

Whitepaper

The age of agentic

A complete guide on agents in clinical trials

Table of contents

Executive summary..... 3

Section 1 | We’ve reached the limits of human-only clinical development..... 4

Section 2 | From generative to agentic, what’s actually new..... 5

Section 3 | The infrastructure beneath the agent 5

Section 4 | The human-in-the-loop governance model that gives you control 7

Section 5 | Build versus buy in fielding a clinical-grade agent..... 8

Section 6 | Why clinical expertise is the deciding factor..... 9

Section 7 | Early evidence from the field..... 10

Section 8 | Overcoming the two myths that stall adoption..... 11

Section 9 | A framework for adoption..... 12

Section 10 | Validation and governance 13

Conclusion..... 16



Executive summary

Clinical development is entering a new operating era. For three decades, the industry has improved clinical trial execution by redistributing human effort, moving from source data verification to risk-based monitoring, and from paper to electronic systems.

Agentic AI is the technology that finally changes that premise. Rather than a tool that responds to a prompt, like generative AI, an agent is a tool that can perceive context across multiple data sources, reason about what it finds, and execute goals that involve multi-step tasks. They can also adapt as new information arrives. Applied to clinical development, agentic systems can continuously monitor trial data, aggregate information across disconnected systems, draft and route documentation, and flag risk well before a human can.

This whitepaper presents that 2026 marks the arrival of the ‘age of agentic’ in clinical development, a period in which agent-based systems, purpose-built ontologies, and standardized connectors converge to eliminate the manual, sequential, and siloed processes that have quietly added months and years to drug development timelines. It examines what agentic AI is and is not, the infrastructure that makes it trustworthy in a regulated environment, why human oversight remains a must, and the build-versus-buy and vendor-selection decisions that determine whether an organization captures this advantage early or falls behind.

Across the sections that follow, three themes recur.

1. Autonomy without infrastructure is fragile.

Agents are only as reliable as the ontology and connector layer beneath them, the semantic and technical plumbing that lets an agent understand and safely act on data it did not create.

2. Autonomy without oversight is unsafe.

The tasks best suited to full AI autonomy are administrative and analytical; the tasks that touch patient safety, regulatory submission, and database lock still require a human decision-maker.

3. Autonomy without clinical grounding is theater.

A model that can reason is not the same as a system that understands a protocol amendment, a TMF filing rule, or the difference between a Phase II and a Phase III trial. That fluency is what separates a proof of concept from a production deployment.

We've reached the limits of human-only clinical development

The complexity of clinical trials has been rising for years, and every measure of that complexity points in the same direction. Research from Tufts CSDD has documented a continuing upward trend across nearly every protocol design variable, including more endpoints, more eligibility criteria, more pages, more investigative sites, and more data points collected per study, a trend most pronounced in Phase II and III trials. IQVIA's analysis of drug development timelines found that new drugs spend an average of 45% of their development time in "white space" (the unproductive gap between one trial phase ending and the next one starting) on the way to regulatory submission.

A second, less-studied form of white space exists inside individual trials themselves: the idle time between protocol finalization and site activation, between first activation and last patient last visit, and between last patient last visit and database lock. Anyone who has run a trial recognizes the pattern: waiting on IRB and site-committee approvals, chasing incomplete feasibility feedback, reconciling data across systems at IDMC review points, resolving queries one by one after the last patient visit concludes.

100.4

Months

Average clinical trial cycle length in 2024, up from 93.4 months in 2020

45%

Share of drug development time spent in white space between trial phases (IQVIA)

14

Months

Growth in inter-phase white space over five years, offsetting gains in trial speed

Data drawn from Statista trial cycle time figures (2020–2024) and IQVIA Pipeline Intelligence and Institute analysis (2024).

These delays are not the product of any single broken process. They stem from three systemic issues that compound across the trial lifecycle. Manual processes create dependency on human availability, since someone has to track a regulatory submission, chase a site for a document, or resolve a query before the next step can begin. Siloed systems fragment information across disconnected platforms, so regulatory, operations, and safety teams frequently work from different pictures of the same trial. Additionally, sequential workflows force activities that could run in parallel to instead wait in a queue. This means protocol amendments cascade through eClinical builds, regulatory approvals must finish before training starts, and database lock cannot proceed until every query and every central lab result is in hand.

The result is a paradox the industry has lived with for years. While process and technology improvements have shaved roughly seven months off actual trial durations over the past five years, the white space between trials grew by fourteen months over the same period, essentially erasing the gain. The pattern points to a structural conclusion: the industry has scaled human effort about as far as it can go, and the next gains have to come from shifting the burden of coordination itself onto technology built to carry it.

