

Generative AI Can Harm Learning

- GenAI tools can undermine learning, not just support it
- Designed for educators, instructional designers, and school leaders
- AI is not universally harmful; specific uses carry measurable risks
- Focus: how and when GenAI actually harms learning



The Central Tension



- Better output in less time doesn't equal better understanding.



- AI raises homework scores while lowering unassisted exam scores.



- Productivity $\neq$ learning: completion is not the end goal in education.



- Why risk exists, what the evidence shows, how to design safer learning.



- The full penalty can take months or two years to emerge.



Cognitive Offloading and Desirable Difficulties



- AI offloads core cognitive work, weakening encoding and retention



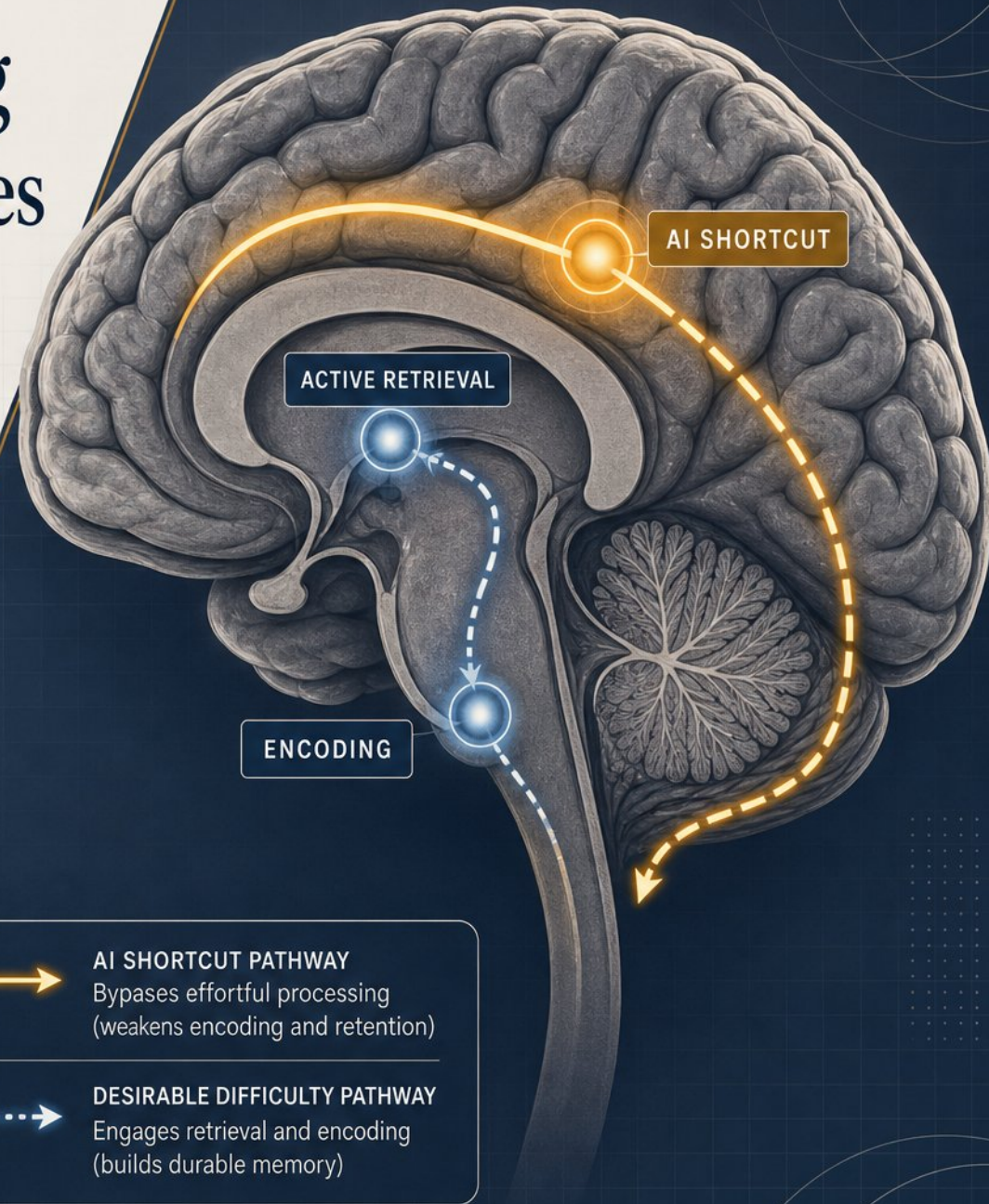
- Struggle and retrieval are features, not bugs: desirable difficulties build durable memory



