

LTEC Intelligence Brief

Learning, AI, Cybersecurity, and Digital Innovation

Friday, September 25, 2026

Published by LTEC with Lance

AI Agents Are Moving From Tools Into Markets, Security Operations, and Infrastructure

Excerpt: The latest 24 to 48 hours show agentic AI moving into domains where decisions have real economic, security, and governance consequences. Anthropic tested agents negotiating on behalf of people in a controlled market and separately proposed measurable public indicators for AI-led research and agent oversight. GitHub demonstrated an autonomous fuzzing pipeline while explicitly warning that agent-selected commands should run only inside disposable environments. GitHub also introduced enterprise controls for how future Copilot features are enabled. CISA is pushing the CVE program toward a new quality framework as vulnerability discovery scales, while actively exploited WSO2 and Adobe flaws add immediate remediation pressure. In education, EDUCAUSE concluded a two-week Teaching with AI program centered on practical integration, engagement, personalization, and academic integrity.

Executive Brief

Anthropic published Project Swap on September 24, a controlled experiment in which 201 employees used Claude-powered agents to trade books on their behalf. After a short preference conversation, agent rankings matched participant rankings on 61 percent of book pairs. Anthropic reports that the agents generally negotiated effectively, while the larger limitation was incomplete information about the people they represented.

Project Swap matters beyond books. Anthropic identifies unresolved questions about whether an agent truly understands a person's preferences, what rules should govern agent marketplaces, what happens when deals fail, and how much agent activity should be visible to participants. Those questions map directly to procurement, scheduling, commerce, hiring, and other delegated decision workflows.

Anthropic also published a measurement framework for frontier AI development. The company proposes tracking the share of AI research and development performed by AI, coverage and latency of agent oversight, escalation rates, and the share of research compute allocated to safety. Anthropic says it plans to give independent third-party evaluators access to internal processes, systems, and data comparable to its internal risk teams.

GitHub Security Lab published an AI-powered fuzzing workflow on September 24. The autonomous pipeline can identify targets, analyze builds, write fuzzing harnesses, run AFL++, inspect coverage, improve harnesses, triage crashes, and draft vulnerability reports. GitHub explicitly warns that the current taskflow can run agent-selected build commands on the host and should therefore be used only in a disposable environment without elevated privileges.

GitHub also introduced a new global default policy for generally available Copilot features in Business and Enterprise. Administrators can choose whether eligible current and future features are enabled, disabled pending approval, or delegated to organization administrators. The policy takes effect October 22, while explicit administrator decisions remain preserved.

CISA's newly reported CVE Quality Era framework focuses on four dimensions: program governance, ecosystem participation, data infrastructure, and CVE record content. The initiative responds to a vulnerability ecosystem in which scale, automation, and AI-assisted discovery increase the need for timely, consistent, actionable records.

CISA's Known Exploited Vulnerabilities catalog also added actively exploited flaws affecting WSO2 products and Adobe Commerce and Magento on September 24. Current reporting identifies CVE-2026-5430 and CVE-2026-71362, with a September 27 federal remediation deadline.

EDUCAUSE's Teaching with AI program concluded September 24 after live sessions across two weeks. The program is designed for higher education faculty, instructional designers, and support staff and focuses on practical AI integration, student engagement, personalization, educational outcomes, and academic integrity.

Top Developments

AI agents are beginning to negotiate and transact on behalf of people, creating new requirements for preference fidelity, market rules, transparency, and recourse.

Frontier AI oversight is becoming measurable through monitoring coverage, review latency, escalation rates, and third-party verification.

Autonomous security agents can expand defensive capacity, but their execution environment must be treated as potentially hostile.

Enterprise AI governance is moving into product controls that determine whether new agent features are enabled by default.

Vulnerability management infrastructure is being redesigned for higher volume, better data quality, and increasing automation.

Higher education AI development continues to shift from awareness toward practical course integration, assessment, and academic-integrity decisions.

Artificial Intelligence and Emerging Technology

Project Swap provides an unusually concrete look at delegated AI decision-making. Anthropic reports that stronger models produced more efficient markets and that model choice affected negotiating outcomes more than changes to agent instructions in its repeated simulations.

The central limitation was representation. An agent can negotiate competently and still optimize for an incomplete or incorrect understanding of the person it represents. That creates a governance problem distinct from model capability.