“

The industry has scaled human effort about as far as it can go, and the next gains have to come from shifting the burden of coordination itself onto technology.

”

Section 2

From generative to agentic, what's actually new

“AI” has become a catch-all term, which makes precision worth the effort. Artificial intelligence is the broad science of building intelligent machines; machine learning is the subfield in which systems learn from data and identify patterns with minimal human intervention.

Generative AI, in the context of clinical trial operations, refers to AI models (usually large language models or multimodal models) that can create new content such as text, data structures, translations, summaries, and even code, rather than just classifying or predicting from existing data.

Agentic AI is a different category of system. Rather than simply responding, an agent can autonomously break a goal into sub-tasks, plan a sequence of actions, execute them using external tools and data sources, and adapt as new information arrives, with minimal human intervention at each step. A useful, if informal, comparison is the difference between asking for driving directions and handing the keys to a driver who can also stop for gas, reroute around traffic, and adjust the plan when a bridge is closed. In a clinical trial, that might look like an agent that pulls enrollment figures from a CTMS, cross-references them against safety data in an EDC, flags a site falling behind, drafts a follow-up email to the coordinator, and updates a dashboard, all without a human manually opening each system in turn.

Section 3

The infrastructure beneath the agent

(ontology, connectors, and MCP)

An agent is only as trustworthy as the layer of infrastructure beneath it. Three components make agentic AI viable in a regulated clinical environment, a semantic layer that lets disparate systems mean the same thing, a connector layer that lets agents reach live data without displacing systems of record, and a governance layer that keeps every action auditable. This section addresses the first two; human governance is addressed in Section 4.

The ontology problem

(why systems don't speak the same language)

Clinical data is fragmented not just across systems but across vocabularies. One EDC records “myocardial infarction”; another abbreviates it “MI”; a third simply logs “heart attack.” A “patient” in one system is a “subject” in another, and a “site” is structured differently depending on which CRO built the database. Multiply this across the more than four data sources the average sponsor now relies on, and the scale of the translation problem becomes clear, since industry surveys report that the large majority of organizations experience active friction from these mismatches. Traditionally, harmonizing data in order to take advantage of AI tools would take a large, resource and time-intensive effort that's often measured in years. Now, this is no more.

Clinical and non-clinical AI are two different bets

The clinical trial technology landscape is being reshaped along two dimensions: the pace of AI innovation and the depth of clinical expertise behind it. Clinical incumbents bring deep domain knowledge, but many are constrained by legacy architectures and slower AI development. At the other end of the spectrum, horizontal or general AI platforms are advancing rapidly, but lack the clinical context, workflows, and regulatory understanding required to meaningfully transform trial execution.

The industry increasingly needs both, a fusion of AI innovation that's grounded in deep clinical expertise. This is where Medable is positioned, bringing clinical depth and modern AI capabilities together in a platform purpose-built for clinical development. This enables Medable's clinical trial platform to move beyond general-purpose productivity and into the complex decisions, workflows, and processes that drive trials.

As AI becomes more embedded in trial operations, the strongest position belongs to platforms that can innovate at the pace of AI while grounding every capability in clinical context, trusted data, and real-world trial workflows.

An ontology solves this by acting as a semantic translator, a formal, machine-readable specification of a domain's concepts and the relationships between them, built from standard identifiers, shared vocabulary, precise definitions, and logical rules that let software reason over the data rather than just store it. The life sciences industry has decades of experience building these systems, from the gene ontology that unified genomics research in the late 1990s to today's SNOMED CT, MedDRA, CDISC/SDTM, LOINC, and NCI Thesaurus standards, which respectively anchor clinical concepts, adverse event terminology, trial data tabulation, lab observations, and oncology and regulatory nomenclature.

Historically, integrating data across these vocabularies meant a manual mapping exercise, with data engineers reverse-engineering proprietary schemas, negotiating field-by-field translations, and producing brittle pipelines that broke with every vendor update. That work could consume six to twelve months before a single AI model ever saw a usable dataset. A clinical ontology layer changes the economics of that work by creating a semantic mapping (grounded in standards such as CDISC, CDASH, and the Pistoia Alliance's clinical development ontology) that lets an agent traverse multiple systems as though they already shared a common vocabulary.

The data stays where it lives, rather than requiring data to be moved or cleaned before an agent can use it, and the understanding travels with the agent.

In short, what used to take a team of data engineers years to normalize can now happen in weeks, and in some cases days, because the mapping itself can be proposed agentially, with a confidence score attached so a human reviewer can accept, refine, or reject it.

