

The Multiply Method

Our learning-science framework — the tangible answer to what “learning-science-based” actually means.

Every claim in this document points at something you can check: a named researcher, a mechanism in the training modules, or a behaviour in the platform. That’s deliberate. Anyone can say “based on learning science” — this document exists so you can test whether we mean it.

THE ANSWER, IN ONE PARAGRAPH

Training pays only if people work differently afterwards — and that’s where almost all AI training fails: people finish knowing what AI can do, and keep working the old way. Learning science is the peer-reviewed answer to why — decades of research on what makes skill stick rather than fade, transfer to real work rather than stay in the classroom, and survive the weeks where new habits normally die. Whether training changes behaviour is decided by how the learning is *built*, not how good the content is. **The Multiply Method is that research engineered into the programme: ten principles — four that carry the outcome, six that shape every module.** The four: a **closed personalisation loop** — every end-of-module reflection is assessed against what that person’s own module taught, and the next module is rebuilt around the gaps and strengths it finds; **zero transfer distance** — from the first task, the exercises are the person’s actual job, so nothing has to “transfer” to real work later; **one-to-one at scale** — every module is generated individually for one person, the closest scalable answer to Bloom’s finding that one-to-one mastery tutoring beats classroom teaching by two standard deviations; and **engineering for the implementation dip** — the confidence drop where behaviour change normally dies, named, expected, and supported at the point it bites. Beneath those sit six structural principles governing how every module is built — from cognitive load management to a sequence where each module opens on a problem the learner has just hit for real. **And we’re as deliberate about what we leave out:** points, fluency drills, multiple-choice tests and learning-style tailoring don’t change adult behaviour, so none of them appear.

WHY THIS IS MORE THAN A GENERIC CLAIM

1. Named, public provenance

Every principle traces to peer-reviewed research — Bloom, Sweller, Ausubel, Vygotsky, Bruner, Bandura, Fullan, Black & Wiliam, Thorndike, Mezirow, Kapur, Bjork. Established science in the public domain, not a borrowed proprietary method. We didn’t invent it; we selected, combined and applied it deliberately.

2. The differentiator is adoption — and it’s shown, not claimed

Most AI training is a content dump. The hard problem is the **intention-action gap**: professionals learn what AI can do and still don’t change how they work. Plenty of providers use adoption language too — the difference is that ours is inspectable. Ask where any principle sits and we point to it: in the module text, in the platform’s behaviour, in the mechanism that runs on every learner.

3. The discard list carries operational provenance

The principles we rejected aren’t strawmen — they’re methods this team has run at scale in AI-powered K-12 education, where they belong. We kept them there and set them aside here because adult behaviour change is a different problem from exam performance. Knowing first-hand why each one works for children and fails for professionals is the rigour behind the selection.

1. The closed personalisation loop — formative assessment, automated

BLACK & WILIAM · FORMATIVE ASSESSMENT — WITH BJORK'S RETRIEVAL PRACTICE AND EBBINGHAUS'S SPACING DELIVERED THROUGH IT

The science. Formative assessment — assessing understanding *during* learning and feeding it back into instruction — is among the highest-impact interventions ever measured in education research. The loop also delivers the two most robust memory findings: retrieval practice (recalling beats re-reading) and spacing (revisiting at an interval combats forgetting).

In the product. At the end of every module the learner writes a reflection that opens with understanding checks answered *from memory* — retrieval, by design. That reflection is then assessed **against the personalised module the person actually received**, not a generic rubric: what landed, and the one or two gaps that would affect real-world application. The next module opens with a section built from that assessment — acknowledging demonstrated strengths, re-teaching the gaps before they compound, and bridging into the new material. The revisit happens at the natural interval between modules — spacing, delivered. By Module 7 the system holds the person's full documented arc. It's formative, not graded: no scores, no pass marks — the diagnosis works invisibly, repairing rather than blocking, which is the right design for adults.

“Most training ends with a feedback form. Ours **reads what you wrote, checks it against what your own module taught you, and rebuilds your next module** around what it finds.”