Organizations considering agents for procurement, scheduling, purchasing, or negotiation should therefore test preference fidelity separately from task performance. The relevant question is not only whether the agent can make a deal, but whether it makes the deal the principal would actually want.

Anthropic's frontier-lab measurement proposal adds another layer. The company argues that the public should be able to track how much AI is involved in building future AI systems, how agent actions are monitored, and how development resources are allocated.

The broader pattern is a move from qualitative claims such as 'human oversight exists' toward operational measures that can be audited, compared over time, and eventually verified by independent parties.

Learning and Digital Education

EDUCAUSE's Teaching with AI program concluded September 24. Its stated learning goals include understanding AI's implications for higher education, particularly academic integrity, and using practical strategies to integrate AI into curricula.

The program's format is significant for faculty development: a two-week sequence, multiple live sessions, self-paced modules, interactive activities, and a year of access. This blends short synchronous engagement with sustained asynchronous learning rather than relying on a one-time workshop.

For instructional designers, the current direction is clear. AI professional development should move from tool demonstrations into course-level decisions about learning outcomes, student agency, academic integrity, accessibility, personalization, and evidence of learning.

The Project Swap findings also have instructional implications. If an agent can act on a learner's behalf, educators need to distinguish assistance from representation. A system that drafts, searches, negotiates, or completes multistep tasks may be doing more than helping a learner think.

Assessment policies should therefore define not only whether AI is allowed, but what forms of delegation are permitted, what evidence the learner must personally produce, and how instructors can verify that the submitted work represents the learner's knowledge and judgment.

Cybersecurity and Digital Trust

GitHub Security Lab's autonomous fuzzing taskflow demonstrates both the promise and risk of agentic security automation. The system can perform a substantial sequence of vulnerability-research tasks without continuous human direction.

GitHub's own warning is important: the current workflow can execute arbitrary build commands chosen by the LLM on the host. A prompt-injected agent could therefore act with the permissions of the user. GitHub recommends a disposable Codespace or throwaway virtual machine without elevated privileges.

That warning captures a durable security rule for agents. Model safeguards are not a substitute for containment.

High-consequence permissions should be restricted through operating-system, container, network, credential, and identity controls outside the model.

CISA's CVE Quality Era framework addresses the information layer of cyber defense. As vulnerability discovery accelerates, defenders need records that are not merely numerous but complete, timely, consistent, and usable by automated risk-management systems.

The immediate operational priority remains exploitation. CISA added WSO2 and Adobe Commerce and Magento vulnerabilities to the KEV catalog on September 24. Organizations using affected products should verify exposure, follow vendor remediation guidance, and investigate for evidence of compromise rather than treating the task as patching alone.

AI Governance and Accessibility

Anthropic's Project Swap makes delegated authority a governance issue. If an agent represents a person or organization, the system needs rules for authorization, preference capture, disclosure, transaction limits, dispute handling, and visibility into what the agent did.

Anthropic's separate oversight measurements suggest a useful governance vocabulary: coverage, review latency, and escalation rate. These measures can turn broad supervision commitments into evidence that can be monitored and audited.

GitHub's new Copilot default-feature policy illustrates governance at the enterprise product layer. Administrators can decide whether future generally available capabilities are automatically enabled, automatically withheld pending approval, or left to organization administrators.

That type of control is increasingly necessary because AI platforms evolve continuously. Governance cannot depend on a one-time product approval if new models, tools, agents, and permissions can materially change the system after deployment.

No material new W3C accessibility standard was identified in the immediate research window. WCAG 2.2 remains the current WCAG 2 Recommendation. Agentic systems should also be tested for accessible status updates, keyboard control, understandable confirmations, error recovery, reversibility, and equivalent access to information about automated actions.

Workforce Development

Agentic work changes the skill requirement from operating a tool to supervising delegated activity. Workers need to define goals, communicate preferences, set constraints, inspect evidence, recognize exceptions, and intervene when the agent's method or outcome diverges from intent.

Project Swap illustrates a specific workforce risk: a capable agent may negotiate well while representing the user imperfectly. Employees who delegate work to agents need to know how to specify preferences and verify whether the system's choices still align with business priorities.

GitHub's fuzzing workflow shows the complementary technical skill. Security professionals need to understand not only what an agent can automate, but where it executes, what permissions it holds, what tools it can invoke, and what evidence it produces.