How connectors and MCP give agents somewhere to act

A model with no access to live systems can only reason about data it is shown. The Model Context Protocol (MCP) and purpose-built connectors solve the access problem, giving an agent permissioned, live visibility into the systems that actually run a trial (EDC platforms, CTMS, eTMF, IRT, safety and pharmacovigilance databases, feasibility tools, collaboration platforms, data warehouses, and identity systems) without requiring any of those systems to be rebuilt, replaced, or exported.

A mature clinical connector layer typically spans five categories of systems.

Layer	Representative systems	What agents gain
Clinical systems of record	EDC (Veeva, Medidata Rave, Oracle InForm), CTMS, eTMF, IRT, safety databases, feasibility tools	Live, permissioned access to the data that defines trial execution
Collaboration & productivity	Email, chat and video platforms, cloud storage, project and documentation tools	The ability to draft, coordinate, and notify inside existing workflows
Data & intelligence	Semantic/ontology layers, cloud data warehouses, search and code utilities	Continuous analysis grounded in the full data landscape, not a single silo
Infrastructure & identity	DevOps tooling, identity and access management, platform configuration	Auditable, access-controlled actions aligned to enterprise security policy
Global trial support	e-signature and consent workflows, translation and localization services	Consistent, compliant execution across languages and regions

The practical effect is a shift from isolated tools that sit on top of a workflow and merely summarize or answer, to systems that participate directly in execution, retrieving data, updating records, and maintaining context as they move between systems. An agent can pull enrollment data from an EDC, compare it against safety signals, generate a document, notify the right stakeholder, and update a dashboard, all inside one continuous, auditable sequence that mirrors how a trial actually runs.

Together, the ontology and connector layers create a compounding effect rather than three independent improvements. Agents accelerate the discovery of relationships across data; connectors make live integration possible without re-architecting existing systems; and the ontology makes each semantic decision durable and reusable, rather than a one-off judgment call that has to be remade every time. None of the three is transformative alone. Combined, they collapse integration timelines that used to stretch across quarters into weeks, and sometimes days.

Section 4

The human-in-the-loop governance model that gives you control

Agentic systems are powerful, but they are not inherently suited to full autonomy in every regulated context, and the organizations seeing durable results are not the ones racing to remove humans from the loop. Current-generation models still require active human supervision to meet the standard regulators expect, and the consequences of getting this wrong are severe, including product delays, financial penalties, and in the worst case, harm to patients.

The practical answer is not to slow adoption but to be precise about where autonomy belongs. Two categories of work sit on either side of that line.

Where AI autonomy is appropriate today

Repeatable, rules-driven, high-volume work is where agents deliver the most immediate value, particularly given that trial data is often scattered across fifteen or more disparate systems, an aggregation burden that is exhausting for people and trivial for a connected agent.

- Administrative and follow-up work, such as chasing outstanding site documents, tracking data-entry and milestone timelines, monitoring protocol deviations, managing routine correspondence, generating standard reports, and coordinating scheduling.
- Data processing and analysis, including initial data cleaning and validation, preliminary statistical summaries, draft enrollment reports, routine database queries, and standardized interim safety summaries.

Where human judgment remains essential

Mission-critical activities that carry direct regulatory or patient-safety consequences require a qualified human decision-maker, not because AI cannot assist, but because the final call carries accountability an algorithm cannot hold.

- SAE reporting, where clinical interpretation of a safety signal determines what gets communicated to regulators and how quickly.
- Regulatory submissions (IND and CTD sections and other filings) where completeness and consistency with evolving guidance require expert review.
- Protocol amendments, which must be assessed by qualified medical and regulatory professionals for clinical and regulatory impact.
- EDC query creation, medical coding, and final database lock, where clinical understanding determines which discrepancies matter and which do not.
- Audit response, risk assessment, and CAPA planning, all of which depend on organizational context and judgment that sits outside the data itself.

Effective programs formalize this split through governance rather than leaving it to individual judgment call by call, through clear decision boundaries, risk-proportionate approval hierarchies (senior review for higher-stakes actions, lighter-touch review for routine ones), defined escalation paths, and ongoing training so human reviewers stay sharp evaluators of AI output rather than passive approvers. This is not a brake on adoption, but the mechanism that lets autonomy expand responsibly. As teams build a track record with a given class of task, the human-in-the-loop model becomes a structured way to widen the scope of what agents are trusted to do next, informed by evidence rather than assumption.

Build versus buy in fielding a clinical-grade agent



Big corporations can't rely on their internal speed to match the transformation that is happening in the world. As soon as I know a competitor has decided to build something itself, I know it has lost.

CEO of Sanofi



These candid sentences from Sanofi's CEO, showcase one of the most common questions that's at the forefront of every pharmaceutical company's mind; whether to build or buy your way into the agentic and generative AI revolutions.