2. Zero transfer distance — trained on the day job

THORNDIKE · IDENTICAL ELEMENTS — THE TRANSFER-OF-TRAINING PROBLEM, DISSOLVED

The science. The documented failure mode of corporate training is **transfer** — skills demonstrated in the classroom that never survive contact with real work. Transfer rises with the similarity between the training task and the performance task; Thorndike's identical-elements principle, a century old and never overturned, says it's maximised when they share the same elements. Ours are not similar — they are the same task. This also delivers Bandura's strongest source of confidence (mastery experiences — wins earned on one's own real work) and Ausubel's schema activation (new skills attach to work the learner already understands deeply).

In the product. Onboarding captures the person's actual working world — their responsibilities, the people they work with inside and outside the organisation, the rhythm of their week — and pushes for depth. The first exercise names the client and project they typed in minutes earlier. From there, every exercise is a real deliverable — “use real work, always; never fictional scenarios” is a design rule of the curriculum. The module hours produce work product, not homework.

“The exercises are not homework beside the day job. **They are the day job.**”

3. One-to-one at scale — Bloom's two-sigma, answered

BLOOM · THE 2-SIGMA PROBLEM AND THE MASTERY-LEARNING TRADITION

The science. Bloom's famous finding: students tutored one-to-one with mastery methods perform two standard deviations above conventional classrooms — better than 98% of conventionally-taught students. His “2-sigma problem” was that nobody could afford it at scale. Personalised, mastery-based, individually-paced instruction is the most powerful configuration education research has found; it just couldn't be scaled — until AI.

In the product. Every module is generated individually, for one person, from their profile — the whole module rebuilt around their role, their people, and their demonstrated progress, not a template with a name inserted. Personalisation is architected from the first minute: at signup the system pre-plans how all seven modules will personalise for that person, then enriches the plan with every reflection. The pedagogy is protected while the scenarios adapt — the teaching sequence and objective of every exercise is preserved by design. Progression is genuinely sequential, and mastery is handled by **detection and repair**: the loop finds understanding gaps and re-teaches them inside the next module, the adult-appropriate form of Bloom's mastery learning. This is the tradition our founder built in before Multiply — AI-personalised mastery learning delivered at K-12 scale.

“Bloom showed one-to-one mastery tutoring beats the classroom by two standard deviations — and that no one could afford it. **Every module here is generated for one person.** That's the answer to his problem.”

4. Engineered for the implementation dip — where adoption usually dies

FULLAN · THE IMPLEMENTATION DIP — WITH BANDURA'S SOURCES OF SELF-EFFICACY

The science. When people apply new skills to real work, confidence drops sharply two to six weeks in before improving. Most interpret the dip as personal failure and quietly revert — which is where AI adoption normally ends. What carries people through is self-efficacy: built strongest by earned mastery experiences, and second-strongest by watching *similar* peers succeed.

In the product. Every module carries a Human Element section — the psychology addressed at the exact point it bites. Module 2, when new patterns feel confusing: brain rewiring is energetically expensive, and it gets easier. Module 4, at the dip's onset when workflows meet real work: the dip named and normalised, success redefined as consistency rather than speed, plus an optional forcing-function challenge. Module 6: the leapfrogging effect. Module 7: the reverting-under-pressure pitfall, confronted directly. Confidence is built the way Bandura says it must be — wins earned on the learner's own deliverables, acknowledgement that is specific and tied to actual performance. And for teams, the dashboard applies his second source as product: everyone sees what peers in their own team and role are doing that they aren't. It doesn't just measure adoption; it spreads it.