- Dependent offloading delegates core thinking; autonomous offloading scaffolds it



- Both feel equally productive in the moment, but outcomes diverge sharply



How Learners Actually Use AI



- High-risk patterns: copy-paste, prompt-only writing, unverified answers, fast high-scoring homework



- Automation bias fuels overtrust in AI output, weakening critical thinking



- Same tool, two outcomes: outsourced work vs. scaffolded learning



- One in five AI users spend comparable time to peers—and keep learning

SPECIMEN: DELEGATE

The Delegate

OUTSOURCES THE WORK



- ✗ Copy-paste
- ✗ Prompt-only writing
- ✗ Unverified answers
- ✗ Fast high-scoring homework



High risk. Weak learning.

SPECIMEN: SCAFFOLDER

The Scaffolder

USES AI TO LEARN BETTER



- ✓ Clarify and ask questions
- ✓ Seek feedback and critique
- ✓ Verify and reflect
- ✓ Build understanding step by step



Lower risk. Stronger learning.

Key Research Findings

- - GPT Base: unassisted exams down 17% (Bastani et al., 2025)
- - Homework scores up 18%; monthly exams down 20% (Strömberg et al., 2026)
- - Retention odds down 25% on proctored items
- - Positive outcomes possible with structured design (40.4% of studies)
- - Outcome depends on design, not the tool itself