The instinct to build in-house is understandable, since a capable foundation model plus a working prototype can feel like most of the job is already done. It rarely is.

External research on enterprise AI adoption has been unusually consistent on this point. One widely cited MIT study found that only a small fraction of custom generative-AI tools survive the transition from pilot to production, while generic off-the-shelf assistants perform well on trivial tasks but stall the moment a workflow demands real context and customization. A separate IDC review of enterprise AI proofs of concept found that fewer than one in eight made it to wide-scale deployment. The gap between an impressive demo and a validated, GxP-compliant, audit-ready production system is where most internal programs lose momentum and money.

Seven recurring obstacles explain why a do-it-yourself clinical agent stack is harder than it looks. This is because:

- 01** Model limitations, since even capable models can hallucinate or obscure their reasoning, and any unexplainable output undermines trust in a regulated setting.
- 02** Integration fragility, as agents depend on external systems whose APIs, latency, and permissioning were never designed with an autonomous consumer in mind.
- 03** Coordination complexity, since shared context and synchronized state become their own engineering problem as soon as more than one agent is involved.
- 04** High-stakes autonomy, where even small automation errors can touch patient safety, raising the bar for auditability to match a human expert's standard.
- 05** Data precision, since trial data is fragmented and inconsistent by nature, so every retrieval has to be validated against a trusted source, not a best guess.
- 06** Security and compliance load, as GxP, HIPAA, and 21 CFR Part 11 obligations turn what looks like a feature into its own multi-year program.
- 07** Long-tail integration cost, since connectors to CTMS, eTMF, IRT, eCOA, EDC, safety, and document systems each need to be built, validated, and maintained through every vendor version change.

A do-it-yourself build typically requires more than two years and several specialized teams (product, engineering, DevOps, QA, data science, security, and clinical operations) before a single validated agent goes live, with every additional agent type and connector adding further months. Partnering with a platform built specifically for clinical development compresses that timeline substantially, because the compliance architecture, connector library, and clinical grounding already exist rather than needing to be assembled from scratch.

The decision is not binary in practice. Many organizations adopt a hybrid approach, deploying a ready-to-use agent to capture early value quickly, then building custom agents for the areas of greatest differentiation, all on a modular, model-agnostic architecture that avoids re-building the same regulated infrastructure study after study. Whichever path an organization takes, the operative question is not "build or buy," but "how fast can we scale responsibly," because in a regulated field, speed without validation is not actually speed.

Why clinical expertise is the deciding factor

Sponsors and CROs evaluating agentic AI today are effectively choosing between two imperfect defaults. One option is a fast-moving, AI-native company with powerful models but no lived experience running a trial, no fluency in CDISC, CDASH, ICH E6(R3), GxP, protocol deviation workflows, TMF triage, or CRA closeout, a body of domain knowledge that takes years to build and cannot be patched in with a model update. The other is an enterprise-first incumbent with deep clinical roots but a slower pace of AI innovation, often bolting agents onto legacy architecture in ways that risk vendor lock-in and limited interoperability.

Internal market research with sponsors, CROs, and study teams surfaces this gap directly, a majority of respondents report that AI and technology vendors do not speak the language of clinical trials, leaving the sponsor to fill in the gaps themselves at real cost in time, risk, and trust. One buyer described a scientifically credentialed health-tech vendor who nonetheless needed to be taught the difference between a Phase II and a Phase III trial before becoming useful. That is not a knowledge gap a more finely tuned model closes on its own.

Three capabilities separate a system built at the intersection of clinical context and AI-native architecture from either default option alone.

- 01** A clinical semantic layer, with CDISC, CDASH, ICH, and GCP ontologies built in from the start, which becomes exponentially more valuable once agents are reasoning across a trial rather than just displaying data.
- 02** GxP by design, not by retrofit, with an audit trail on every agent action, role-based controls, and a validation framework that is native to the architecture rather than compliance work added after the fact.
- 03** Model-agnostic orchestration, an architecture that lets an organization swap the underlying model as the field advances, so the value of the platform compounds independently of which foundation model happens to be winning at any given moment, and sponsors are never locked into a single vendor's roadmap.

In practice, this shows up as concrete proof points that a generic AI vendor cannot yet claim, including day-one data readiness with no multi-year harmonization project, inspection-ready governance rather than governance retrofitted after a near-miss, production deployment measured in weeks rather than another pilot, out-of-the-box agents for clinical monitoring, TMF management, digital data flow, and anomaly detection, and the ability to work alongside a sponsor's existing systems rather than requiring them to be replaced. That last point matters more than it may first appear, since disruption cost is one of the most common objections to any new clinical technology, and a platform that fits into the existing ecosystem removes that objection before it is raised.

