



The AI Capability Crisis in Pharma

Why 'Trained Once' and 'Trained for Everyone' Both Fail in Regulated Work - and What Holds

By Dr. Andrée Bates, Founder, Eularis

For leaders in biotech and pharmaceutical organisations of every size – from small to large - and from single-function teams to global functions.

The AI Capability Crisis in Pharma

Why 'Trained Once' and 'Trained for Everyone' Both Fail in Regulated Work - and What Holds

For leaders in biotech and pharmaceutical organisations of every size – from small to large - and from single-function teams to global functions.

Dr. Andrée Bates
Founder, Eularis

Executive Summary

The pharmaceutical industry is investing heavily in AI training. Workshops are full. Post-session surveys are glowing. Learning leaders are reporting scores to sponsors with confidence.

Here is what the reporting misses.

The training works - and then it erodes.

Not all at once, not for everyone, and not evenly. Within roughly ninety days, part of what a team learned has faded and part of what they still remember has gone out of date. Neither of those is a failure of the training. Both are the predictable consequence of buying an event to solve a maintenance problem.

There are two ways AI training fails in pharma, and the second is now the more common.

The first is the one this paper is named for: training that worked, and then faded.

The second is training that never fitted the job. An organisation-wide AI literacy programme, delivered on an ongoing basis and fully documented, addresses the first and can leave the second entirely untouched.

That matters more than it used to, because the AI literacy duty in the EU AI Act is not framed around whether training happened. It is framed around measures appropriate to the context in which a person uses the system. And this is not just relevant to EU companies. The Act makes this clear.

Section 1.1 takes this on directly, because it is the position most larger biotech and pharmaceutical organisations are now in.

The two failures are not evenly distributed by size of organisation.

- Larger organisations have mostly solved the first with an ongoing company-wide programme and are now exposed to the second.

- Smaller biotechs rarely have the second problem at all, because they never bought the company-wide programme - they have the first problem in its purest form, and usually no one whose job it is to notice.

The remedy is the same at both ends, and section 5 runs the arithmetic for a fifteen-person team as well as a five-hundred-person function.

This paper makes a direct argument: trained once a year is not a capability strategy. In pharmaceutical and biotech settings it can produce something more dangerous than a skills gap - a false sense of readiness. It raises usage without reliably sustaining competence, and it can introduce compliance risk at precisely the point an organisation believes the risk has been addressed.

The erosion is not a matter of opinion. It is one of the better-measured phenomena in organisational psychology, and two meta-analyses cover different halves of it. On the procedural side - knowing how to run a workflow - a 2025 meta-analysis in Psychological Bulletin pooled 1,344 effect sizes from 457 reports and found that half the initial gain in accuracy-dependent procedural skill is lost at around six and a half months of non-use, with speed-based skill holding roughly twice as long.¹ On the judgement side - deciding whether an output is good enough to rely on - an earlier review of 189 effect sizes across 53 studies found that below ninety days of non-use the evidence is mixed, while beyond ninety days decay appears consistently, and that cognitive, accuracy-dependent work erodes far faster than physical or speed-based work.² Applied AI work in pharma spans both halves, and loses the more consequential half faster.

That is the ordinary curve, and it applies to every kind of training. AI adds a second erosion on top of it, because the subject itself keeps moving. Together they produce a working half-life for applied AI knowledge that I estimate at about ninety days in pharmaceutical settings.

Drawing on that evidence, the pace of change in AI tools and governance, and the operational realities of regulated pharmaceutical work, I argue that capability must be maintained, not periodically repurchased. I examine the hidden costs of episodic AI learning and set out what a sustained capability model looks like in practice: designed around the short shelf life of AI knowledge, shifting regulatory expectations, and role-specific application in real workflows.

The organisations that lead in AI-enabled pharmaceutical functions will not be those that purchase the most training. They will be those that build, govern, maintain and document capability over time, by function, not simply general AI fluency.

The two most useful questions a pharma leader can ask an AI training provider are not “how good is the session?” They are “what happens in month four?” and “what does this look like for a regulatory writer rather than a brand manager?”

1. Introduction: The Room That Empties

Every learning leader in pharma knows the pattern - and in smaller organisations the same pattern shows up without anyone whose job it is to notice - the workshop happens, the enthusiasm is real, and there is no one holding the follow-through. An AI webinar, keynote or workshop goes well. Attendance is high. The discussion is active. The post-session feedback is positive: high ratings, strong comments, visible enthusiasm. The sponsor receives the results and the investment appears justified.

Then the room empties. Over the following weeks, some of the learning does too.

This is not a new phenomenon, and it is not a criticism of the training. Learning science has shown for decades that trained capability declines without reinforcement.

What is new in the AI context is that knowledge does not only fade from memory; it also loses relevance. Tools change. Interfaces evolve. Policies tighten. Use cases mature. Regulations change. Even when people remember exactly what they were taught, what they learned may no longer be sufficient for the work in front of them.

In pharmaceutical and biotech environments that carries greater consequence than in most sectors. Weak AI practice in regulated workflows does not simply produce mediocre output. It can create documentation errors, inappropriate tool use, inconsistent review practice, audit vulnerability, and a particular form of confidence without competence.

For that reason the core AI question in pharma is not a training question. It is a capability-design question. The distinction matters. Training is an event. Capability is a system.

The industry already understands this principle everywhere else. Pharmacovigilance, regulatory writing, GxP compliance and field medical development are not managed through isolated learning interventions. They are supported by repeatable structures: standards, reinforcement, oversight, practice and accountability. Yet AI is still frequently approached as though a system-level challenge could be closed with a workshop.

This paper examines why that approach underdelivers, where its hidden costs emerge, and what a maintained AI capability model looks like instead.

1.1 If you already have an ongoing AI programme

Many pharmaceutical organisations now do: a company-wide AI literacy programme, refreshed periodically, open to everyone, with completion tracked and recorded. If that describes your organisation, you are ahead of most of your peers, and the section above will have read as though it were about somebody else.

It is worth pausing there, because that programme solves one problem completely and another not at all.

What it solves is availability. Content does not sit ageing on a shelf. New joiners can be enrolled without restarting anything. There is a record. Measured against the workshop model, that is a real advance.

What it does not solve is fit. A company-wide curriculum is, by construction, the same curriculum for everyone. It teaches what AI is, what it does well, where it fails, and what the acceptable use policy says. Those things are worth knowing and they are identical for a medical writer and a supply chain analyst. What is not identical is the work. The

judgement a regulatory writer needs in order to decide whether a generated summary faithfully represents its source has almost nothing in common with the judgement a brand manager needs in order to decide whether a generated claim is substantiated. Neither is taught by a general course about AI in pharma.

The regulatory position is sharper than most organisations realise.

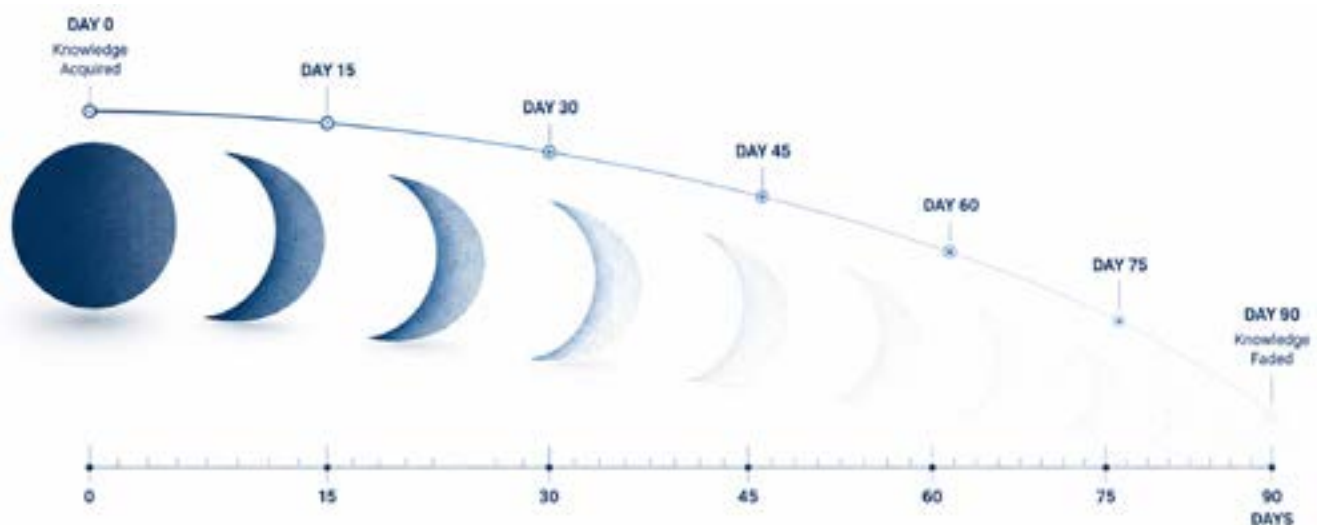
The AI literacy duty under Article 4 of the EU AI Act is not framed around whether documented training took place. It is framed around measures appropriate to people's skills, experience and education, and to the context in which systems will be used, and guidance indicates that generic AI training sessions are unlikely to be sufficient against a standard framed that way. Section 3.3 sets out the current position in full, including what changed in July 2026.

