large amounts of text data so they can recognize patterns and generate meaningful responses (Hugging Face, n.d.). Since the launch of OpenAI's Chat Generative Pre-trained Transformer (ChatGPT) in late 2022, interest in artificial intelligence has surged worldwide. Economists predict that AI will help people think more critically and increase their productivity (Buchanan et al., 2024). LLMs such as ChatGPT can assist academic researchers in a variety of ways. These tools can make it easier to search the literature, summarize research findings, and draft systematic reviews (Chelli et al., 2024). This advancement of computer science brings new problems, such as AI hallucinations. This happens when LLMs provide us with answers that sound real, but have no factual basis (Erdem et al., 2025).

Hence, there are challenges in generating accurate and reliable citations, particularly with the increasing use of LLMs. Key concerns include citation bias, invalid digital object identifiers (DOIs) and hallucinated references (Kim et al., 2025). ChatGPT's tendency to generate fabricated references remains unclear, highlighting the need for systematic investigation and mitigation

(Frosolini et al., 2023). Erdem et al. (2025) found in that chatbots vary widely in citation accuracy for financial literature, with ChatGPT-4o and o1-preview showing hallucination rates of about 20%, while Gemini Advanced had a much higher rate of 76.7%, and errors tended to increase for newer topics. An evaluation of 15 articles generated using ChatGPT 4 and ChatGPT o1 found that only 60.4% of references were accurate and 22.2% of citations were highly relevant, with no significant improvement across different training conditions (Cheng et al., 2025). However, a gap has been observed in previous studies regarding whether incorporating additional large language models into the evaluation process can reduce hallucination rates in LLM-generated financial references.

The purpose of this study is to determine whether incorporating a multi-LLM evaluation framework improves the accuracy and reliability of LLM-generated citations. In this research, 150 article references will be generated across 15 financial-literature topics using a GPT-5.2 prompt. Each reference will undergo manual verification using authoritative indexing sources, specifically Google Scholar and Scopus, to establish a baseline accuracy metric. The same set of references will then be evaluated by four different LLM models: Gemini 3 Flash, Claude Sonnet 4.6, DeepSeek-V3.2, and GPT

5.2 to examine whether the hallucination rate decreases under this verification-based procedure. The responses from all four LLMs will be collected and statistically analyzed to compare hallucination rates. This study aims to improve trust and accountability in AI-generated references and will guide researchers in ensuring a more robust reference validation mechanism.

References

- Buchanan, J., Hill, S., & Shapoval, O. (2024). ChatGPT Hallucinates Non-existent Citations: Evidence from Economics. *The American Economist*, 69(1), 80–87. <https://doi.org/10.1177/05694345231218454>
- Chelli, M., Descamps, J., Lavoué, V., Trojani, C., Azar, M., Deckert, M., Raynier, J.-L., Clowez, G., Boileau, P., & Ruetsch-Chelli, C. (2024). Hallucination Rates and Reference Accuracy of ChatGPT and Bard for Systematic Reviews: Comparative Analysis. *Journal of Medical Internet Research*, 26, e53164. <https://doi.org/10.2196/53164>.
- Cheng, A., Nagesh, V., Eller, S., Grant, V., & Lin, Y. (2025). Exploring AI Hallucinations of ChatGPT. *Simulation in Healthcare: The Journal of the Society for Simulation in Healthcare*, 20(6), 413–418. <https://doi.org/10.1097/SIH.0000000000000877>.
- Erdem, O., Hassett, K., & Egriboyun, F. (2025). Hallucination in AI-generated financial literature reviews: evaluating bibliographic accuracy. *International Journal of Data Science and Analytics*, 20(5), 4501–4510. <https://doi.org/10.1007/s41060-025-00731-0>.
- Frosolini, A., Gennaro, P., Cascino, F., & Gabriele, G. (2023). In Reference to “Role of Chat GPT in Public Health”, to Highlight the AI’s Incorrect Reference Generation. *Annals of Biomedical Engineering*, 51(10), 2120–2122. <https://doi.org/10.1007/s10439-023-03248-4>.
- Kim, E., Kipchumba, F., & Min, S. (2025). Geographic Variation in LLM DOI Fabrication: Cross-Country Analysis of Citation Accuracy Across Four Large Language Models. *Publications*, 13(4), 49. <https://doi.org/10.3390/publications13040049>.
- Hugging Face. (n.d.). Natural language processing and large language models. Retrieved March 15, 2026, from <https://huggingface.co/learn/llm-course/chapter1/2>.

Author Index

Ashok, Ragavi Pobbathi.....	24
Balmaki, Behnaz.....	11
Bhattacharjee Ankur.....	25
Deptula, Adrianna.....	20
Erdem, Orhan.....	24, 25
Friess, Erin.....	7
Gu, Kai.....	9
Gurunathan, Deepeka.....	15
Hunter, Paul.....	20
Nguyen, Thuan Luong.....	12
Orhan, Mehmet.....	13
Orhan, Zeynep.....	13
Pellegrini, Mason.....	20
Penn, Stephen.....	22
Rea, Ashley.....	2
Rostami, Masoud.....	11
Rowe, David.....	20
Sabnam, Sonali.....	17
Shi, Weishi.....	9
Stevens, Aleksander.....	9
Suribhotla, Shreyes.....	9
Wang, Bryan.....	9
Warren, Scott J.....	18
Williams, Eric.....	4
Zhang, Jiaxin.....	5