-17%

Unassisted
Exam Score



+18%

Homework
Score



-25%

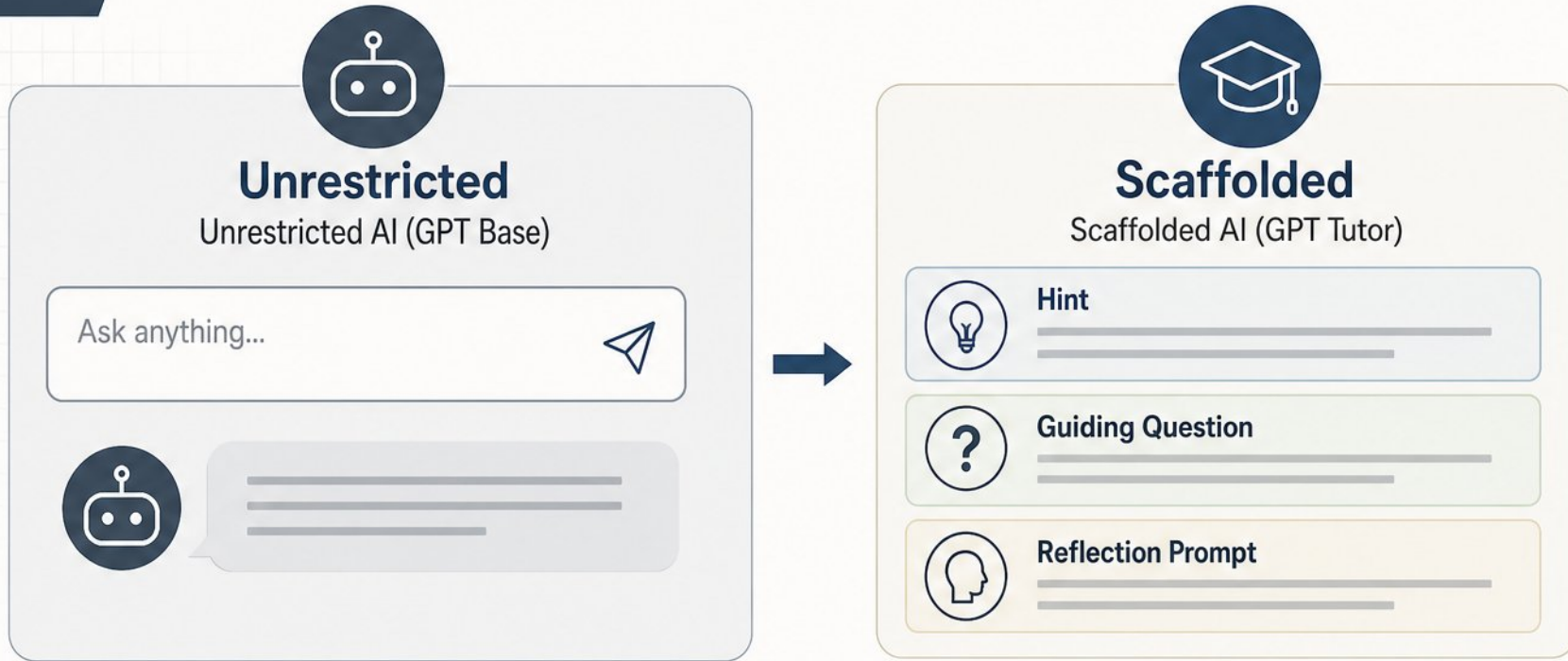
Proctored
Retention Odds



Evidence suggests mixed outcomes.
Design and context matter.



The Role of Guardrails and Scaffolding



Scaffolding Structure



- Teacher-designed hints in Bastani et al. (2025) largely mitigated learning harm
- Unrestricted AI (GPT Base): copying solutions; unassisted grades fell 17%
- Scaffolded AI (GPT Tutor): offered hints and support; harm essentially eliminated
- Structured pedagogical integration turns potential harm into neutral or positive outcomes

Where the Harm Is Greatest

- ▶ Social sciences hardest hit: exam scores fall 27%
- ▶ Junior, high-achieving, and male students face larger penalties
- ▶ Foundational skills and problem-solving practice most vulnerable
- ▶ Highest risk: K-12 and early undergraduate foundational stages
- ▶ Workplace training risks emerge when AI replaces core competency practice



Designing Safer Learning Experiences



- Use process-based tasks: drafts, staged submissions, oral defenses



- Make thinking visible: annotations, reflections, and reasoning



- Position AI as feedback support, not answer generator



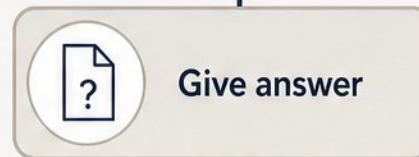
- Anchor tasks in local contexts to reduce AI's generic edge