Faculty and instructional-design development is moving in the same direction. EDUCAUSE's current programming emphasizes practical application and academic-integrity judgment rather than generic familiarity with AI.

High-value workforce development should therefore assess supervised performance in realistic workflows. The worker should demonstrate that they can delegate appropriately, verify results, protect data, manage permissions, and exercise stop authority.

Practical Takeaways

- Test whether an agent accurately represents user or organizational preferences before allowing it to negotiate, purchase, schedule, or transact.
- Measure AI oversight with concrete indicators such as monitoring coverage, review latency, escalation rate, blocked actions, and human intervention.
- Run autonomous security and coding agents in isolated environments with minimal privileges and separate credentials.
- Establish an enterprise default for newly released AI features instead of allowing governance to depend on product defaults.
- Treat CISA KEV additions as immediate investigation and remediation priorities.
- Design faculty development as sustained practice with real course artifacts, not only one-time AI demonstrations.
- Define acceptable AI delegation in assessment policies, including what the learner must personally know, decide, explain, or demonstrate.

- Continue using WCAG 2.2 as the accessibility baseline and explicitly test dynamic agent status, confirmation, reversal, and error states.

Watchlist

- Further evidence about how well agents represent human preferences in higher-stakes markets such as purchasing, employment, healthcare, and procurement.
- Whether frontier AI developers adopt comparable public metrics for AI-led research, oversight coverage, review latency, and safety-resource allocation.
- Independent verification of frontier-lab oversight measurements and the design of third-party access to internal AI-development systems.
- Adoption of autonomous vulnerability-research agents and the security controls used to contain prompt injection and unsafe execution.
- Enterprise responses before GitHub's October 22 Copilot default-feature policy takes effect.
- CISA implementation of the CVE Quality Era framework and measurable improvements in CVE completeness, timeliness, infrastructure reliability, and global participation.
- Exploitation activity involving the newly added WSO2 and Adobe Commerce and Magento KEV entries.
- Evidence that higher education AI programs improve assessment quality, faculty confidence, accessibility, and student learning rather than only tool adoption.

Final Thought

The next phase of AI is delegation.

Agents are beginning to negotiate, research, test software, and make multistep decisions on behalf of people. That makes capability only one part of the evaluation.

The more important questions are whether the agent understands the person it represents, whether its authority is bounded, whether its actions are observable, whether its environment is secure, and whether a human can intervene before a mistake becomes consequential.

Complete Original-Source Links

[Anthropic - Project Swap: What happens when agents trade for us? - September 24, 2026](#)

[Anthropic - Measurements for understanding the pace of AI development inside frontier labs - September 2026](#)

[GitHub Security Lab - AI-powered fuzzing with the GitHub Security Lab Taskflow Agent - September 24, 2026](#)

[GitHub - Default Enablement of Copilot Features for Copilot Business and Enterprise - September 24, 2026](#)

[CISA - Known Exploited Vulnerabilities Catalog](#)

[CISA - BOD 26-04: Prioritizing Security Updates Based on Risk](#)

[EDUCAUSE - Teaching with AI - September 15-24, 2026](#)

[Google - Behind Project Suncatcher, our moonshot to put AI in space - September 24, 2026](#)

[Google - 5 ways to upgrade your study habits with Chrome - September 24, 2026](#)

[W3C Web Accessibility Initiative - WCAG 2 Overview](#)

Source and Analysis Note

This edition prioritizes primary and authoritative sources published, updated, or operationally relevant during the previous 24 to 48 hours. Anthropic's Project Swap results and frontier-lab measurements are company-reported research and are identified as such. GitHub's agentic fuzzing description and security warning come directly from GitHub Security Lab. CISA's KEV catalog is treated as the authoritative exploitation-priority source; contemporaneous reporting was used to confirm the September 24 additions where the searchable CISA page did not expose a dedicated item in the research interface. Educational event pages document current professional-development activity and are not presented as evidence of measured learning outcomes. Analysis, implications, and recommendations are LTEC with Lance interpretations based on the cited material.

LTEC with Lance

LTEC with Lance publishes practical analysis at the intersection of learning technology, artificial intelligence, cybersecurity, digital accessibility, and workforce development. [Visit LTECwithLance.com](https://ltecwithlance.com) for additional analysis, resources, toolkits, apps, and practical guidance.

LTEC with Lance | Learning, Technology, and Thoughtful Change.