Which makes the documentation point uncomfortable rather than reassuring. A completion record for a generic programme documents, precisely, that a generic programme was completed. It is evidence of a measure taken. Whether it is evidence of a measure appropriate to the context is a separate question, and it is the one that matters now.

For a leader who has done the sensible thing and put an ongoing programme in place, the uncomfortable version is this. The tick in the box is real. **It is just a thinner tick than Article 4 describes.**

2. The Ninety-Day Half-Life: How AI Knowledge Actually Decays

Trained once underdelivers for structural reasons, not executional ones. Two forces drive it. First, people lose what they do not use. Second, AI tools, interfaces and governance expectations change quickly - a Copilot agent may have one interface in one week and a materially different one after an update. Together, these forces sharply reduce the useful life of a single training event.



2.1 The first force: how trained capability erodes

Learning science has shown consistently that capability fades when it is not reinforced through practice, retrieval and application. A strong workshop creates awareness, energy and real short-term competence. On its own it rarely creates durable capability.

The most recent evidence is precise about the shape of that decline, and precise about its scope. A 2025 meta-analysis in *Psychological Bulletin* examined procedural skill retention - the how-to layer - pooling 1,344 effect sizes from 457 reports and, for the first time in this literature, treating the retention interval as a continuous variable rather than a set of bands. The result is a gradient rather than a cliff: half the initial gain is lost at around six and a half months of non-use for accuracy-dependent skill, at roughly thirteen months for speed-based skill, and at about eleven months for mixed skill. The coded task categories include medical and dental procedures, so this is not a laboratory artefact.¹

Read that carefully, because it is not what most commentary on corporate training claims. It does not say the training failed. It says the gain was real, and then eroded, at a measurable rate, from the day the session ended.

That covers the procedural half. It does not cover judgement, and judgement is where pharma's exposure sits. For that, the relevant work is earlier. Pooling 189 effect sizes across 53 studies, Arthur and colleagues found two things that matter most for this industry.²

The first is about timing. Below ninety days of non-use, the evidence is mixed: some studies find decay and others do not. Beyond ninety days it appears consistently.²

Ninety days is therefore not the point at which knowledge disappears. It is the point after which you can no longer assume it is there.

Their second finding is the uncomfortable one. Decay is not uniform across kinds of work. Physical and speed-based skills hold up comparatively well. Cognitive and accuracy-based skills - those where being right matters more than being fast - decay around three times faster.²

Now separate what pharma teams actually do with AI into those two halves. Running the workflow - supplying the right source material, structuring the request, working a draft into usable shape - is procedural, and it decays on the curve the 2025 paper measures. Deciding whether a generated summary faithfully represents its source, whether a drafted response is defensible, or whether an output is good enough to put in front of a reviewer, a payer or a regulator is not procedural at all. That is cognitive, accuracy-dependent judgement, and it sits in the fastest-eroding category either paper identifies. An organisation therefore loses the two halves at different rates - and loses the half that carries the regulatory consequence faster.

2.2 What this looks like from inside a real programmes

The pattern repeats across programmes, and its two halves always arrive together.

Two months after a programme, messages are still coming in from the room: evidence synthesis that had taken weeks coming back in hours, a dossier rebuilt, someone who has redesigned an entire workflow around it. The training has clearly taken.

In the same period, from the same room, other people write asking for a reminder of how to do something specific they were shown - and say they would welcome another session. That request is not a criticism of the training. It is the most accurate read of the situation anyone in the room offers.

Both are true at once, and neither cancels the other. That is what erosion looks like in practice: not a cliff everyone falls off together, but an uneven decline that takes different capabilities from different people at different rates. The daily habits hold. The occasional workflows go first. And note the interval - two months, well inside the window in which the literature says retention should still be likely.

This is not one unusual room. I see it in regulatory affairs teams, in commercial functions, in market access, and in rooms of external clinicians: the messages about what is working and the requests for a refresher arrive from the same programme, in the same time period, from different people.

Which points to the mechanism that matters most commercially. Decay is driven by non-use - both meta-analyses identify how often a capability is actually exercised as a key determinant of whether it survives.^{1 2} The tasks a team performs daily will hold. The tasks they perform quarterly will not.

Now look at where AI creates the most value in pharma. The annual payer dossier. The periodic safety update. The submission that comes round twice a year. The launch that happens once. Those are precisely the workflows where AI saves the most time - and precisely the ones a team will have lost the working knowledge for by the time they next need them.

The highest-value use cases erode fastest. That is not a training failure. It is arithmetic.

Client examples in this paper are drawn from multiple programmes and generalised. No organisation or individual described in them is identifiable.

2.3 The second force: how AI moves

Everything above is the ordinary curve. It applies equally to project management, negotiation or aseptic technique, and every learning function in pharma already knows about it.

AI adds a second erosion on top, and this one has no published literature yet, because nothing has moved this fast before.

Ordinary decay assumes the object of the capability holds still. Aseptic technique does not change while you forget it. AI does. Models change. Features evolve. Interfaces improve. Capabilities that did not exist when a curriculum was designed appear in the tools people already have open. Governance expectations shift as organisations adapt policies, controls and standards.

The result is not only capability loss. It is relevance loss. Teams may remember precisely what they learned and still find that parts of it no longer apply. In practice this shows up as:

- prompting techniques tuned to older model behaviour
- workflows built around features that have since changed
- assumptions about model strengths and limits that no longer hold
- guidance that no longer reflects the current control environment

This is what makes AI capability uniquely perishable: retention declines at the same time as the target keeps moving. And it produces a failure mode that ordinary decay does not - a team that retained its training perfectly can still be wrong, because what they retained is no longer true.

2.4 Why ninety days, and what I can and cannot claim

I want to be exact about this figure, because there are different erosions for different elements of the training.

The research puts memory-side erosion alone at roughly six and a half months to lose half the initial gain for accuracy-dependent procedural skill, and faster than that for cognitive, accuracy-dependent judgement.^{1 2} That is the first term, and it is measured.

The second term is the one section 2.3 describes, and it has no equivalent evidence base. Nobody has measured how quickly applied AI knowledge stops being true, because the phenomenon is too new and the thing being measured does not hold still long enough to be studied. What can be observed is how often the underlying tools, interfaces and governance expectations change materially. In pharmaceutical settings over the past two years, that interval has been closer to a quarter than to a year.

Put the two together. Memory-side erosion begins the day the session ends, and runs faster for judgement than for procedure. Relevance loss runs alongside it on a shorter cycle, and it does not wait for anyone to forget anything. The point at which a manager should stop assuming that what was taught is both still remembered and still correct therefore arrives well before either curve would reach it alone. Ninety days is where I place that point.

That is a practitioner's estimate, not a published finding, and I will happily revise it against better data. What is not in question is the shape: two erosions, compounding, in the fastest-declining category of work, concentrated on the tasks that matter most and are performed least often.