“People learn what AI can do and still don't change how they work — the intention-action gap. So the adoption layer isn't a pep talk at the end; **it's engineered into every module, at the exact week the research says confidence drops.**”

THE STRUCTURAL SIX — THE ENGINEERING INSIDE EVERY MODULE

	PRINCIPLE	RESEARCHER	WHERE TO POINT
5	Cognitive load, managed — and taught	Sweller	Full worked examples early, content chunked, examples fading from Module 3. And the mirror: Module 2 <i>teaches this same science to the learner</i> — managing an AI's context window is cognitive load theory applied to the machine. We manage your cognitive load while teaching you to manage your AI's.
6	The four-principle spiral	Bruner	<i>Context is everything · shared brain · AI creates outputs · human as supervisor</i> recur in all seven modules at rising depth — never taught once and dropped. Every module closes by naming how each principle just deepened.
7	Fading scaffolding on a fixed schedule	Vygotsky · Kalyuga	80% direct instruction (Modules 1-2) → 60% (3-4) → 40% (5-7). Exact-steps setup in Module 1; outcome-only briefs by Modules 5-6. What helps novices hinders the proficient, so support is withdrawn deliberately.
8	Advance organisers at every level	Ausubel	Every module and every part opens above the content's abstraction level, hooking new material onto what the learner already knows — and the loop builds one <i>personally</i> for each learner at the start of each module.
9	Experience before explanation	Mezirow · Engelmann	Three knowledge types, taught three ways: procedures by direct instruction; the paradigm shift by experience first — the Module 1 “magic moment” lands <i>before</i> the principles are named (“let's understand why that worked”); adaptive expertise through varied practice later.
10	The barrier-driven sequence	Kapur's mechanism, at programme level	Each module opens on a failure the learner has just genuinely hit in their own work — context fills up after Module 1; the workspace turns messy after Module 2. The gap is felt, and instruction lands as the answer to a question the learner has already formulated.

The three-way distinction in row 9 matters: procedural knowledge (taught directly), paradigm knowledge (experienced first), and adaptive expertise (built through varied practice) each need different teaching. Most training treats all three the same — which is why tool tutorials produce so little behaviour change.

These aren't methods we overlooked — several are methods this team ran at scale in K-12, where they work. We set them aside here because adult behaviour change is a different problem from exam performance:

- ✗ **Extrinsic motivation (XP, points, penalties)** — powerful for children; adults are motivated by watching their own work improve.
- ✗ **Automaticity / fluency drills** — built for recall *speed* (times tables); paradigm shift is about consistent *choice* of a new approach, not speed.
- ✗ **Interleaving of problem types** — an academic-assessment technique with limited transfer to workflow learning.
- ✗ **Multiple-choice / percentage testing** — we assess understanding through written reflection on demonstrated real work, not quiz scores.
- ✗ **Standardised-exam validation** — designed to predict test results; irrelevant to professional adoption.
- ✗ **Learning-style tailoring** — the research is clear that learning styles don't predict outcomes. Everyone learns better when both channels are engaged, which is why module videos pair spoken teaching with on-screen demonstration — dual coding, not “videos for visual learners.”

ASK US TO SHOW YOU

Name any of these and you're describing something real in the product — ask, and we'll show you where it sits:

- **The closed loop, end to end.** Reflection with from-memory understanding checks → assessed against that person's own module → gaps and strengths recorded → the next module opens on them. Every learner, every module.
- **Personalisation architected at signup.** The profile pre-plans how all seven modules will personalise for this person, then enriches with every reflection.
- **Whole-module regeneration, pedagogy protected.** The entire module is rebuilt for the individual; the teaching sequence and exercise objectives are preserved by design.
- **The four-principle spiral** across all seven modules, with each closing summary naming the deepening.
- **The “WHY” pattern.** Principle → why it matters in this context → the observable outcome — never an abstract statement.
- **Advance organisers** at module and part level throughout, plus the personalised bridge built per learner.
- **Fading scaffolding** on the fixed 80/60/40 schedule.
- **A Human Element section in every module** — the dip addressed at its point of onset.
- **Sequential progression** — each module locked until the previous is complete and reflected on.
- **Reflective self-checks** — a checkpoint in every module asking the learner to verify capability against concrete demonstrations before moving on.

The invitation stands. Every claim in this document points at something inspectable. If you're evaluating AI training, put the same question to every provider: “show me where it sits.” We'll always have an answer.