- Add verbal debriefs to replace AI-completable written reports



Task Design Choices



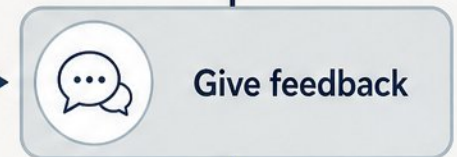
Give answer



One step



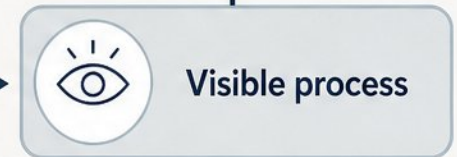
Hidden work



Give feedback



Staged submission



Visible process



SHORTCUT

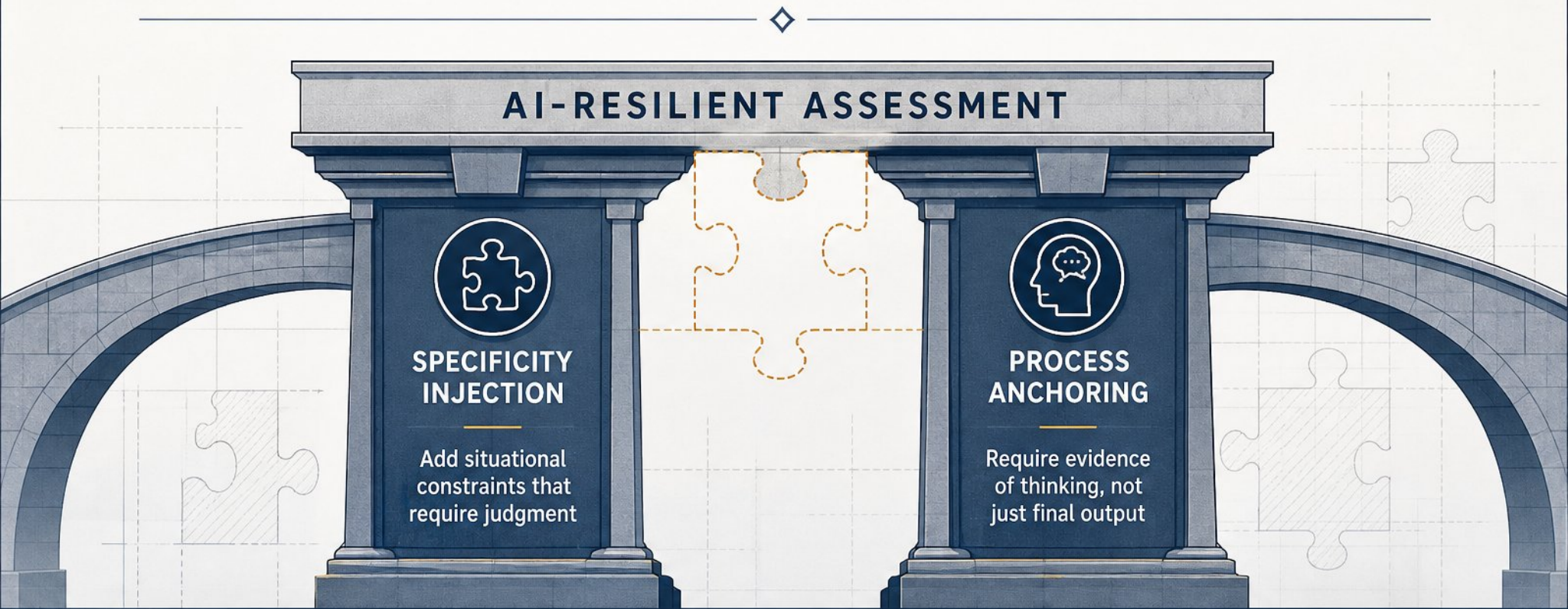
Easier for AI.
Weaker learning.



PATH

Harder for AI.
Stronger learning.