Nor is all AI knowledge equally perishable. It is more useful to think in three layers:

- The tool layer - specific features, the exact way to elicit a useful result from a given model, current limitations and workarounds - can be stale within weeks.
- The applied layer - how AI fits a particular workflow, what it is useful for in a given role - has a longer useful life.
- The foundational layer - how these systems behave, where they fail, how to evaluate outputs critically - lasts longest of all.

This has a clear implication. Effective AI capability building should refresh tool knowledge frequently, update workflow use cases on a regular cadence, reinforce foundational judgement over time, and embed learning into day-to-day work.

A single session does the opposite. It often concentrates effort in the fastest-changing layer and provides no follow-through. That is why it generates genuine enthusiasm and genuine early results, and still does not produce lasting capability.

And none of it feels like anything from the inside. The people were trained. They have the certificate. They are confident. Confidence erodes slowest of all - which is the part that should concern a regulated business, because the failure mode is not someone who knows they are unsure. It is someone who is certain, and out of date.

3. Why Pharma Bears a Disproportionate Cost

Every industry invests in AI training, and every industry loses value when that capability fades. Pharma is different because the downside extends well beyond productivity. AI is being used inside regulated workflows where weak judgement, incomplete review or outdated practice can create compliance exposure, quality risk and documentation gaps.

3.1 Regulated, consequential, and embedded in core workflows

Pharma teams do not use AI only at the edges of the business. They increasingly use it inside workflows that are regulated, documented and tied to high-consequence decisions - clinical operations, pharmacovigilance, regulatory affairs and submissions, medical information, MLR review, market access and payer materials, commercial content development, and cross-functional scientific documentation. And this is not just in big companies. In a fifty-person biotech running a first submission carries the same regulatory exposure per document as a global company, with none of the internal review depth to catch a problem.

In many industries, weak AI use leads to rework. In pharma it can also lead to unsupported or weakly substantiated claims, flawed summaries or incomplete interpretation of evidence, errors embedded in regulated documentation, increased review burden for already constrained teams, and downstream compliance exposure.

3.2 Trained once can increase exposure, not reduce it

The hidden risk in AI training is not low adoption. It is false confidence.

A well-run session succeeds on its own terms: it raises awareness, lowers resistance and gets people using the tools. In regulated environments that is necessary but not sufficient. If usage rises faster than durable judgement is sustained, the organisation can become more exposed, not less.

That gap shows up in familiar ways: outputs accepted because they appear polished; human review becoming procedural rather than critical; teams assuming prior training remains sufficient as tools and policies change; managers reading usage as evidence of readiness.

This is not an argument against training. It is an argument against treating training as a completion. A single session is the beginning of capability development, and without reinforcement it can create a misleading picture of organisational readiness.

3.3 The moving regulatory landscape

Capability decay is more costly in pharma because the external standard does not stand still - and as of 2026, one part of that standard addresses capability directly.

Article 4 of the EU AI Act requires providers and deployers of AI systems to take measures supporting the development of AI literacy among staff and others operating AI on their behalf, taking account of those people's skills, experience, education, and context in which the systems will be used.³

That is new. The duty itself has applied since 2 February 2025, but Article 4 was rewritten in full in July 2026 by the Digital Omnibus on AI, in force from 27 July, replacing a duty to ensure, to their best extent, a sufficient level of literacy, with a duty to take measures supporting its development, and adding an explicit clarification that no organisation must guarantee any particular individual's level. National supervision and enforcement powers took effect the following month. At first reading

the rewrite looks like a relaxation. It is better understood as a change of kind. What it lightens is the result, and what it leaves is the effort.⁴

An obligation to ensure a sufficient level, even qualified by best efforts, is oriented to a result. It can, in principle, be discharged by a single intervention that demonstrably raised capability. An obligation to take measures supporting the development of literacy is an obligation of effort - and effort is not a state an organisation arrives at. It is something demonstrated continuously, by what is done and by what can be shown to have been done.

Three features of it matter more than the headline.

First, its scope. Most obligations under the Act attach to high-risk systems. **Article 4 does not.** It applies to any AI system an organisation provides or deploys, **at any risk level** - which for most pharmaceutical companies means the general-purpose AI already in daily use across every function, not a narrow set of validated clinical tools. The same Omnibus that softened the literacy standard deferred the high-risk obligations sitting alongside it: standalone Annex III systems to December 2027, embedded Annex I systems to August 2028. Article 4 was not deferred. For an organisation whose AI use is general-purpose rather than validated high-risk systems, it is an obligation under the Act that is enforceable now.

Second, its reach. Organisations established **outside the EU fall within scope** where the output of their AI systems is used inside it, which captures most global biotech and pharmaceutical operations regardless of where the work is performed.

Third, what satisfies it. The context framing survived the rewrite intact: measures must still account for the skills, experience and education of the people concerned, and for the context in which the systems are used - which is where role enters, since the context in which someone uses an AI system is not separable from the job they do. A generic AI awareness session is a measure. Against a standard framed that way, it is a thin one. There is no requirement to measure literacy levels, and now an explicit statement that no individual's level need be guaranteed. What remains is the expectation that an organisation can show what it has actually done.

Read together, those features describe an obligation discharged by sustained, context-specific measures and evidenced by a record of them. A historical training certificate is evidence of one measure, taken once. It is not evidence of effort maintained.

The rest of the regulatory picture reinforces the same logic from different directions. Article 14 requires high-risk systems to be overseen by natural persons with the necessary competence, training and authority to intervene - again a present-tense capability rather than a historical one, though the obligations it attaches to now apply from December 2027 for standalone high-risk systems.⁵ Article 50, by contrast, has applied since August 2026 and reaches far further into everyday work: it governs the labelling of synthetic content and disclosure when a person is interacting with an AI system, which for pharmaceutical organisations touches AI-assisted promotional and patient-facing material directly.⁶ The FDA's guidance on AI supporting regulatory decision-making is built around demonstrating that a model is credible for its specific context of use rather than in general; as at August 2026 it remains in draft, and is expressly non-binding.⁷ In January 2026 it was complemented by ten guiding principles issued jointly with the EMA, which state that AI should support rather than replace human regulatory decision-making.⁸ The EMA's reflection paper places responsibility for validating, monitoring and documenting model performance with the marketing-authorisation holder.⁹ In each case the burden sits with the organisation, continuously.

The practical questions are increasingly operational: where AI is being used, what level of human oversight is applied, what documentation exists for AI-assisted work, who is accountable for decisions made within regulated workflows - and now, what evidence exists that the people doing that work are equipped to do it.

A historical training record does not answer any of them. It proves only that training once occurred. A team trained months ago may still be operating with outdated assumptions about tool behaviour, incomplete understanding of acceptable use, inconsistent review practice, weak documentation discipline, and limited awareness of how expectations have evolved. From an inspection perspective, the question is not whether people attended training. It is whether current AI use is controlled, defensible and aligned with present expectations.

Regulatory note. *This paper describes the regulatory position as at August/early September 2026 and is provided for general information only. It is not legal advice and it is not a compliance assessment of any organisation. The EU AI Act, the guidance issued under it, and the FDA and EMA positions referred to here all continue to develop, and how any of them applies depends on an organisation's own systems, jurisdictions and circumstances. Readers should confirm their own obligations, and how best to meet them, with their internal legal, regulatory and compliance functions. No training or capability programme, including the Pharma AI Enablement Institute, can of itself make an organisation compliant with Article 4 or with any other provision of the Act.*

3.4 The evidence layer: what you would have to show

Capability decay becomes most costly where it intersects with evidence. Two distinct records are involved, and organisations routinely maintain the first while neglecting the second.

The first is the record of AI use itself. AI increasingly leaves a trail an organisation may need to explain, defend or reconstruct - AI-assisted drafting in submissions and labelling; support for protocol development and evidence synthesis; safety narratives and case documentation; AI-assisted content entering MLR review; payer-facing materials supported by AI-generated analysis. In each case the organisation may need clear answers to basic questions: how AI was used, what human review took place, how the output was validated, whether use complied with internal policy, and what documentation standard applies.