“

Vendors do not speak the language of clinical trials, leaving the sponsor to fill in the gaps themselves at real cost in time, risk, and trust.

”

Section 7

Early evidence from the field

The theoretical case for agentic AI is strong, but the more useful evidence is operational. Two areas of clinical development (trial master file management and site monitoring) have generated the earliest and most measurable results, because both are high-volume, high-friction, and well-suited to the blend of structure and repetition that agents handle best. Take the below results from a Medable case study with a top-five pharma.

99%

Documents correctly classified by an eTMF agent in real-world deployment

60%+

Fewer quality findings vs. manual TMF processing

62%

Reduction in CRA task time for data review and back-office monitoring

Data drawn from Medable case study results with a top-five pharmaceutical sponsor.

On the monitoring side, tasks that used to take an hour dropped to roughly 23 minutes across an estimated 40–70% of a typical CRA's workload, freeing capacity for the site relationships, risk interpretation, and judgment calls that no agent should be making. Just as notably, the humans doing the work reported a near-perfect ease-of-use score, and the large majority who tried the tool asked to keep using it, a signal that the gain is not just a number in a slide deck but a tool clinical teams actually want inside their day-to-day work.

This pattern, displacement of tactical burden rather than displacement of the role, is consistent with what a broader look at clinical monitoring's history would predict. The industry has been through two prior shifts in how monitoring is performed, moving from near-exhaustive source data verification in the 1990s and early 2000s to risk-based, centralized monitoring, codified in ICH E6(R2) in 2016. Both shifts redistributed human effort more efficiently. Neither changed who was doing the work. Agentic AI is the first shift that changes that premise, assuming continuous data surveillance, signal detection, issue triage, and pre-visit preparation at a pace no human team can match, while leaving judgment, site relationships, and risk interpretation with the CRA.

The effect compounds beyond any single study. At the study level, earlier signal detection compresses the distance between enrollment completion and database lock. At the portfolio level, liberated CRA capacity can be redeployed across more sites rather than absorbed by added headcount, travel intensity falls, and lower burnout in monitoring-intensive roles improves retention and site-relationship continuity, outcomes that stack into a durable operating advantage rather than a one-time efficiency gain.

“

Agentic AI is the first shift that changes that premise, assuming continuous data surveillance ... at a pace no human team can match, while leaving judgment, site relationships, and risk interpretation with the CRA.

”

Overcoming the two myths that stall adoption

Given the scale of the opportunity, the more common failure mode is not choosing the wrong agentic AI vendor. It is not starting at all. Two assumptions show up in nearly every boardroom conversation about AI adoption, and both are wrong in ways that cost sponsors real time while the gap between early movers and everyone else keeps widening.

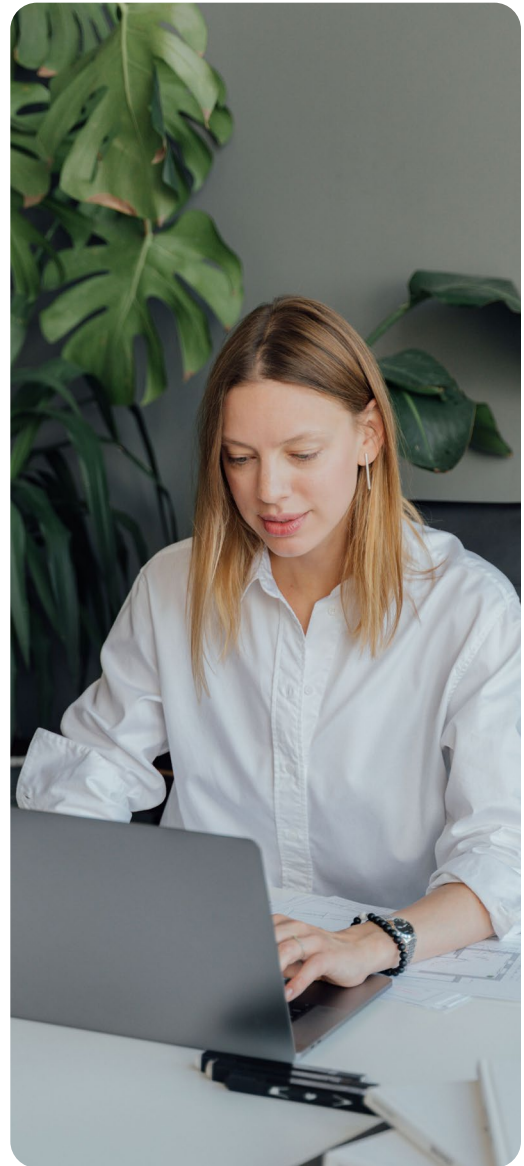
The “we already have a platform” myth

EDC, CTMS, and eTMF systems are systems of record. They store trial data, run studies, and hold years of validated history, and none of that needs to be replaced or disturbed to add an agentic layer. A properly built agent platform sits across the existing stack, connecting to those systems through purpose-built connectors, reasoning across the data they generate, and acting on the organization’s behalf (querying enrollment figures, flagging risk signals, summarizing TMF completeness) without touching a single underlying record. There is no migration project, no new system of record to manage, and no dependency on the AI vendor as a primary data store.