AI-Resilient Assessment Strategies



- **Specificity injection:** add situational constraints that require judgment



- **Process anchoring:** require evidence of thinking, not just final output



- **Examples:** disagreement analysis, staged submissions, verbal debriefs



- **Most fixes change what is asked**—not the whole assessment

Designing for Visible Thinking



- Require drafts, reflections, and process documentation



- Anchor tasks to class discussions and local contexts



- Ask learners to explain their choices and reasoning



- AI can't fake a messy, authentic thinking process



Policy and Guardrails for Organizations



- Define allowed uses, disclosure rules, and prohibitions



- Align policies with culture and incentives



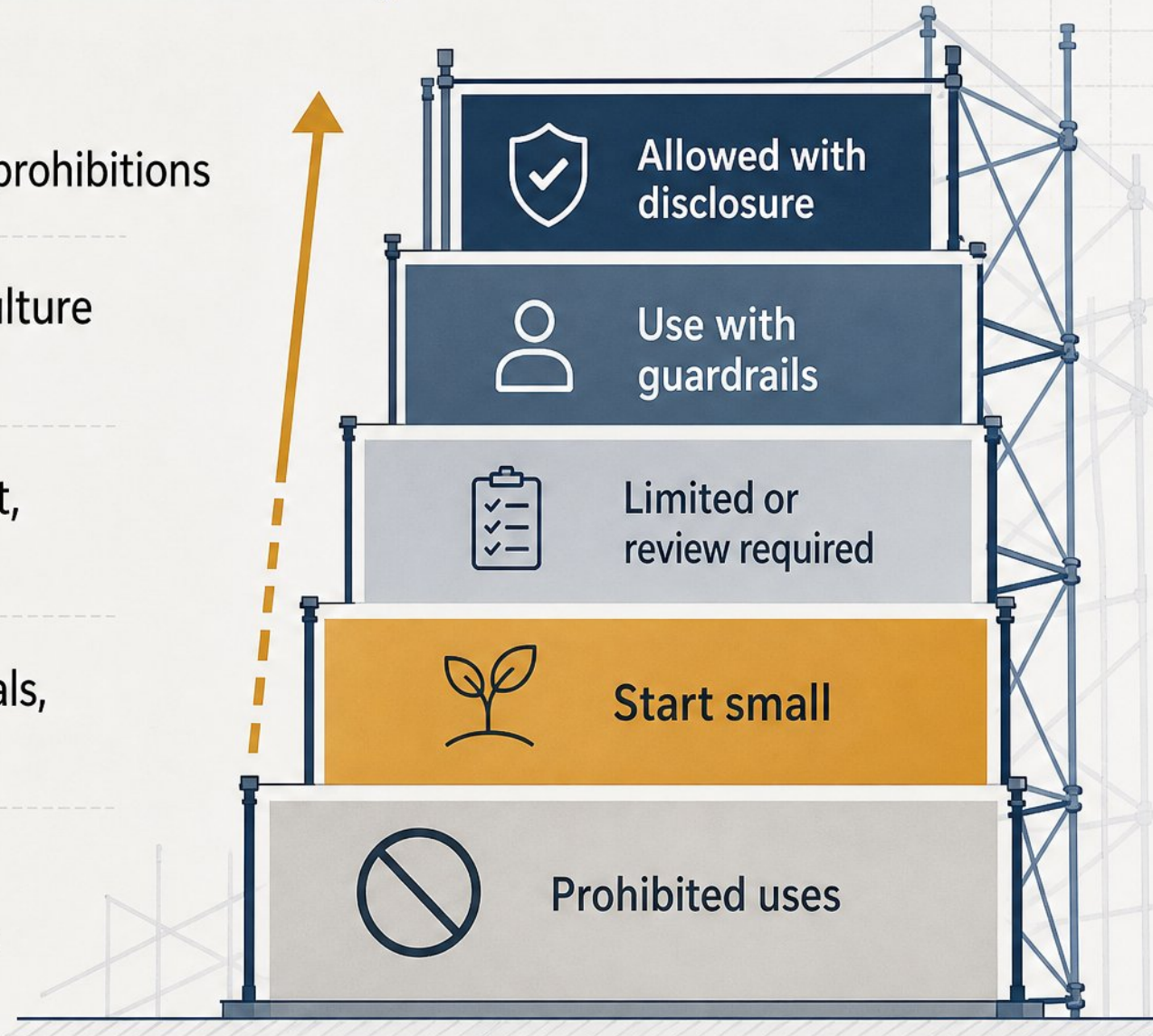
- Phase adoption: pilot, evaluate, then scale