The second is newer, and thinner in most organisations: the record of what the organisation has actually done.

An obligation of effort is not proved by asserting that people are competent. It is proved by showing the measures taken - which in practice means a per-person training record, held at a granularity that lets a functional sponsor see the state of their own team rather than an organisation-wide completion percentage. It is a specific and unglamorous piece of infrastructure, and most companies discover they do not have it at the moment someone asks for it.

The two records answer different questions. The first asks whether a given piece of work was done properly. The second asks whether the people doing it were equipped to do it. An organisation can be strong on the first and still unable to answer the second - and it is the second that Article 4 speaks to.

When teams work from months-old training without updated guidance, a gap opens between actual practice and documented expectation. That gap does not appear in completion rates or satisfaction scores. It appears when work is reviewed under scrutiny.

3.5 What the route to reliable judgement actually looks like

The case so far has been structural. It is worth setting one practitioner's experience beside it, because the sequence people actually move through with these tools is not the sequence training is usually designed around.

Dr Myriam Cherif spent fourteen years in large pharma in affiliate and above-country roles - Novartis in the UK, then Takeda and GSK in Dubai - moving from cardio-metabolic and respiratory medical affairs into oncology and haematology leadership. Latterly she was Director, Regional Medical Affairs Lead for oncology across GSK's Emerging Markets,

accountable for medical strategy and activity spanning Latin America, Turkey, the Middle East, Africa and India, including publication development, advisory boards, promotional material review and approval, and the medical training of affiliate teams. She now runs Kalyx Medical, a medical affairs advisory practice working with individual professionals and with teams, where AI workflows form part of what she builds. Describing her own path with AI, she is clear that it was not a before-and-after. It was five stages.¹⁰

1. **Dislike.** AI seemed too generic to be worth the effort.
2. **Learning.** Better prompts produced better outputs - followed by a plateau: what else can I actually do with this?
3. **Daily use.** Exploring, trying different tools. It clicks, and she loves it.
4. **Over-reliance.** Loving it so much that, without noticing, she began outsourcing her thinking to it. A couple of weeks.
5. **Sparring partner.** Knowing what to keep and what to ignore - and, in her own words, still learning.

Four features of that sequence matter for this paper.

Stage one is a finding about training design, not about the person. AI was dismissed as too generic to be worth using - a judgement about the content she had encountered rather than about the technology. Section 6.2 returns to this point: generic AI literacy programmes underperform not because participants lack curiosity, but because generic content gives a specialist no reason to persist past the first disappointment.

The same pattern shows up from the other side. Medical affairs professionals across a range of companies, having worked through function-specific AI guidance rather than a general programme, have said the internal training covered basics like prompting technique competently but offered nothing specific to their function - and that the use cases were what made the difference. That is a handful of people, not a study. It is the mechanism in stage one described by the people it happened to, and it identifies what the generic programme was missing: not depth on the technology, but a starting point in the work.

Stage two is where most organisational programmes stop. The plateau she describes - prompting improved, and then what? - falls precisely where a workshop ends and where a reinforcement structure would have to begin. She moved past it through sustained self-directed daily exploration. That is not a realistic expectation of an entire business unit with a day job.

Stage four is this paper's central risk, reported from the inside. Over-reliance did not arrive through neglect. It arrived through enthusiasm, at the point of peak confidence and peak usage, and it was invisible while it was happening - she names it only in retrospect. That is the confidence-without-competence failure mode described in 3.2 and priced in 4.3, with one detail the abstract version cannot supply: a person in this stage does not feel uncertain. They feel excellent. On a post-session survey they will report high satisfaction, high usage and high confidence - the three indicators most organisations currently read as success.

Stage five is a maintained capability, not an arrival. Knowing what to keep and what to ignore is a precise description of the cognitive, accuracy-dependent judgement that Section 2 identified as the fastest-eroding category of skill in the literature. It took extended daily engagement to build, it is held in place by continued use, and she is explicit that it is still developing. No single session produces it, and nothing preserves it automatically.

The stage that carries the most risk is also the stage that feels the best.

It is worth being precise about what her route required: months of daily, self-directed practice, and a detour through over-reliance that she detected and corrected herself. A maintained capability programme is an attempt to engineer that same trajectory deliberately - faster, across a whole function, and with the fourth stage named in advance rather than discovered afterwards.

Five-stage progression shared by Dr Myriam Cherif, PhD - founder, Kalyx Medical; formerly Director, Regional Medical Affairs Lead, Oncology, Emerging Markets, GSK - and used with her permission.¹⁰

4. The Three Hidden Invoices

Most organisations cost AI training by visible spend: session fees, participant time, and where relevant travel and logistics. Those costs are real but they are not the full economics. The larger costs sit elsewhere, and they determine whether training compounds or simply repeats.

4.1 The repurchase invoice

Because capability erodes, an event-based model creates a repurchase problem. Organisations buy fresh interventions to restore ground that previous sessions could not hold on their own.

Over time this produces a weak economic pattern: budget spent repeatedly against the same underlying gap, each intervention producing a lift that fades, and limited cumulative value carried between cycles. Over a multi-year horizon, repeated introductory workshops can consume budget comparable to a sustained programme while retaining considerably less.

4.2 The generic-content invoice

The second hidden invoice is content that was never about the job. Training can succeed as an event, and be perfectly maintained thereafter, and still not change how work is done.

Generic examples do not translate into function-specific use cases. That single sentence accounts for more wasted AI training spend in pharma than every other cause combined. People may leave more aware, more interested and more capable, but if the learning is not applied to their actual tasks, within their actual workflows, under their actual constraints, and then reinforced, behaviour reverts. Teams return to established ways of working. Prompts are never integrated into daily work. Managers see participation but not execution change.

This invoice is the one an always-on generic platform does not clear. Enrolment is high, completion is high, the reporting is excellent, and the submission still gets written the way it was written before. Satisfaction scores can show a programme was engaging and useful. They cannot show whether the organisation now operates differently.

4.3 The confidence-without-competence invoice

The third invoice is the cost of unmanaged confidence. Training increases willingness to use AI. If competence, guidance and oversight do not keep pace, usage expands faster than control.

In most industries that misalignment creates productivity risk. In pharma it creates governance risk. The danger is not that people use AI imperfectly. It is that they use it with the appearance of readiness inside regulated workflows - producing greater review burden, inconsistent validation, documentation gaps, and remediation costs when issues surface later.

The five-stage progression in 3.5 sharpens why this invoice is so hard to see coming. The period of heaviest over-reliance was also the period of greatest enthusiasm, and from the inside it produced no signal at all. There is no post-session survey question that catches it.

5. The Reframe: From Training Events to Maintained Capability

The issue is not poor intent. It is a category mistake: asking an event-based investment to produce a maintained capability. That is the wrong mechanism for the job.

There are, in fact, two axes here rather than one, and separating them explains why organisations that have already fixed the obvious problem are still not getting what they paid for. The first axis runs from episodic to maintained. The second runs from generic to role-specific.

An episodic generic intervention is the workshop. An episodic role-specific one is a bespoke consulting engagement: excellent, expensive, and gone when the consultant leaves. A maintained generic programme is the company-wide platform subscription, which is where most of the industry has now arrived and where most of it has stopped. Only the fourth combination, maintained and role-specific, produces capability that both persists and applies.

Most organisations have made one of the two moves and treated it as the destination. Moving from episodic to maintained solves availability. It does not, on its own, change what anybody does on a Tuesday.

5.1 Pharma already knows how to do this

Pharma does not validate a system once and consider the work complete. It does not run a single pharmacovigilance check and assume the function is covered. It does not train an MSL once and assume the capability holds indefinitely.

The operating logic is consistent across every critical function: important capabilities are built deliberately, applied to real workflows, maintained as conditions change, and governed so they remain defensible. That is not a philosophical preference; it is an operating requirement.

AI belongs in the same category. It is cross-functional, fast-changing, and increasingly embedded in work affecting quality, speed, judgement and compliance. Yet one of the fastest-moving capability areas in the business is still routinely assigned one of the lightest maintenance models. In a regulated environment, that is a structural gap.