The “our data isn’t ready” myth

This is the myth most likely to quietly kill an initiative before it starts. The underlying concern is real, since the same concept genuinely does mean different things across systems, but the assumption that this has to be solved before AI can begin is outdated. A clinical ontology layer solves the mismatch without a data-transformation project, building a semantic mapping, grounded in recognized standards, that lets an agent traverse systems as though they already share a vocabulary, while the data stays exactly where it lives. What once required a team of data engineers months or years of calendar time can now happen in weeks, because the mapping itself can be proposed agentially, with a confidence score attached so a human reviewer decides what to accept.

For an organization on the fence, the responsible way to test this is with a narrow, low-risk deployment rather than a full commitment, using read-only connections to EDC, CTMS, and eTMF, alongside the communication tools a team already uses, to generate immediate and visible value without changing a single record. Trust for broader, write-enabled access is then earned incrementally, through demonstrated performance rather than assumed upfront, the same risk-proportionate logic that governs human oversight in Section 4.



“

A clinical ontology layer solves the mismatch without a data-transformation project, building a semantic mapping that lets an agent traverse systems as though they already share a vocabulary.

”

A framework for adoption

Organizations that succeed with agentic AI tend to follow a recognizable sequence rather than a single leap. The following framework synthesizes the common threads across successful deployments.

01 Change management: Secure early stakeholder and executive alignment

Involve executives, functional leaders, and end users from the outset to align on vision, priorities, and success metrics before implementation begins. Early engagement builds sponsorship and surfaces objections while they are still cheap to address.

02 Define scope and objectives clearly

Distinguish explicitly between using agents as augmentors of human roles versus automating a role outright, and communicate agent boundaries and interactions transparently so no stakeholder is surprised by what the system does or does not do.

03 Establish regulatory and ethical alignment first

Adopt a compliance-first posture aligned to ISO 27001, SOC 2, HIPAA, GDPR, and 21 CFR Part 11, with clear accountability, defined permissions, and reliable audit trails as a precondition, not an afterthought.

04 Pilot on low-risk use cases, then expand

Start with contained, reversible deployments that demonstrate value quickly. Flexible deployment tools and human-in-the-loop controls make it possible to expand agent autonomy incrementally as confidence and evidence accumulate.

05 Identify champions and address resistance directly

Internal advocates accelerate adoption and lend credibility to the effort. Equally, address skepticism openly with real use cases and realistic expectations rather than glossing over legitimate concerns.

06 Frame adoption as augmentation, not replacement

Transparent reasoning and explainable outputs help teams trust agent recommendations, and framing agents as collaborators rather than replacements sustains the human engagement the model depends on.

07 Choose a partner with genuine vertical depth

Evaluate any vendor's grounding in the specific vertical, and ideally the specific therapeutic area, in which the organization operates. A generic AI company can offer a platform; a specialized partner offers a framework that is pre-validated for the regulatory nuances of the field.

08 Monitor, measure, and optimize continuously

Real-time performance analytics and continuous feedback loops let organizations refine their approach over time, turning agent adoption into a compounding capability rather than a one-time deployment.

Validation and governance

Every question a compliance team asks about clinical AI eventually reduces to the same concern. Not whether the agent works on the day it goes live, but whether it will keep working, accurately, safely, and traceably, for the entire life of a study. That is a governance question, not a technology question, and it deserves a governance answer rather than a reassurance.

A different kind of software, and a different kind of validation

Traditional validated clinical systems are deterministic. The same input always produces the same output, which is exactly what makes one-time validation sufficient. Demonstrate correct behavior once, lock the system, and behavior stays correct. The validation event is effectively the governance.

Clinical AI agents do not work that way. They are adaptive by design, and that adaptability is precisely what makes them valuable over the life of a trial. As models improve, connected data sources change, and workflows evolve, agent outputs shift with them. This is not unique to AI in principle (even deterministic software goes through UAT, regression testing, and re-validation after releases and configuration changes), but AI extends the principle further. Trust cannot be established once and assumed thereafter. It has to be established continuously, through ongoing monitoring, validation, and oversight as the system evolves.