- Monitor learning signals, not just output

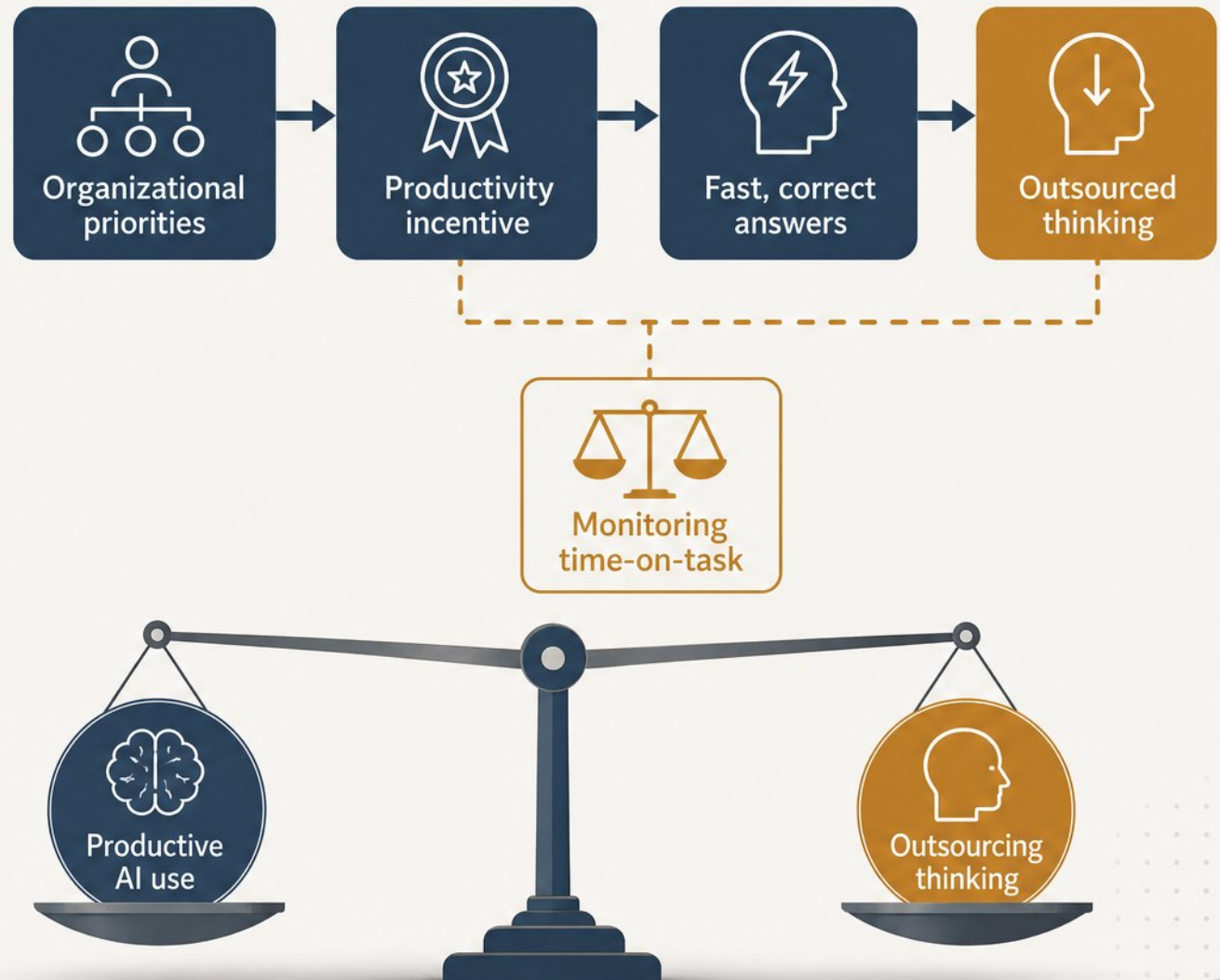


- Start small with clear success metrics

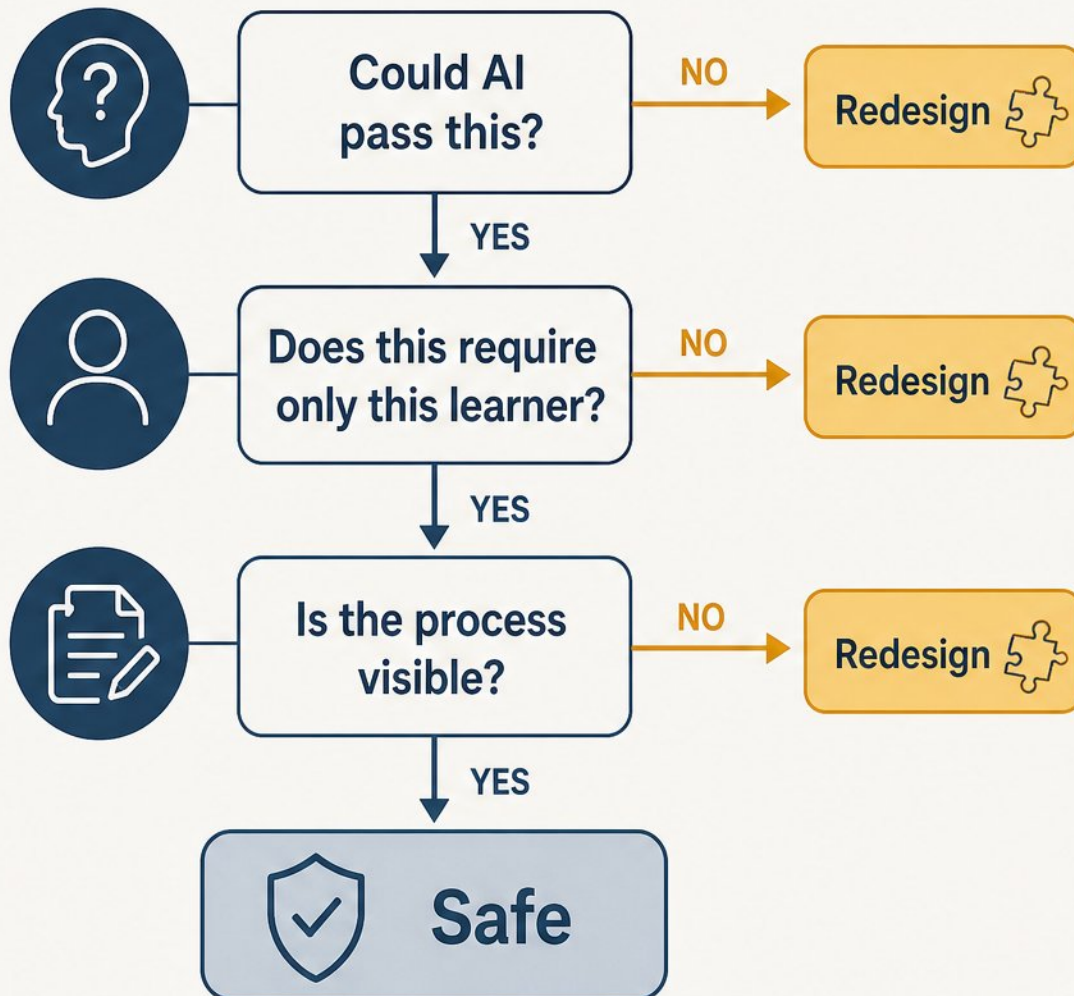


Leadership and Culture

- AI adoption is an organizational and incentive challenge, not a technical one
- Monitor time-on-task and frequent low-stakes quizzes, not just homework scores
- Explain to stakeholders: productive AI use vs. outsourcing thinking
- Students outsource when fast, correct answers are rewarded over effort



Decision Guide for Evaluating AI Use



- Ask: Could AI complete this and pass? If yes, redesign.
- Does the task require something only this learner can provide?
- Is the thinking process visible through drafts or annotations?
- Use specificity injection and process anchoring to protect learning.

Action Commitments and Next Steps



- Prioritize process over product: visible thinking, drafts, and reflection



- Monitor authentic signals: time-on-task and frequent low-stakes checks



- Distinguish dependent offloading (crutch) from autonomous offloading (scaffold)



- Redesign vulnerable assessments to require learner-specific, defensible reasoning



- Commit to engagement design: AI amplifies or substitutes based on how it is used



SHORTCUT

Easy now,
weak later.

PATH

Hard now,
stronger later.

CRUTCH

Replaces effort
and limits growth.

SCAFFOLD

Supports growth
and independence.



*Choose the path that builds capability.
Use supports that develop independence.*

PersonWise.

PersonWise • AI digital-human course platform



Scan the code or click anywhere to watch this course live, with voice Q&A

personwise.ai/t/a8a278ca/public