5.2 The economic case, with the arithmetic shown

Once AI capability is understood as maintenance rather than event purchasing, the comparison changes. The figures below are illustrative and should be replaced with your own; the point is the shape, not the precision.

Consider a team of 50 people. Under an event model, an organisation typically commissions an introductory programme roughly twice a year to restore momentum, plus internal delivery time and the coordination cost of arranging each one. Under a maintained model, the same budget buys a continuous programme with live sessions, a library that is added to monthly, and ongoing governance updates.

On the return side, the threshold is lower than most leaders expect. If a maintained programme helps each of those 50 people recover 50 minutes a week through better drafting, synthesis, search or meeting preparation - and that time is redeployed to higher-value work - the annual recovered time is roughly 2,000 hours across the team. At a fully loaded cost of \$100 an hour, that is around \$200,000 of redeployed capacity a year.

The arithmetic scales in the direction that helps larger teams and tests smaller ones. A 500-person function clears the same bar at around five minutes per person per week. A 15-person biotech needs proportionally more from each person - but it is also buying for fifteen people rather than five hundred, and in a company that size the same fifty minutes a week is frequently recovered by one person on one workflow.

What this means: over a three-year horizon, a sustained programme frequently costs less than the accumulated event spend it replaces, and it is the only one of the two that produces compounding capability. Run the comparison against your own numbers before you accept it - the calculation is straightforward, and it is the one your finance partner will ask for.

This is not primarily a technology ROI story. It is a productivity story, with one condition: the capability has to persist long enough to change how work gets done.

5.3 The question that changes the buying decision

The wrong question is: how good is the session? The better questions are: what happens in month four, and what does this look like for a regulatory writer rather than a brand manager?

If the answer is simply another session, the organisation is still buying activity. A credible model should be able to explain how capability is maintained after the initial event - how teams stay current as tools, policies and use cases evolve; how general awareness becomes role-specific workflow change; how governance and acceptable-use guidance are reinforced over time; how new joiners are brought in without restarting the programme; and how progress is measured in adoption, output quality and risk control rather than attendance.

That structure is where the durable value sits. Not in the polish of the opening session, but in the system that keeps capability current, applied and governed after it.

6. What Maintained AI Capability Looks Like in Practice

6.1 What to require

A maintained programme is not a longer workshop. It is a different structure, designed around the dynamics of AI knowledge and around how each function actually uses it, rather than around the convenience of event-based procurement. It has to be specific to the business unit and to the role.

The role-specificity is structural, and it exists before any client is assessed. The curriculum is not one syllabus with worked examples swapped in. It is a foundational layer covering what is genuinely common across regulated pharmaceutical work, and then a separate track for each business unit, built around the work that unit actually performs. Routing is by function, not by an AI knowledge assessment. Everyone starts with foundations, and then a medical affairs team does medical affairs, a regulatory team does regulatory, a supply chain team does supply chain. No assessment is required to work that out, and none is used for it.

On that structure sits the core training: the foundational and applied function layer giving teams genuine first competence in the context of their specific roles. This is what most organisations currently purchase as the whole programme. It is the necessary beginning, not the complete answer.

The useful thing about that structure is that it can be specified before you know who is going to build it. What follows is a set of requirements rather than a description of any particular offering. Apply them to an internal programme, to a supplier already in place, or to a proposal on the desk.

Routing by function, not by an AI knowledge assessment. Everyone starts with foundations. After that a medical affairs team should be doing medical affairs, a regulatory team regulatory, a supply chain team supply chain. No assessment should be needed to work that out.

A foundational layer that is common because the work is common, not because it is introductory. A great deal is genuinely shared in a regulated business, and it goes well past what AI is and where it fails. GxP validation applies wherever AI touches a qualified process. Why initiatives stall on data rather than models is the same problem in discovery and in commercial. Security, confidentiality and the unsettled question of who owns AI-generated output do not vary by department. Nor do bias, explainability and hallucination management once outputs reach a patient or a regulator, or the judgement required to evaluate work you did not write. That layer should be deep, and it should be built for pharma rather than adapted to it. The test is not how much sits in the foundation. It is whether what sits there is common because the regulatory and technical reality is common, or common because it was cheaper to write once.

A standing route to bring live questions. The gap between understanding a technique and using it on the submission in front of you is where most training quietly fails. The structure should give teams a standing route to bring the questions that arise from that work.

A library kept close to the regulatory and technical reality it describes. Recorded content that ages is the same problem as a workshop that fades, in a format that hides it better. What matters is whether the library is actively maintained, not whether any particular item was updated in a particular month.

Role-specific prompt guidance, kept under review as the models change. Guidance optimised for a previous generation of model can produce materially worse results on the current one, and it fails silently, so someone has to be watching for it.

A per-person record at a granularity the functional sponsor can act on. Not an organisation-wide completion percentage. Who has done what, in which role, and how recently, visible to the person accountable for that team. A record that captures only attendance answers a weaker question than the one being asked of it.

Tool assessment from a source with no stake in the answer. The market for AI tools sold into pharma is growing fast and making large claims, and most published comparisons are sponsored by someone in them.

The requirements, as a checklist:

- Routing by function, with foundations common to all
- A foundational layer common because the work is common, not because it is introductory
- A standing route to bring live questions
- A library kept close to the reality it describes
- Prompt guidance kept under review as models change
- Per-person records at sponsor granularity
- A record that shows more than attendance
- Independent tool assessment

6.2 Why role-specificity is not a refinement but the mechanism

Role-specificity is usually treated as a design question - a refinement that makes training feel more relevant to the people sitting in it. Under Article 4 it is closer to a compliance one: the duty to support the development of AI literacy must still account for people's skills, experience, education and the context in which they use the system - wording that survived the July 2026 rewrite intact. A single company-wide introduction, however well attended, is one measure taken once: thin evidence of effort against a standard framed around context.

The supporting evidence points the same way. Research conducted by Centiment for Docebo across 2,000 enterprise respondents in six countries - 1,000 employees at VP level or below and 1,000 learning leaders - found that 85% of employees said the training they received did not help them use AI in their role.¹¹ The study is not pharma-specific, and Docebo is a learning-platform company with a commercial interest in the finding; both caveats are worth stating. But the direction is consistent with what the transfer-of-training literature has held for decades: whether training changes practice depends heavily on the environment and the specificity of application, not on the quality of the event alone.^{12 13}

The use cases that matter differ sharply by function, and so do the governance questions attached to them. For example:

- **Discovery and R&D** - target identification, generative chemistry, protein design, literature-based hypothesis generation, and the question of what in-silico evidence is sufficient to justify wet-lab spend.
- **Clinical development and operations** - protocol authoring and amendment reduction, site selection, enrolment forecasting, medical coding, statistical programming, and CSR support under a GxP audit trail.
- **Medical affairs** - standard response documents, publication planning, RWE synthesis, MSL insight processing, and the growing question of what AI systems tell clinicians about your products before anyone from your company does.

- **Regulatory affairs** - submission writing, labelling review, response drafting and regulatory intelligence, where the output is read by an authority.
- **Pharmacovigilance and safety** - case processing, signal detection and narrative drafting, where a missed event is not a productivity problem.
- **Quality and manufacturing** - deviation handling, batch record review, validation documentation and predictive maintenance inside a qualified environment.
- **Supply chain** - demand forecasting, serialisation and shortage management across markets with different data availability.
- **Market access and HEOR** - dossier preparation, health economic modelling support and payer evidence synthesis.
- **Market research and insights** - where synthetic respondents and simulated panels are the most oversold technique in the toolkit and the one leadership is most likely to ask about by name.
- **Commercial** - the largest function in most companies and the one most often served by a single overview. Field teams need pre-call synthesis and compliant call documentation; marketing needs modular content architecture that turns promotional review from a bottleneck into a gate; commercial operations needs incentive design, territory modelling and an honest account of what syndicated and claims data can and cannot support; and every brand now needs to know what models say about it and its competitors when a prescriber or patient asks.
- **Corporate functions** - finance, legal, HR and IT, which are adopting AI fastest of all and are frequently outside the scope of any function-specific programme.