Traditional validated software	Clinical AI agents
Same input, same output, always	Outputs shift as models, data, and workflows evolve
One-time validation is sufficient to start, though updates still need re-verification	Continuous validation is required
Behavior is fixed at release	Behavior is adaptive by design
Governance is a one-time event per release	Governance is an ongoing operation

Framing this adaptability as a defect to be suppressed leads to the wrong governance posture. Adaptability is a design characteristic, and the right response is not to resist it but to make it visible, measurable, and correctable. That is what a continuously governed clinical AI platform is built to do.

What actually causes an agent's behavior to shift

Four sources of change affect any deployed clinical AI agent over time, and none of them are problems to prevent. They are realities to be governed.

- Model updates, where upstream providers push changes that can alter response patterns without any action by the sponsor
- Knowledge base changes, as study protocols, SOPs, regulatory guidance, and connected data sources evolve over the course of a trial
- Workflow and connector changes, as tools, triggers, and integrations mature alongside the clinical operation
- Prompt and configuration changes, the intentional improvements a team makes that still require controlled, documented change management in a regulated environment

The difference between governance that anticipates and governance that reacts shows up clearly in a simple scenario. Let's look at a scenario where an informed consent review agent scores 94% accuracy at baseline. In week six, a model provider pushes an update, and the agent's accuracy on consent document queries quietly drops to 71%. No one made a change, no alert fired, and no evaluation ran. Without continuous monitoring, that degradation stays invisible until a site coordinator surfaces a wrong answer after the damage is done, and the compliance team learns about it through a complaint rather than a monitoring report. With native performance monitoring, the same drop is caught at the next scheduled evaluation, before any user is affected, and the full event (detection, root cause, remediation, re-qualification) is documented automatically.

Continuous performance validation, always on from the moment of publish

Agent performance monitoring should begin the moment an agent goes live, running on configurable schedules (ad hoc, hourly, or daily) and producing scored results across every dimension relevant to the agent’s clinical function. Results are AI-assisted but human-reviewed before publication, with a human always in the loop. Configurable thresholds let a team set the bar for its own clinical context, with a warning firing at the first threshold and a critical alert firing if performance drops further, set independently per dimension and per deployment.

A well-built evaluation framework covers several core dimensions across the full behavioral surface of a clinical agent.

Dimension	What it measures
Response accuracy	Returns accurate, complete answers consistently across clinical workflows
Behavioral consistency	Maintains consistent responses and catches subtle reasoning drift early
Grounded responses	Stays within validated instructions, governed tools, and controlled knowledge boundaries
Safety	Maintains appropriate scope and content boundaries for regulated clinical use
Tool usage	Calls the correct tools, in the correct sequence, for each multi-step workflow
Latency	Responds within operationally acceptable time windows
Task delegation	Leverages connected agents and facilitates downstream workflows as designed
Task completion	Confirms the task was completed as specified
User criteria	Extensible to organization-specific workflows and compliance requirements
User sentiment	Captures direct user feedback on individual responses

The most important property of a framework like this is that it is generated from the agent’s actual configuration, its skills, connectors, instructions, and workflows, rather than applied as a generic benchmark. A generic evaluation tool assessing a Clinical Monitoring Agent has no way of knowing what that agent is supposed to do. It can measure general accuracy, but it cannot tell whether the agent called the right tools in the right order for a site query workflow, or whether it stayed inside the knowledge boundaries defined for that specific study. Evaluation built from the agent’s own configuration can. This is one of the foundational trust properties of the system, governance that actually knows the agent it is governing, and it should produce audit-ready QMS reporting by default rather than something assembled under pressure once an audit is already underway.



Complete, immutable traceability

Versioning closes a compliance gap that would otherwise be unanswerable. If a prompt is improved, a connector is updated, and a knowledge source changes, a regulated environment needs a reliable answer to a basic question, what was the agent doing on a given date, what changed, and who approved it. Without versioning, that answer requires reconstruction, and it tends to surface at the worst possible moment, mid-audit.

The underlying principle is that published agent configurations are immutable. Once an agent goes live, no one can quietly edit it. Any change to any component creates a new version with full attribution and a timestamp, while the prior version is preserved untouched. A change to a prompt never creates ambiguity about what else may have shifted, and teams can roll back to the last known good version in one click, with the rollback itself entering the version history rather than erasing it. Combined with continuous performance monitoring, this gives a compliance team something more valuable than a point-in-time snapshot, a complete and continuous account of the agent's configuration and behavior across the entire deployment period, produced automatically rather than assembled after the fact.

A lifecycle, not a layer

Performance monitoring and versioning work best as two parts of one governed lifecycle rather than two separate tools bolted together.

Phase	What happens	What agents gain
Build	Skills, connectors, and triggers are configured. Evaluation probes are auto-generated from that configuration, then reviewed and approved by a human before go-live	Baseline performance scores, human approval record
Operate	Scheduled and ad hoc evaluations run automatically, with dimension scores tracked against configurable thresholds and alerts triggered on warning or critical drops	Timestamped evaluation scores, alert history, QMS reports
Govern	Every component change creates a new version, with one-click rollback available and reports exportable for QMS filing at any time	Immutable version history, full audit trail, exportable QMS-ready reports

The distinction between native and bolted-on governance is operational, not cosmetic. A governance tool integrated from outside the platform has no access to the agent's actual configuration, so it applies generic benchmarks to a specific clinical workflow, introduces a new vendor relationship and integration surface, and creates a gap between the governance record and the agent record. Native governance knows the agent because it runs on the same platform, which means the audit trail and version history live in one place with no seam where risk can enter. When an auditor asks for evidence, the answer should not be spread across three systems. It should be a single, exportable, timestamped record that has been accumulating since the agent was published.

Why building this yourself carries the same risk as building the agent itself

The build-versus-buy calculus explored earlier in this paper applies with particular force to governance, arguably more so than to the agent itself. Testing non-deterministic systems is a specialist capability that traditional software QA was never built for. Because agents are adaptive by design, the same input can produce a different output on every run, not incorrectly, but inherently, which means there is no test suite that can assert a single fixed output. The right approach is evaluation, scored dimensions and AI-assisted judgment, rather than pass or fail assertions, and building that capability requires deep expertise in evaluation methodology and ongoing calibration as underlying models change.

There is also a blind spot that teams consistently underestimate. In a regulated environment, the governance system itself is a GxP software system, and it is not exempt from validation. That applies even to custom-built evaluation agents designed to monitor other agents. Teams routinely budget for the engineering effort to build an evaluation framework, but not for the validation workstream the framework itself requires (the IQ, OQ, and PQ protocols, the change control procedures, the re-validation triggered by every update). Without that second layer of validation, a homegrown governance system provides no real regulatory cover no matter how sophisticated its evaluation logic looks. With it, a team is quietly running two regulated software lifecycles in parallel, one for the agent and one for the system that watches it, and that second lifecycle is a permanent operational commitment rather than a one-time project. It requires a dedicated team to maintain evaluation frameworks as agents evolve, update the system as models change, and keep validation documentation current, for the life of every agent in production.

What this looks like when it works

Consider a CRA-facing agent answering site management questions. After a protocol amendment in week eight, response accuracy on enrollment threshold queries begins to slip, not because anyone changed the agent, but because the underlying knowledge base updated and the model's grounding shifted with it. The next scheduled evaluation flags the dimension below its warning threshold, the clinical operations team investigates, updates the knowledge base configuration, re-runs evaluations, and re-qualifies the agent before any site coordinator ever encounters a degraded response. The entire event, detection, root cause, remediation, re-qualification, and human approval, is already documented in the audit trail. Nothing needs to be reconstructed, because the evidence was produced automatically as the events happened.

That is the operational outcome continuous governance is meant to deliver. Not just compliance confidence after the fact, but the freedom to deploy with certainty rather than deploy and hope, because the team is not managing a governance process alongside the trial. Governance is simply running, in the background, for the life of the study.

Conclusion

From downtime to breakthroughs

The persistent white space inside clinical trials is a solvable problem, not an inevitable cost of doing business. Three decades of process innovation improved how efficiently people did the work; agentic AI is the first technology capable of changing who (or what) does it, while leaving the decisions that require clinical judgment exactly where they belong, with qualified people.

The path to capturing this advantage runs through infrastructure as much as ambition. An ontology layer that lets systems understand each other. A connector layer that lets agents reach live data without disrupting the systems of record an organization has spent years validating. A governance model that keeps humans firmly in the loop on the decisions that matter. And, a vendor selection built on genuine clinical fluency, not just technical capability (because in a regulated industry, a partner who has to be taught the difference between a Phase II and a Phase III trial is a partner an organization cannot afford).

The age of agentic is not a distant forecast. It is a present-day operating decision. The sponsors and CROs who begin building on agentic infrastructure now (starting with what they already have, expanding trust as performance earns it) will be the ones whose advantage compounds with every study that follows. Those who wait for perfect conditions will find that readiness was never the prerequisite. Instead, it was always the byproduct of starting.



Learn about how Medable's AI-enabled solutions are solving every day trial issues, [here.](#)

