

FT: Agentic AI

2026 FT: Agentic AI

Card Type	Future Technology Possibility
Series	Immersive Futures Guild — Vision 2035
Layer	1 — Atomic Foresight Object
Status	Active
Confidence	Medium
Workshop	Circle of Scholars — January 2026
Facilitator	Circle of Scholars Workshop Team
Tags	agentic-AI autonomy pedagogy layer1 ft
Tally.so Form	https://tally.so/r/ilrn-if-ft-agentai-2026

Agentic AI systems capable of planning multi-step actions, pursuing goals across time, and operating with minimal human oversight are moving from laboratory research into deployed applications. In immersive learning, agentic AI raises questions about who controls the trajectory of a learning experience, how agency is shared between learner, educator, and system, and what happens when learning agents pursue optimization targets that do not align with human educational values.

Key Drivers / Contributing Conditions:

- AI capability scaling enabling multi-step goal pursuit
- Commercial deployment of agentic systems in consumer and enterprise contexts
- Research on AI-assisted learning path generation and adaptive scaffolding

Tensions Carried Forward to Part II:

- Who is accountable when an agentic AI system makes a pedagogically harmful decision?
- Can learner agency be preserved in an environment where AI agents are continuously optimizing?

Ways of Knowing: Tree · Garden · Lantern

PART II — COMMUNITY EVIDENCE & DIALOGUE TRACK | FT: Agentic AI | H2 2026 — Living

T	COMMUNITY CONTRIBUTION FORM — FT: Agentic AI Submit case examples, methodological challenges, cultural perspectives, and proposed evidence criteria via: https://tally.so/r/ilrn-if-ft-agentai-2026
---	--

Part II — Scope and Instructions
This section collects community responses, case examples, and challenges to the Part I foresight snapshot above.
It opens July 1, 2026 and undergoes synthesis review in September 2026, November 2026, and January 2027.
Contributions are submitted via the Tally.so form above and appear in the registers below after editorial review.
The Part I text is not modified in response to Part II contributions; it is versioned at the Annual Handoff review.
Contribution categories: Case Example Methodological Challenge Cultural/Community Perspective Proposed Evidence Criterion
Ways of Knowing accepted: Tree (evidence) Garden (practice) Lantern (futures)

Tensions Open for Community Response:

- Who is accountable when an agentic AI system makes a pedagogically harmful decision?
- Can learner agency be preserved in an environment where AI agents are continuously optimizing?

Contributor / Date	Category	Way of Knowing	Contribution Summary
[Awaiting contributions — form opens July 1, 2026]			

Revision #1

Created 25 May 2026 20:23:54 by Jonathon Richter

Updated 25 May 2026 20:24:18 by Jonathon Richter