The leadership dimension is often underserved. Functional heads who understand how to lead AI-enabled teams, evaluate vendors, set meaningful KPIs and disclose AI use appropriately are not a luxury. They are the governance layer that makes everything else defensible.

6.3 The governance integration layer

In pharma, capability and governance must move together. A programme that builds practitioner skill without maintaining governance awareness recreates the confidence-without-competence risk in a more sophisticated form: the team knows how to use the tools and is not current on what is expected of that use.

Programmes designed for regulated environments therefore run a governance-readiness component alongside skills development - regular updates on regulatory expectations, clear guidance on where AI use must be documented and disclosed, and recurring attention to where rising usage is creating new obligations. This is not an additional burden. It is what makes the skills safe to deploy.

6.4 What this means for your AI champions and AI Task Force Teams

Many organisations have already reached for an internal answer: an AI taskforce, a champions network, a community of practice. Sometimes it spans the company. Just as often it sits inside a single function - a market access team that has organised its own network, a regulatory group that meets monthly to compare notes. If you have one at either scale, you are ahead of most of your peers and you should protect it. The instinct is exactly right: capability needs owners inside the business, and the people who volunteer are usually your best - curious, generous with their time, and trusted by colleagues in a way no external expert can be.

Two questions are worth asking them privately.

First: how much of this are they doing on top of their actual job?

Almost every internal AI network in pharma is a volunteer effort layered onto full-time roles, and the content, the sessions and the keeping-up-with-what-changed all compete with the work they are paid and measured on.

Second: when they run a Q&A, what happens to the hardest questions?

Not because your champions lack ability, but because they are enthusiasts rather than full-time specialists - and some questions need someone who does nothing else. Those questions get parked, guessed at or quietly dropped, and they are reliably the ones carrying the risk.

The conclusion is not that a champions network is the wrong answer. It is that a group of volunteers has been asked to outrun a decay curve in their spare time. A maintained capability programme should sit underneath that network and make its work easier - supplying the content, the expert answers and the governance watch that no side-of-desk team can sustainably produce, while the champions keep ownership of the rollout and the internal relationships that only they can hold.

7. What Learning Leaders Should Measure Instead

If episodic AI training has a single systemic enabler, it is post-session satisfaction as the primary measure of success. Satisfaction scores make the investment look good. They do not indicate whether anything changed.



7.1 What satisfaction scores actually capture

Satisfaction scores have a legitimate function: they surface facilitation problems, content gaps and participant experience issues. A session generating consistently low scores has problems worth addressing. But a session generating fours and fives tells you participants valued the experience - not that capability was retained or applied months later.

What they capture is the peak. That peak is real, and it is exactly what a well-run session is designed to produce. What it cannot tell you is what remains at the ninety-day mark. That requires a different measurement at a different time horizon.

7.2 The metrics that matter

Measuring AI capability properly is harder than running a post-session survey. It requires follow-up at ninety days, structured assessment of actual workflow change, and the discipline to distinguish adoption from attendance.

The questions driving evaluation should be direct. Are teams using AI in their real workflows ninety days after training? Has cycle time on specific tasks improved, and by how much? Can the team demonstrate safe, compliant AI use in an actual work context rather than a workshop exercise? Have AI-related quality events, errors or compliance queries moved in the expected direction?

These are the questions a CFO or chief compliance officer would ask on close inspection. They are also the questions a sustained programme is positioned to answer, because it generates longitudinal data that a single event structurally cannot.

7.3 Repositioning the budget line

For many organisations the barrier is not the argument. It is where the money sits.

AI training is typically booked as discretionary development spend, competing with every other one-off initiative for a fixed budget. That classification does more damage than it appears to. It sets the expectation that the spend is occasional, it puts the decision on an annual cycle, and it forces the item to compete against things that genuinely are one-offs.

Two moves change it, and which one applies depends on the size of the organisation.

Where a learning and development function exists, the move is reclassification: shifting AI capability out of the episodic expense category and into the standing capability investment line - alongside pharmacovigilance competence, quality systems training and clinical data management expertise. Those are not one-off investments. They are the recognised ongoing cost of maintaining a capability the organisation depends on. AI capability belongs in the same category, with the same budget treatment and the same expectation of continuous maintenance. This requires no additional budget, only a different classification of budget that already exists.

Where no learning and development function exists, there is no line to reclassify, and the move is the opposite: create one, at the point AI enters the work rather than after something goes wrong. In a fifty-person biotech the practical question is not how AI capability is categorised - nothing is categorised - but who owns it and whether it survives the next quarter. The answer that holds is the one already applied to the quality management system or the pharmacovigilance contract: a named owner, a standing cost, and a renewal date, rather than a training decision revisited whenever someone remembers to raise it.

Smaller organisations usually reach this point later than large ones. They reach it with fewer people standing between the AI use and the regulated output.

The Act anticipates this. The rewritten Article 4 adds a second paragraph tasking the Commission and Member States with supporting and facilitating the efforts of providers and deployers, with particular regard to **small and medium-sized enterprises** - a legislator's acknowledgement that smaller organisations are squarely in scope and will need help getting there.

AI literacy is now a regulatory obligation, not a development ambition - and notably, it is an obligation that was not deferred when the Act's high-risk deadlines moved back. It should be budgeted where the organisation budgets its other compliance costs, from the outset, rather than competing for discretionary funds it will lose to something more urgent every quarter.

The distinction matters more than it sounds. In most organisations compliance spend and discretionary development spend are governed by different rules, approved by different people, and defended differently when budgets tighten. A capability that is genuinely required is rarely the first thing cut. A capability that is merely valuable frequently is.

This is, in practice, the single most useful structural change available to a reader of this paper.

In a large organisation it is a reclassification of budget that already exists.

In a smaller one it is a decision, taken early, to fund the obligation rather than the aspiration.

8. AI Capability as a Strategic Differentiator

The case so far has been made largely in terms of risk: compliance exposure, audit vulnerability, the cost of decay. Those are compelling. They are not the only reasons.

Pharmaceutical and biotech functions - discovery and R&D, clinical, regulatory, medical, supply chain, market access, commercial, finance, legal, HR - are in a period of genuine differentiation based on AI capability. Functions that build real, maintained, role-relevant competence now are not merely reducing risk. They are building an operational advantage that compounds.

Technology transitions across pharma's history - from electronic data capture in clinical trials to evidence-based reimbursement frameworks in market access - have shown the same pattern: organisations that built and maintained the underlying competence, not just the tools, captured advantages that took late movers years to close. AI is no different. The differentiator is not budget. It is structure.

For most functional leaders, AI currently sits on the leadership agenda as a risk to be managed. That is an accurate description of the present, not a permanent condition. The shift from "AI is a risk we are managing" to "AI is a capability our function is known for" is not primarily a technology shift. It is a design shift - and teams that make it answer questions about AI use with confidence rather than defensiveness.

9. Four Practical Moves

The argument resolves into four moves, each within reach of functional leaders, L&D heads and compliance officers, and none requiring board approval or significant new budget.

9.1 Audit what is already happening

There are two audits here, and most organisations only think to run the first.

The first is of what has been purchased. How many AI training sessions has the function attended in the last twenty-four months? What follow-up structure existed after each? What was measured, and at what horizon? What proportion of the content delivered is still current, given how the tools and the regulatory landscape have moved?

That audit usually reveals the same pattern: repeated investment against the same gap, with satisfaction scores as the primary evidence of value. Naming the pattern is normally enough to make the case for a different approach internally.

The second audit is of what people are already doing without being asked to - and it is the one worth running even if the first returns nothing. Which tools are in daily use, including the ones nobody approved. Which workflows they have quietly entered. What proportion of that use touches regulated, confidential or patient-related work. Whether anyone could currently answer, for a given piece of AI-assisted output, how it was produced and what was checked.

For an organisation that has never bought AI training, that second audit is the whole exercise. For one that has bought five workshops, it is still the more revealing of the two, because purchase history describes what the organisation intended and usage describes what it actually has.

How the question is asked determines the answer. Run as a compliance investigation, it produces denial and drives use further underground. Run as a capability question - what are you using, what is it good for, where do you not trust it - it produces a remarkably candid picture, because most people are not hiding anything. They are using tools that were already on their desktop for work that was already on their list.

This is also what the AI literacy obligation now requires in practice. Literacy appropriate to the context in which a person uses an AI system cannot be scoped without knowing which contexts those are. An organisation that cannot describe its own AI usage cannot demonstrate that its training addresses it.

9.2 Ask the two questions

When evaluating providers, make the structure of post-training support the primary criterion rather than the quality of the opening session. Ask specifically: what happens in month four? What keeps the team current as the tools change? How does the training differ by context and role? How is the governance layer maintained alongside the skills? How does the programme drive behaviour change rather than awareness?

A provider who answers by describing another session is offering the expensive model in economical packaging. The answer to look for describes a structure: a refresh cadence, applied working sessions, an evolving library organised by role, and a governance checkpoint built into the rhythm.

The second question is the one a maintained generic programme cannot answer, and it should be asked in exactly these terms: what does this look like for a regulatory writer, and how is it different for a brand manager?

A provider with genuine role-specificity answers immediately, and differently, for each. A provider with a generic catalogue will describe the same content twice and call the difference a learning path. Ask for the two curricula side by side. The gap between them is the whole answer, and it takes about a minute to see.

9.3 Change what you measure and report

Retire satisfaction scores as the headline metric - or at minimum supplement them with ninety-day capability measures: workflow adoption, efficiency gains on specific tasks, demonstrated safe use in regulated contexts, governance readiness. These are harder to gather and they are the only ones that provide a defensible evidence base for continued investment. The internal conversation changes when the metric changes. An executive shown post-session satisfaction scores will reasonably ask whether the investment was worthwhile. One shown ninety-day adoption data, time savings and governance evidence has a different conversation to evaluate.

9.4 Reposition AI capability as a standing investment

Work with finance and procurement to move AI capability out of the episodic L&D budget and into the standing capability category. The internal case rests on three components: the hidden cost analysis, the three-year comparison run against your own numbers, and the risk argument - that confidence without competence in regulated workflows is a liability rather than a missed opportunity.

That case is available to make now, and the functions that make it this year will spend the next two operating from a materially stronger base than those that continue to buy the moment.

Conclusion: The Most Expensive Training Ends When the Room Empties

There are two ways to overspend on AI training in pharma, and neither is obvious at the point of purchase.

The first is the repurchase cycle: buying the same introductory session year after year, watching capability erode after each one, and funding a recurring lift that produces no cumulative asset.

The second is subtler. It is the session that works - that genuinely raises usage and confidence - without anything in place to sustain the competence underneath. That session gets the glowing scores and is, in the moment, money well spent. Ninety days later, usage is still elevated and the judgement supporting it has quietly eroded. In a regulated workflow, that is the more expensive of the two outcomes, and the satisfaction data provides reassurance precisely where reassurance is least warranted.

Both share a root: training treated as an event that ends when the room empties, rather than a capability that is built, applied, maintained and governed over time.

The fix is not more spending. It is a different shape of spending - designed around how AI knowledge actually behaves, how regulated environments actually operate, and how capability is actually sustained. That means a regular refresh cadence, applied working sessions, always-on resources, integrated governance, and measurement that looks ninety days out rather than five minutes after the session ends.

Trained once is not a capability plan. It is the illusion of one.

Pharma already knows how to build capability that lasts. It does it for pharmacovigilance. For regulatory affairs. For clinical operations. For every function where the cost of getting it wrong is consequential. The logic is identical for AI. The only thing missing is the decision to apply it.

PART TWO

How Eularis Applies This: The Pharma AI Enablement Institute

Part One set out the problem and the evidence. This part describes how Eularis has built a programme around it. It is included so that readers who want the applied answer have it in one place; the argument above stands independently of it. Nothing in this part is legal advice, and nothing in it should be read as a representation that enrolment satisfies any regulatory obligation - see the regulatory note at the end of section 3.3.

What a year of maintained capability can look like

The picture below is an illustration rather than a commitment. It describes the direction a function tends to move in when capability is maintained rather than repurchased. How far it moves, and how quickly, depends on the team, on the workflows it chooses to work on, and on what it puts in.

Over time, AI can stop being the agenda item that follows the risk register and start becoming a capability the function is known for.

A team that has been asking whether it is safe to use a given AI tool for a given task may find that question settling, its use of the tools becoming both more fluent and more defensible and AI output is judged effectively and compliantly. Which work stops bottlenecking depends entirely on the team. For one it might be the submission section that consumed a week of drafting; for another the literature synthesis that swallowed three days of a scientist's time, or the evidence pack that sat half-assembled, or the competitive intelligence brief that never quite arrived by Friday. Which of those moves, and by how much, depends on what the team brings and what it chooses to work on.

The friction around AI-assisted work can settle - the uncertainty about what standard applies, the back-and-forth before anyone will put their name to an output, the quiet anxiety about whether the team's use of these tools would survive scrutiny. For a functional leader in a regulated workflow, that might show up as review processes that run with less rework. For an L&D leader, it can mean having more than anecdote to draw on when someone asks what has actually changed in how people work.

And the capability question can move further down your personal worry list - not because you have stopped watching it, but because there is a structure alongside you that is watching it too. Surfacing new model capabilities and shifts in regulatory expectation is what the sessions are for. When someone joins, they onboard into a live programme rather than a slide deck from last spring.

What enrolment includes

The quarterly refresh

Four times a year, the first live session of the quarter looks at what has changed across the field since the last one. What is covered depends on what has actually moved in the period; the kind of ground it can cover includes new model capabilities, tools entering the market, shifts in the regulatory landscape, and emerging practice in how pharma teams are

applying AI. These sessions run once across the Institute rather than separately for each function, and they are not repeat performances of introductory material.

The intent is that attending one saves a team some of the work of scanning the field themselves, distilled and contextualised for the industry. Sessions are led by Dr Andrée Bates and, where the subject calls for it, by specialist contributors.

Monthly live office hours

Once a month there is a live session built from questions submitted in advance and answered anonymously: a regulatory team working out how to use AI appropriately in a submission workflow; an L&D manager building the internal case for sustained investment; a commercial leader whose team has adopted a tool but is not using it well; a market access team unsure whether a vendor's claims will hold up. Anonymity is what allows the session to be shared across organisations without anyone having to describe their internal workflows in front of peers.

These are not lectures. Sessions are built around the problems teams actually bring, and they are the standing route for the live questions and challenges described in section 6.1 - the applied, problem-specific practice most conspicuously absent from event-based models.

Every licence begins with Foundations. A business unit takes the foundational tier before its own track opens, so that the function-specific work starts from a common baseline rather than rebuilding one. Private organisational-specific office hours are also available.

A recorded library, organised by function

Between live sessions, enrolled teams have unlimited access to the library, with new content added every month. It is organised by foundational, function and Leadership: a regulatory affairs professional does not search through content built for commercial teams. Each covered business unit has a track built around the work that unit actually performs.

Role-specific prompt guides, updated regularly

A continuously updated library of tested prompt guides built around the actual tasks each function performs - showing a regulatory writer how to approach submission drafting, label review or regulatory intelligence synthesis; a clinical operations lead how to approach protocol support, literature review and CSR assistance; a market access director how to structure dossier preparation and payer evidence synthesis. Kept under review as the underlying models change in ways that affect how they should be used.

Underneath all of it sits a delivery layer that produces the evidence described in section 3.4: per-person lesson completion and progress, visible to the functional sponsor for their own team rather than only in aggregate. Assessment and certification are part of the same record as the platform develops. This is not a feature list. It is the difference between believing your team is trained and being able to show it.

Commercial terms

- One price per business unit, banded by the size of the unit. No per-seat charges
- Everyone on the licence receives the full business-unit library, the complete Foundations curriculum, all content added during the term, and the live sessions
- Sponsor visibility of usage and completion
- PO-based onboarding and procurement documentation
- Annual term. No mid-term cancellation and no automatic renewal

- Available as add-ons: portal theming to your own branding, private sessions for your teams alone, and delivery into your own learning management system. Each is quoted on request. LMS integration is quoted per system connected, because every connection is built for that specific environment

Also on the roadmap

The following is planned but not yet part of what enrolment includes at launch. It is described here so that the direction of the programme is clear.

Independent vendor evaluations

Independent evaluations are planned, covering the tools in use in the functions enrolled: what those tools actually do well, where they fall short, what the regulatory and compliance considerations are, and whether the claims in the pitch match the reality in the product, written for functional leaders in the language of the work.

When they arrive, Eularis will take no referral fees, revenue share or promotional arrangements from any vendor it evaluates, and no vendor pays to be included in an evaluation or left out of one. That is not how most analyst coverage is funded, which is worth knowing before relying on any comparison. That independence is not incidental; it is the foundation on which the value of these evaluations rests.

About Eularis

Eularis has been working inside pharma on AI since 2003 - twenty-three years, making it one of the longest-established AI consultancies working exclusively in the pharmaceutical and biotech sector, with no platform to sell and no incentive other than results. Dr Andrée Bates has delivered AI solutions and training for the majority of the world's largest pharmaceutical companies, and more than forty senior pharma executives have publicly endorsed her training.

The Institute exists because the gap between the future described above and where most functions sit today is not a talent gap or a budget gap. It is a structure gap. The people are capable. The tools exist. What has been missing is a programme designed around how AI knowledge actually behaves and what a regulated environment actually requires.

What participants said

These are public LinkedIn endorsements from participants in Eularis training programmes. More are in the [recommendations section of my profile](#).

Having spent Section 7 arguing that satisfaction proves nothing, I should be consistent about what these are and are not. Some describe what changed in the work. Others describe the session itself, and those are an account of the teaching rather than evidence that capability held. That distinction is the entire reason the Institute exists.

The training was engaging, dynamic and highly practical. Rather than focusing only on AI theory, it provided real-world applications and actionable techniques that can be immediately applied in daily work. One of the most valuable aspects for me was the advanced prompting section, where I learned how to build and use AI agents tailored to my own Regulatory Affairs activities. These approaches have already helped me work more efficiently and think differently about how AI can support decision-making in a regulated environment.

Jovana Stankić · Regulatory Affairs Professional, Dechra Pharma

Dr Bates conducted a workshop that was an incredible hit. It provided a comprehensive overview of AI applications in pharma and biotech and left audience members with a clear understanding of how AI can drive efficiencies and greater outcomes across all areas of the industry, including marketing, sales, forecasting, pricing and regulatory. The workshop was fun, engaging, interactive, and left me with many actionable takeaways.

Raca Banerjee · Associate Director, Marketing, Novartis

Her presentation covered a spectrum of AI trends and practical use cases, and demonstrated AI's impact on enhancing R&D productivity in the pharmaceutical industry. Andrée seamlessly navigated through examples showcasing practical application across the drug development lifecycle, and shed light on emerging AI vendors worth considering.

Sotirios Perdikeas · Leader, Predictive Modeling & Data Analytics, Research and Early Development, Roche

Dr Andrée Bates delivered a most astounding and insightful lecture on AI for our Board Directors Programme. Her research is current and her capacity to make a most complex subject easy to absorb - especially for busy senior professionals - is enviable. She held the audience on the edge of their seats and stimulated the greatest degree of discussion of all presenters.

Andrew Kakabadse · Professor of Governance and Leadership, Henley Business School

Where to go next

If the structure described above is one you want for your function, the next step is a twenty-minute call with Dr Bates - not a pitch, but a look at where your team currently sits, what would need to be true for this to work in your organisation, and whether it is the right fit at all. The programme is then shaped around your function workflows rather than a generic AI syllabus.

See the programme and book a call » eularis.com/institute

Prefer to start by email? Reply to the message this paper arrived in, or write to training@eularis.com, and we will answer directly.

Eularis | training@eularis.com | eularis.com/institute

References

1. Tatel, C. E., & Ackerman, P. L. (2025). Procedural skill retention and decay: A meta-analytic review. *Psychological Bulletin*, 151(6), 696–736. <https://doi.org/10.1037/bul0000481>
2. Arthur, W., Bennett, W., Stanush, P. L., & McNelly, T. L. (1998). Factors that influence skill decay and retention: A quantitative review and analysis. *Human Performance*, 11(1), 57–101. https://doi.org/10.1207/s15327043hup1101_3
3. Regulation (EU) 2024/1689 (the EU AI Act), Article 4 (AI literacy), as amended by Regulation (EU) 2026/1744. Applicable from 2 February 2025; national supervision and enforcement from August 2026. See European Commission, AI literacy questions and answers, digital-strategy.ec.europa.eu.
4. Regulation (EU) 2026/1744 (the Digital Omnibus on AI), amending Regulation (EU) 2024/1689. Published in the Official Journal 24 July 2026; in force from 27 July 2026. Article 1, point 5 replaces Article 4 in full.
5. Regulation (EU) 2024/1689, Article 14 (human oversight of high-risk AI systems). Applicability of standalone Annex III high-risk obligations deferred to 2 December 2027, and Annex I embedded systems to 2 August 2028, by Regulation (EU) 2026/1744.
6. Regulation (EU) 2024/1689, Article 50 (transparency obligations). Applicable from 2 August 2026. Transitional arrangements apply to synthetic-content systems placed on the market before that date.
7. U.S. Food and Drug Administration, Considerations for the Use of Artificial Intelligence To Support Regulatory Decision-Making for Drug and Biological Products. Draft Guidance for Industry and Other Interested Parties, January 2025, docket FDA-2024-D-4689. Not for implementation; contains non-binding recommendations. Still in draft as at August 2026; no final guidance issued.
8. U.S. Food and Drug Administration and European Medicines Agency, joint paper ‘Guiding Principles of Good AI Practice in Drug Development’, January 2026.
9. European Medicines Agency, reflection paper on the use of artificial intelligence in the medicinal product lifecycle.
10. Cherif, M. (2026, July). Five-stage account of individual AI adoption in medical affairs, shared with the author and reproduced with permission. Related discussion published on LinkedIn: https://www.linkedin.com/posts/myriam-cherif-medical-affairs_i-used-to-worry-ai-would-get-things-wrong-activity-7483098263005499393-g6Uc/
11. Docebo (2026). The AI Readiness Gap: The 2026 Enterprise Learning Wake Up Call. Research conducted by Centiment across 2,000 enterprise respondents in the US, UK, Canada, France, Germany and Italy, comprising 1,000 employees at VP level or below and 1,000 learning leaders.
12. Baldwin, T. T., & Ford, J. K. (1988). Transfer of training: A review and directions for future research. *Personnel Psychology*, 41(1), 63–105.
13. Jackson, C. B., Brabson, L. A., Quetsch, L. B., & Herschell, A. D. (2019). Training transfer: A systematic review of the impact of inner setting factors. *Advances in Health Sciences Education*, 24(1), 167–183.



About the author: **Dr. Andrée Bates**

Dr. Andrée Bates is the founder and CEO of Eularis, and one of the world's foremost authorities on AI strategy in pharmaceutical and life sciences — having pioneered AI applications in the industry back in 2003, two decades before it became mainstream.

She has led hundreds of AI implementations across the pharmaceutical value chain for companies including Pfizer, Novartis, Roche, GSK, Merck, and AstraZeneca, delivering outcomes that include 1,000% efficiency gains, billions in additional revenue, and products that achieved blockbuster status in their first year. Her work spans drug discovery, clinical trials, regulatory affairs, market access, and commercial operations across six continents.

She holds a unique position as one of the few experts who can bridge the gap between data science teams and business leadership — the disconnect responsible for 87% of AI project failures. Dr. Bates serves as guest lecturer at INSEAD, Fordham University's Gabelli School of Business, Henley Business School, the University of Rhode Island, and St. Joseph's, and is an assessor for Innovate UK, evaluating AI-focused startup grants. She also hosts the podcast AI for Pharma Growth and has been recognised in Leading Women to Watch 2025 and the Women in AI Initiative.

Check out the Podcast for all the latest insights:

