

*White Paper – SIM Boston Life Sciences: CIO Roundtable*

# AI Starts with Plumbing: Governing Artificial Intelligence as a Sociotechnical Capability in Regulated Life Sciences.

Fernando Bardella <sup>1,\*</sup>, Stacey Fernandes <sup>2</sup> and James Bilotta <sup>3</sup>

<sup>1</sup> Research Fellow, Disruptive Technologies Lab (DTL) at Nuclear and Energy Research Institute (IPEN-CNEN);

<sup>2</sup> Chief Technology Officer (CTO), Versetal Information Systems;

<sup>3</sup> Subject Matter Expert (SME), Independent Advisor;

\* Correspondence: bardella@ipen.br (F.B.);

## Executive Summary

Artificial intelligence (AI) in regulated life sciences is moving from experimentation to enterprise deployment across clinical, regulatory, quality, commercial, safety, and operational workflows. The strategic question is no longer only which model, vendor, or platform to select, but whether the organization can use AI safely, repeatedly, and defensibly. This whitepaper argues that AI should be governed as a sociotechnical enterprise capability rather than a standalone technology tool. Responsible scale depends on trusted data, secure architecture, prompt and configuration governance, named accountability, human oversight, operational validation, lifecycle monitoring, and clear value-chain impact. Drawing on regulatory guidance, standards, academic literature, and CIO roundtable insights, the paper provides a practical operating model for life sciences leaders. Its central call to action is for CIOs to convene executive teams and boards around AI readiness, with governance designed to protect patients, preserve trust, and create measurable enterprise value.

**Keywords:** Artificial Intelligence; AI Governance; Regulated Life Sciences; Sociotechnical Systems; Data Governance; Human Oversight; Lifecycle Management; Patient-Centered Value

Published: 06/22/2026

**Copyright:** © 2026 by the author(s).  
This whitepaper is self-published on the NucleAI platform and distributed under the terms of the Creative Commons Attribution 4.0 International (CC BY 4.0) License.



## 1. Introduction

The current AI discussion in life sciences can too easily narrow around models, platforms, and vendors: which model to buy, which platform to standardize on, or which vendor is “ahead.” These questions matter, but they are not sufficient. In a regulated environment, the more important question is whether the organization can translate AI capability into reliable, auditable, secure, and value-producing operations.

The January 2026 FDA and EMA Guiding Principles of Good AI Practice in Drug Development mark an important shift in regulatory expectations. The document describes AI as system-level technology used to generate or analyze evidence across the drug product lifecycle, including nonclinical, clinical, post-marketing, and manufacturing phases [1]. Its principles

include human-centric design, a risk-based approach, adherence to standards, clear context of use, multidisciplinary expertise, data governance and documentation, model design and development practices, risk-based performance assessment, lifecycle management, and clear information for intended audiences [1]. In practical terms, this means AI cannot be governed only as an IT tool. It must be governed as an enterprise capability that connects technology, science, quality, privacy, legal, regulatory, business ownership, and human judgment.

The EMA reflection paper on AI in the medicinal product lifecycle reinforces the same direction. It states that sponsors, marketing authorization applicants or holders, and manufacturers are responsible for ensuring that algorithms, models, datasets, and data-processing pipelines are fit for purpose and aligned with legal, ethical, technical, scientific, and regulatory standards [2]. This creates a clear accountability signal: regulated organizations cannot outsource AI responsibility to vendors, data scientists, or innovation teams. The accountable institution remains responsible for the evidence it generates, the decisions it supports, and the risk it creates.

The strategic implication is simple but significant. AI advantage in life sciences will not come primarily from choosing one frontier model over another. Model performance will continue to change rapidly. The more durable advantage will come from building the organizational “plumbing” that allows AI to be safely adopted, monitored, scaled, corrected, and retired. That plumbing includes certified data products, access controls, lineage, metadata, secure development patterns, lifecycle validation, human oversight, cost governance, and decision logs.

This sociotechnical framing is also consistent with NIST’s treatment of AI bias, which emphasizes that AI systems are shaped not only by mathematical and computational constructs, but also by datasets, human interaction, organizational behavior, and social context [3]. For life sciences organizations, this is not an abstract distinction. AI-enabled workflows can influence evidence generation, operational decisions, patient identification, quality processes, safety monitoring, commercial execution, and ultimately trust in the institution.

This whitepaper was informed by discussions at the BoAF SIM Boston Life Sciences CIO Roundtable, held in Boston, Massachusetts, on April 29, 2026; additional event details and participant acknowledgements are provided in the Acknowledgements section. The discussion surfaced a recurring executive pattern: AI enthusiasm is high, but enterprise readiness is uneven. Organizations are moving quickly on pilots, copilots, agents, medical writing, clinical operations, commercial analytics, software development, and internal productivity use cases. Yet many still lack the governance structures, data foundations, and operating standards needed to move from experimentation to sustainable production.

The thesis of this paper is that regulated life sciences companies should treat AI as a sociotechnical capability. The model is only one component. The true system includes data, workflows, users, controls, incentives, vendors, policies, validation evidence, cost models, and the human decisions that AI influences.

### 1.1 Technical Context

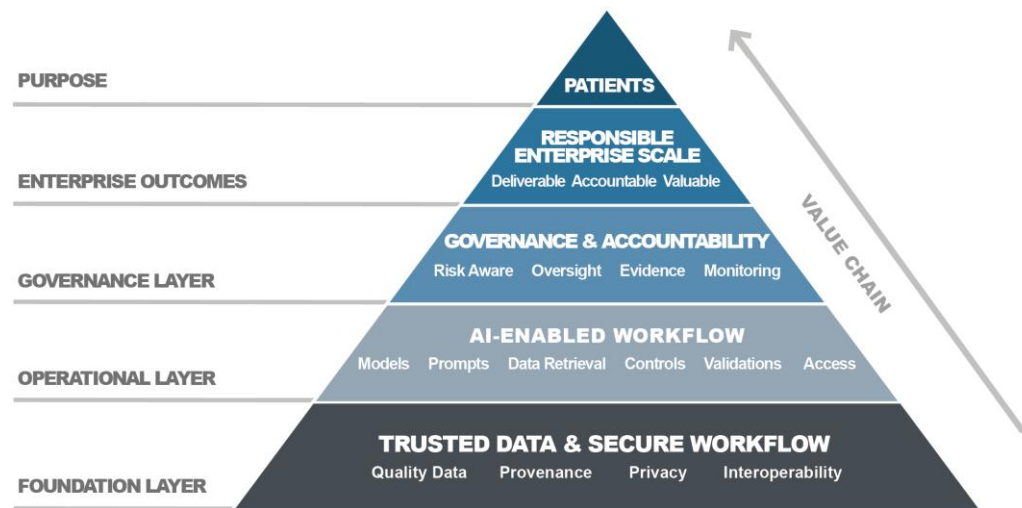
AI systems used in life sciences may include machine learning models, large language models, generative AI applications, agentic workflows, retrieval-augmented generation systems, predictive analytics, optimization models, and AI-assisted software development tools. Their technical behavior differs from traditional deterministic software in important ways.

A traditional application usually follows predefined logic. The same input, processed under the same conditions, generally produces the same output. By contrast, many AI systems infer patterns from data, generate probabilistic outputs, and may behave differently depending on model version, prompt structure, training data, retrieval sources, context window,

system instructions, and human interaction. The FDA draft guidance on AI for regulatory decision-making emphasizes the importance of defining the model's context of use, meaning the specific role and scope of the AI model in addressing a question of interest [4]. This concept is critical because the same model may be low risk in one workflow and high risk in another.

This distinction matters operationally. A generative AI tool used to summarize an internal meeting may pose a different risk profile than a model used to support clinical endpoint interpretation, manufacturing deviation analysis, pharmacovigilance triage, patient identification, or regulatory evidence generation. The model being used is only one part of that risk profile. How the model is instructed also matters: prompts, system instructions, retrieval context, configuration settings, and agent permissions can materially shape model behavior, output quality, and robustness [5]. As prompts are reused or embedded into workflows, they become governable assets requiring proper ownership, version control, testing, and access protection. They also create security risk, since prompt-based attacks can cause unintended system behavior or downstream effects [6]. The question is not only "Is the model good?" The better question is "Is the complete AI-enabled process fit for its intended use, with appropriate controls for its risk?"

Good Machine Learning Practice (GMLP) guidance for medical device development emphasizes safe, effective, and high-quality development of AI/ML-enabled medical devices across the total product lifecycle [7]. Although GMLP was developed in the medical device context, its operating logic is relevant to broader life sciences AI governance because it emphasizes intended use, data quality, performance evaluation, monitoring, and lifecycle control [7]. ISO/IEC 23053 also provides a conceptual framework for describing AI systems that use machine learning, reinforcing the importance of understanding AI systems as interacting components within a broader ecosystem rather than as isolated models [8].



**Figure 1. From AI Capability to Patient-Centered Value.** This figure shows AI value in regulated life sciences as a value chain built from the bottom up. Trusted data and secure architecture provide the foundation for governed AI-enabled workflows; those workflows require governance, accountability, validation, monitoring, and evidence; responsible enterprise scale then enables reliable operations, regulatory defensibility, and measurable value; the ultimate purpose is better outcomes for patients, people, and society. **Source:** Author's synthesis based on FDA/EMA AI principles, EMA lifecycle guidance, NIST AI RMF, NIST Generative AI Profile, ISO/IEC 42001, and ISO/IEC 23894 [1], [2], [9], [6], [10], [11].

NIST's AI Risk Management Framework provides a cross-sector governance structure organized around four functions: Govern, Map, Measure, and Manage [9]. For life sciences organizations, those functions can be translated into practical operating questions. Govern asks who owns AI risk. Map asks where AI is being used and what it affects. Measure asks how performance, bias, safety, security, privacy, and reliability are tested. Manage asks what the organization does when risk exceeds tolerance.

Generative AI adds additional complexity. NIST's Generative AI Profile identifies risks that are novel to or intensified by generative AI, including information integrity, data privacy, harmful bias, information security, intellectual property exposure, and value-chain or component integration risk [6]. These risks are not theoretical in enterprise life sciences settings. They appear when employees use public tools for regulated data, when AI assistants generate software code without adequate review, when copilots surface over-permissioned documents, when prompts leak sensitive business information, or when a model-generated answer is trusted beyond its validated scope. The practical technical context, therefore, is not "model versus model." It is architecture, controls, data readiness, lifecycle evidence, and accountability.

Figure 1 summarizes the paper's central argument: AI value in regulated life sciences is not created by model capability alone. It is built through a value chain that starts with trusted data and secure architecture, moves through governed AI-enabled workflows and accountable enterprise operations, and ultimately serves patients, people, and society. The upward arrow emphasizes that patient-centered value depends on each layer below it being reliable, governed, and fit for purpose.

## 2. Problem Description

AI adoption in life sciences is constrained less by the absence of use cases than by the maturity of the organizational systems required to govern them. Many organizations can identify promising applications for generative AI, machine learning, copilots, and agentic workflows. Fewer can consistently determine which use cases are appropriate, which data are fit for purpose, which controls are proportionate, which outputs are defensible, and which risks remain acceptable over time.

This distinction is especially important in regulated life sciences because AI systems do not operate in isolation. They interact with scientific evidence, business processes, human judgment, regulated records, clinical and commercial workflows, privacy obligations, quality systems, and patient-facing consequences. Recent FDA and EMA principles emphasize that AI should be managed through human-centric design, context of use, data governance, risk-based performance assessment, lifecycle management, and clear information for intended audiences [1]. The EMA reflection paper similarly emphasizes that AI risk depends on the technology, the data, the context of use, the influence on decision-making, and the stage of the medicinal product lifecycle in which the system is used [2].

The problem, therefore, is not simply that AI introduces new technology risk. The deeper issue is that AI exposes weaknesses in data governance, operating accountability, lifecycle discipline, enterprise architecture, and institutional decision-making. The following subsections describe the principal failure modes that life sciences organizations must address if AI is to move from experimentation to responsible scale.

### 2.1. Model Selection Is Being Mistaken for AI Strategy

A common failure mode is to treat AI strategy as a model-selection or vendor-selection decision. In this framing, the organization assumes that selecting a leading model provider or

enterprise AI platform will resolve the strategic question. This view is incomplete because the model is only one component of the AI-enabled system.

In regulated life sciences, the strategic asset is the governed workflow built around the model. That workflow includes the data sources, retrieval mechanisms, access controls, human review points, validation evidence, audit trails, cost thresholds, change-control processes, and business ownership required to make AI use repeatable and defensible. NIST's AI Risk Management Framework supports this broader view by organizing AI risk management around governance, mapping, measurement, and management rather than around model performance alone [9]. NIST's Generative AI Profile also highlights risks related to value-chain and component integration, information integrity, privacy, information security, harmful bias, and intellectual property exposure, all of which extend beyond the model layer itself [6].

Model capability is evolving rapidly, and many organizations will use multiple models over time. The strategic question is whether that heterogeneity is governed by design or created accidentally through disconnected team-level decisions. A defensible architecture should allow models to be evaluated, isolated, retired, or replaced without requiring the organization to redesign the surrounding business process.

One response is to introduce an abstraction layer above the model, either through a cloud-based AI platform or an internal model gateway. Platforms such as Amazon Bedrock, Microsoft Foundry Models, and Google Model Garden illustrate this direction by providing access to multiple models or model ecosystems through a managed platform layer [12]–[14]. This can make model evaluation and routing more practical, but it does not eliminate dependency. It may shift dependency from a model provider to a cloud or platform provider.

Model abstraction must not be mistaken for friction-free model substitution. Platform gateways lower low-level API connection overhead, but they do not eliminate the technical debt of re-validating prompt frameworks, re-tuning system instructions, re-indexing vector databases, or rebuilding evaluation evidence across dissimilar model architectures. In a GxP-regulated context, any underlying frontier model swap represents a major change-control operation, requiring rigorous regression testing and documented confirmation that prior validation evidence remains legally and operationally defensible.

Emerging interoperability approaches such as the Model Context Protocol may help standardize how AI applications connect to external tools, data sources, and workflows [15]. Its relevance is not model hot-swapping, but context, data, and tool integration. For regulated organizations, the governance implication remains the same: portability only creates value when access control, auditability, validation evidence, and lifecycle ownership remain explicit.

Switching costs also differ by use case. General-purpose activities such as drafting, summarization, and exploratory analysis may be relatively model-agnostic. Retrieval-augmented generation workflows, pharmacovigilance signal workflows, clinical document review systems, and validated data pipelines are different. In those cases, the enterprise asset is not only the model. It includes prompt design, retrieval corpus, data permissions, validation evidence, audit trail, monitoring process, and human review workflow. Prompt design itself can materially shape model behavior, output quality, and robustness, which reinforces why prompts should be treated as governable workflow assets rather than disposable instructions [5].

This does not mean that vendor selection is unimportant. It means that vendor selection should be subordinated to context of use, data control, interoperability, monitoring, supportability, concentration risk, and exit planning. In this sense, the question is not only which model performs best today, but whether the enterprise architecture can preserve

accountability and operational continuity as models, vendors, platforms, and regulatory expectations change.

## 2.2. Data Foundations Are Often Too Weak for AI-Scale Trust

The roundtable discussion repeatedly returned to the same point: AI starts with data. This is not a new observation, but AI changes its urgency. Models and agents can reuse, recombine, summarize, and expose information at a speed and scale that make weak data foundations more visible and more consequential.

Traditional data governance often focused on lineage, ownership, quality, and access. AI requires those controls, but it also requires richer contextual evidence. For regulated life sciences, the organization must understand not only where data came from, but whether it is fit for the intended use, representative of the relevant population or process, legally usable, scientifically meaningful, and appropriately separated across training, validation, and testing activities. The EMA reflection paper identifies class imbalance, data leakage, representative data, test datasets, training datasets, validation datasets, interpretability, and model transparency as important AI and machine learning considerations [2].

The FDA and EMA principles emphasize that data source provenance, processing steps, and analytical decisions should be documented in a traceable and verifiable manner, aligned with GxP requirements where applicable [1]. ICH E6(R3) similarly emphasizes data integrity, metadata, traceability, confidentiality, documented computerized-system responsibilities, and appropriate procedures for computerized systems used in clinical trials [16]. This is a higher bar than simply making more data available in a lakehouse or analytics platform. It implies that AI-ready data must be curated, described, governed, protected, and connected to the context in which it will be used.

The FAIR principles define scientific data stewardship around findability, accessibility, interoperability, and reusability [17]. They also emphasize machine actionability, meaning that data should be structured so computational systems can find and use it with reduced manual intervention [17]. For AI in life sciences, FAIR should not be interpreted as a mandate to make all data broadly open. It should be interpreted as a requirement to make governed data and metadata usable by authorized humans and machines under clear access, purpose, and quality controls.

The central issue is therefore not only data availability. It is data trust. AI systems that operate on poorly defined, weakly governed, poorly orchestrated, or contextually inappropriate data may produce outputs that appear authoritative while resting on evidence that is incomplete, biased, stale, over-permissioned, incorrectly sequenced, or not fit for purpose. In AI-enabled workflows, how data pipelines run, when they refresh, how dependencies are sequenced, and whether upstream changes are detected can materially affect the reliability of downstream outputs.

## 2.3. AI Pilots Are Outrunning Lifecycle Governance

Many organizations can produce AI pilots. Fewer can sustain AI products. This gap reflects a lifecycle governance problem rather than a proof-of-concept problem. An AI proof of concept may perform well in a controlled demonstration but degrade when exposed to new users, new workflows, new regions, new source systems, or new model versions. The machine learning literature describes dataset shift as a condition in which the joint distribution of inputs and outputs differs between training and use environments [18]. Concept drift further describes situations in which the relationship between input variables and the target variable

changes over time [19]. These concepts are directly relevant to life sciences AI because clinical practice, patient populations, commercial territories, manufacturing processes, regulatory expectations, and source systems can all change after deployment.

The execution challenge is not only technical. AI systems that move from pilot to enterprise use must be treated as learning platforms, not static tools. They require continuous feedback, monitoring, governance, business ownership, and improvement. This is consistent with enterprise MLOps research showing that production AI adoption is constrained by organizational, technical, operational, and business challenges, not only by model performance [20].

This issue is especially visible in less mature life sciences organizations, where clinical, commercial, quality, regulatory, and external partner data may be distributed across CROs, vendors, spreadsheets, applications, and legacy systems. In these environments, conflicting methods, unclear data ownership, and disagreement over the system of record can prevent a single source of truth from being established. The same system-of-record problem that has long constrained traditional data warehouse programs can also constrain enterprise AI adoption. Technology may evolve quickly, but business processes, ownership models, and data definitions often evolve more slowly.

The EMA reflection paper notes that AI risk may vary throughout the lifecycle and that relevant risks should be systematically managed from early development through decommissioning [2]. The FDA and EMA principles similarly call for lifecycle management, scheduled monitoring, periodic re-evaluation, and drift management [1]. These expectations are consistent with the broader observation that machine learning systems can accumulate hidden technical debt because they combine traditional software maintenance challenges with data dependencies, model dependencies, feedback loops, and changing external conditions [21]. Production-readiness work in machine learning further emphasizes that deployed models require tests and monitoring across data, model behavior, infrastructure, and operational performance [22]. MLOps literature similarly frames production AI as an operational discipline involving workflows, roles, automation, monitoring, and lifecycle management rather than model training alone [23].

The risk is especially high when organizations treat AI as a one-time implementation. AI-enabled systems require monitoring because the environment around them changes. Source systems change. Business definitions change. Patient populations change. Commercial and clinical workflows change. Regulatory expectations change. Model providers update their systems. Retrieval libraries expand. Users adapt their behavior. Costs shift. Performance at launch is therefore not evidence of continued fitness for use.

Without lifecycle governance and business ownership, AI becomes a portfolio of unmanaged experiments. The organization may accumulate models, agents, prompts, integrations, dashboards, and copilots without a clear record of what remains in use, what has changed, what is still monitored, and what should be retired.

#### **2.4. Human Oversight Is Often Too Vague to Be Useful**

“Human in the loop” is often used as if it resolves AI risk by itself. It does not. Human oversight only functions as a meaningful control when the human reviewer has the expertise, authority, time, context, and evidence required to challenge the AI output. The problem is broader than output review. In generative AI and agentic workflows, risk is also shaped by the inputs and configurations that guide system behavior. Prompts, system instructions, retrieval context, role definitions, embedded business rules, and guardrails can materially influence what the AI produces and how users interpret its output. These elements should therefore be treated

as governed assets, not informal user preferences. In regulated life sciences, prompt libraries, context configurations, and reusable instructions may require ownership, version control, access restrictions, testing, approval workflows, and protection from inappropriate reuse or disclosure. NIST's Generative AI Profile highlights risks related to information integrity, privacy, intellectual property exposure, and value-chain integration, all of which can be affected by how generative AI systems are instructed, contextualized, and deployed [6].

This also creates a bias and blind-spot risk. NIST's work on identifying and managing AI bias emphasizes that AI bias is not only a statistical or model-design issue; it can emerge from datasets, human interaction, organizational processes, deployment context, and social assumptions [3]. In generative AI workflows, unmanaged prompts, retrieval context, and configuration settings can therefore propagate bias or skewed assumptions even when the underlying model appears technically capable.

The human-factors literature has long shown that automation can be used, misused, disused, or abused depending on trust, workload, perceived risk, system design, and user behavior [24]. This matters for AI governance because a nominal reviewer may become a weak control if the workflow encourages passive acceptance, if the AI output appears authoritative without sufficient evidence, or if the reviewer lacks an escalation path. Human-AI interaction research further emphasizes the need to support appropriate reliance, communicate system uncertainty, make system behavior understandable, and enable efficient correction when AI systems are wrong [25].

FDA and EMA principles emphasize risk-based performance assessment of the complete system, including human-AI interactions [1]. This is important because AI risk often emerges at the boundary between model output, prompt design, retrieval context, and human decision-making. A model may be technically acceptable in isolation but unsafe or inappropriate if users misunderstand its scope, overtrust its output, apply it outside its context of use, or lack the domain expertise needed to identify error.

In regulated life sciences, the right framing is not merely "human in the loop." It is accountable professional and accountable institution. Human oversight should be understood as an active governance mechanism, not a passive checkpoint. The business function that uses AI should own the outcome, the intended use, and the associated business risk. IT, data, legal, quality, privacy, security, regulatory, and compliance functions should enable, advise, and control, but they should not become substitutes for business ownership.

The governance challenge is therefore organizational. Shared ownership cannot mean no ownership. Cross-functional AI governance is necessary precisely because AI risk cuts across business, technology, legal, quality, regulatory, privacy, security, HR, medical, clinical, and commercial domains. A governance council can help prevent functions from working in silos, but it must also clarify decision rights, escalations, accountability, and ownership of both AI inputs and outputs. This reinforces the broader position that responsible AI governance depends not only on technical controls, but also on diverse human judgment, multidisciplinary review, and institutional accountability across the full lifecycle of AI-enabled systems [26].

## **2.5. Shadow AI and AI-Assisted Development Create New Operational Risk**

Generative AI lowers the barrier to building software, automating workflows, analyzing documents, and creating business-facing tools. This creates legitimate productivity opportunities, but it also expands the risk of uncontrolled AI use outside approved enterprise patterns. A business analyst can now build a prototype that previously required a software team. A developer can ask an AI assistant to refactor code, generate tests, write integration logic, or



optimize database queries. A knowledge worker can connect a copilot to shared files and surface content that had been technically accessible but socially invisible. These patterns accelerate experimentation, but they can also bypass architecture, validation, cybersecurity, privacy, records management, and quality-system controls. For regulated records and electronic signatures, 21 CFR Part 11 applies to electronic records created, modified, maintained, archived, retrieved, or transmitted under FDA records requirements [27]. Part 11 also requires controls such as system validation, record protection, limited system access, secure computer-generated audit trails, authority checks, operational checks, and documentation controls [27]. Even when Part 11 does not apply directly, its control logic remains instructive because regulated organizations need evidence that systems behave as intended, records are protected, and changes can be traced.

FDA's Computer Software Assurance guidance supports a risk-based approach in which assurance activities and evidence are commensurate with the intended use and risk of software used in production or quality-system contexts [28]. NIST's Secure Software Development Framework provides secure software development practices that can be integrated into software development lifecycle models to reduce software vulnerabilities [29]. These sources support a central point for AI-assisted development: generated code should not bypass the assurance, security, testing, and change-control standards that apply to human-written code.

The validation debt problem becomes more serious when AI-generated code is used inside the control infrastructure itself. If AI assistants generate scripts for lineage checks, data-quality tests, MLOps monitoring, validation test harnesses, or automated release gates, the organization must also validate the validator. In these cases, defects in AI-generated control code can create false confidence by allowing flawed data products, model updates, or workflow changes to pass through governance undetected. For regulated environments, AI-generated code used in testing, monitoring, or assurance infrastructure should therefore receive heightened review, independent verification, and change-control treatment.

AI-assisted development also complicates authorship, validation, and accountability. Commercial coding assistants can accelerate software delivery, but they may create validation debt if generated code enters regulated or business-critical environments without the same review, testing, documentation, secure development, dependency management, and change-control standards applied to human-written code. This concern is consistent with the broader machine learning technical-debt literature, which warns that rapid gains from ML systems can create compounding maintenance and reliability costs when dependencies, feedback loops, and system interactions are not managed deliberately [21].

If AI-generated code is accepted into a validated or business-critical environment, the organization still owns the resulting software, evidence, defects, and operational risk. The fact that software was generated quickly does not reduce the need for assurance. In some cases, it increases it.

## **2.6. AI Governance Is Being Separated from Business Value**

Some organizations create long lists of AI use cases without connecting them to enterprise priorities. Others apply restrictive controls without distinguishing between low-risk productivity use and high-impact regulated workflows. Both approaches weaken adoption.

AI governance should not be a bureaucracy that slows all work equally. It should be a decision system that helps the organization allocate attention, risk controls, and investment proportionate to value and exposure. ISO/IEC 42001 establishes requirements for establishing, implementing, maintaining, and continually improving an AI management system within

organizations [10]. ISO/IEC 23894 provides guidance for organizations that develop, produce, deploy, or use AI systems to manage AI-specific risk and integrate AI risk management into organizational activities and functions [11]. The EU AI Act similarly reflects a risk-based regulatory model for AI and includes requirements for high-risk AI systems such as risk management, data governance, transparency, human oversight, and post-market monitoring [30].

The practical implication is that governance should enable prioritization. Low-risk productivity use cases may be scaled broadly when appropriate guardrails are in place. Higher-impact use cases require stronger evidence because they may affect clinical development, regulatory decisions, pharmacovigilance, manufacturing quality, commercial execution, patient identification, or patient-facing interactions.

AI governance should therefore be linked to value-chain impact. A disciplined portfolio should distinguish between AI initiatives that improve local productivity and AI initiatives that change enterprise outcomes, such as reducing clinical trial cycle time, improving safety monitoring, strengthening regulatory readiness, accelerating medical writing, reducing manufacturing investigation time, or improving commercial effectiveness. Each strategic use case should have a value hypothesis, risk classification, accountable sponsor, monitoring plan, and explicit stop criteria. AI governance should protect the enterprise from reckless adoption, but it should also protect the enterprise from endless pilots that never scale. The objective is not control for its own sake. The objective is responsible value creation.

## 2.7. Trust Is Still Attached to People, Not Data Products

In many organizations, users trust a report because they trust the person or team who created it. This trust model does not scale well to AI. AI systems need inputs that are discoverable, governed, documented, and reusable. The organization must move from person-based trust to product-based trust.

A certified data product should have a named owner, business definition, technical lineage, quality metrics, permissible-use rules, access controls, refresh cadence, metadata, known limitations, and monitoring procedures. The same logic should apply to AI components. Internal algorithmic auditing literature supports the use of lifecycle documentation and audit artifacts to reduce the accountability gap in the development and deployment of AI systems [31]. Model cards were proposed to improve transparency around model reporting, including intended use, performance characteristics, evaluation procedures, ethical considerations, and limitations [32]. Datasheets for datasets were proposed to document dataset motivation, composition, collection process, recommended uses, and other properties that help dataset creators and consumers communicate more clearly [33].

In regulated life sciences, these concepts should be adapted into internal evidence standards. The purpose is not to copy academic documentation templates mechanically. The purpose is to create enough evidence for users, auditors, quality teams, regulators, and executive sponsors to understand what a data product or AI component is intended to do, what it should not be used for, what evidence supports it, and what limitations remain.

The goal is not paperwork. The goal is defensibility. When a regulator, auditor, executive committee, or internal quality review asks why a model was used, what data supported it, which subgroup performance was evaluated, who approved the use case, and how drift is monitored, the organization should not depend on memory.

### 3. Solution / Operating Model

Life sciences organizations need a practical AI operating model that connects governance, data, architecture, validation, security, cost management, and human capability. The goal is not to create a separate AI bureaucracy, but to make responsible AI adoption repeatable. In regulated environments, this requires operating mechanisms that are clear enough for business adoption, rigorous enough for quality and regulatory scrutiny, and flexible enough to keep pace with rapid technical change.

#### 3.1. Establish AI Governance as an Outcome-Based Accountability Function

AI governance should be designed as an outcome-based accountability function, not as an adoption accelerator or an IT control layer alone. In regulated life sciences, its primary purpose is to ensure that AI-enabled workflows are safe, compliant, defensible, and connected to meaningful enterprise value. Effective governance must be able to approve, constrain, redesign, escalate, or stop AI use cases when risk, evidence, ownership, or value are insufficient. The objective is not frictionless adoption; it is appropriate control at the right decision points, with clear escalation paths and accountable decisions.

This framing is consistent with AI impact assessment guidance, which emphasizes identifying, evaluating, and documenting the potential impacts of AI systems across their lifecycle [34]. It is also aligned with responsible AI governance literature, which distinguishes high-level principles from the organizational practices needed to operationalize them, including structural, relational, and procedural mechanisms [35].

In practice, AI governance requires executive sponsorship and cross-functional participation, but participation should not be confused with diffuse ownership. IT, data, cybersecurity, privacy, legal, regulatory, quality, clinical, medical, commercial, HR, and business functions each see different dimensions of AI risk and value. However, shared input cannot mean unclear accountability. For each AI-enabled workflow, especially those with clinical, regulatory, quality, safety, commercial, or patient-facing impact, there should be a named accountable owner responsible for the outcome, risk acceptance, and ongoing fitness for use.

A governance council should define the enterprise rules of engagement while preserving clear accountability at the use-case, workflow, or model level. Its role is to establish standards for context of use, risk classification, data readiness, prompt and configuration governance, technical architecture, validation evidence, human oversight, monitoring, and retirement. FDA guidance on AI for regulatory decision-making emphasizes credibility assessment in relation to the model's context of use [4]. The NIST AI RMF Playbook translates governance, mapping, measurement, and management into actions organizations can adapt to their own risk profiles [36].

Risk classification should not be treated as a one-time intake decision. Initial risk tiers can become stale as models, prompts, data sources, retrieval corpora, workflows, users, vendors, and output uses change. Governance should therefore include reassessment triggers when there is a material change in deployment context, data drift, model behavior, prompt configuration, access permissions, regulatory relevance, or business use. This is consistent with risk-based regulatory expectations for ongoing risk management, documentation, human oversight, and post-deployment monitoring [30], [37].

The practical test of AI governance is whether it produces safe, explainable, value-generating, and defensible decisions. A strong governance model makes ownership visible, creates evidence for auditability and regulatory review, and gives the organization the discipline to scale, redesign, or stop AI-enabled workflows when appropriate.

### 3.2. Build a Controlled AI Landing Zone

A controlled AI landing zone should be established before shadow AI becomes the default operating model. The purpose of such an environment is not to constrain innovation, but to make safe experimentation easier than unsafe experimentation. Organizations should provide approved access to AI capabilities while preserving visibility into data movement, model use, identity, access, cost, and output handling.

A practical AI landing zone should bring together approved model access, identity and access management, role-based permissions, data loss prevention, logging, prompt and response retention rules, private retrieval sources, secure integration patterns, human review workflows, token and cost monitoring, and vendor off-ramp options. For sensitive intellectual property or regulated data, organizations should also evaluate company-controlled cloud enclaves or private deployment patterns that preserve enterprise control over prompts, retrieval content, model interaction records, and downstream use.

A controlled AI landing zone should also treat data access as dynamic rather than static. Traditional folder permissions and shared-drive structures were often designed for human navigation, not for AI-enabled semantic search, retrieval, and summarization. When copilots or retrieval systems index legacy repositories, they can expose over-permissioned, outdated, confidential, or context-blind content to users who would not have discovered it through normal workflows. For this reason, access control should be evaluated at the level of the indexed corpus, retrieval layer, user role, data sensitivity, and intended use, not only at the level of legacy file permissions. This dynamic-access approach is consistent with zero trust architecture principles, which emphasize explicit authentication and authorization, granular access decisions, and least-privilege access to enterprise resources [38]. The landing zone should also draw from established security and privacy control catalogs for access control, auditability, configuration management, system integrity, and privacy risk management [39].

This is especially important in life sciences because AI systems may touch patient data, clinical data, adverse-event information, commercial strategy, manufacturing investigations, regulatory correspondence, confidential research, and trade secrets. The organization should be able to answer where data went, who accessed it, which model processed it, what output was produced, and whether that output influenced a business or regulated decision. These questions are not only technical. They are questions of institutional accountability.

For organizations that cannot yet enforce dynamic retrieval-layer permissions, the safer fallback is to narrow the AI indexing corpus rather than expose unmanaged legacy repositories. Legacy shared drives, broad SharePoint libraries, and poorly permissioned file stores should be excluded from AI indexing until access rights, ownership, retention, and sensitivity classifications have been reviewed. In these environments, the practical starting point is a curated corpus of approved data products, controlled repositories, and access-reviewed knowledge sources. This is not the mature-state architecture, but it is a defensible compensating control while the organization remediates legacy data access.

### 3.3. Productize Data Before Scaling AI

The fastest path to AI value is not necessarily to centralize all enterprise data into a perfect platform. That ambition is often too slow and too broad to support near-term adoption. A more practical approach is to identify the data domains most relevant to high-value workflows and convert the most reused datasets into governed data products. This approach is consistent with contemporary data product and data mesh literature, which frames data products as managed, reusable assets designed to satisfy recurring information needs and create

value, while data mesh emphasizes domain ownership, data-as-a-product, self-service infrastructure, and federated governance [40].

A data product is more than a table, report, dashboard, or data extract. It is a governed asset with a defined owner, business meaning, technical lineage, quality expectations, permissible uses, access rules, refresh cadence, known limitations, and monitoring process. For AI use, these characteristics become even more important because models and agents may reuse data at speed and scale, often outside the immediate context in which the data was originally created. This product-oriented approach changes the AI conversation. Instead of asking each project team to rediscover the meaning, quality, permission model, and limitations of a dataset, the organization creates reusable data assets that can support multiple AI-enabled workflows. This is particularly important for clinical, commercial, medical, regulatory, quality, and manufacturing domains, where inconsistent definitions or weak lineage can create downstream risk.

The standard should be pragmatic. Organizations do not need to wait for a perfect enterprise data foundation before scaling AI. They do, however, need to distinguish between data that is merely available and data that is fit for purpose. AI-ready data should have enough metadata, lineage, quality evidence, access control, and business ownership to support its intended use. In regulated life sciences, “good enough” must still be defensible.

### 3.4. Implement an AI Lifecycle Standard

AI should be governed across its lifecycle, not only at the moment of deployment. A proof of concept can demonstrate feasibility, but it does not establish that an AI-enabled process is fit for ongoing use. Lifecycle governance is essential because AI performance, risk, cost, and value may change as data, users, workflows, prompts, models, vendors, and operating conditions evolve.

A practical lifecycle standard should begin with a clear use-case definition. The organization should specify the business problem or desired outcome, intended users, decision impact, expected value, patient or regulatory relevance, and context of use. This definition establishes the basis for risk classification and determines which evidence, controls, approvals, and monitoring obligations are proportionate.

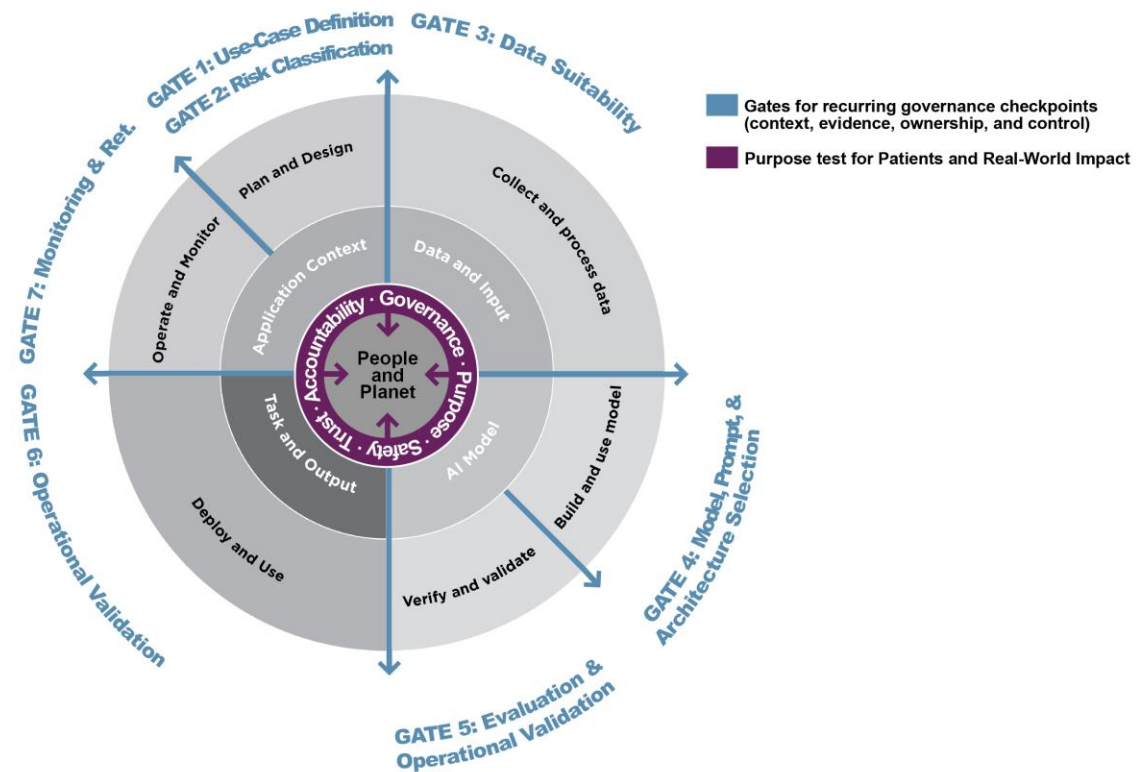
The next step is data suitability. Before model selection becomes the focus, the organization should assess whether the data is appropriate for the intended use. This includes provenance, quality, representativeness, access rights, leakage risk, bias considerations, orchestration, and separation of training, validation, and testing datasets where applicable. The model, model set, prompt design, and architecture decision should then follow the risk profile, not the other way around.

Evaluation and validation should be tied to the intended-use metrics defined at the beginning of the lifecycle, not to generic benchmarks alone. Evaluation should test whether the AI-enabled workflow meets practical acceptance criteria for the task, including subgroup performance, human-AI interaction, robustness, explainability, security, and known failure modes [1], [2]. Validation should confirm that the system performs as intended in its actual operational context, including users, interfaces, data sources, prompts, retrieval logic, downstream decisions, and human review procedures [2], [7], [16]. For adaptive or frequently updated models, the validation plan should define how performance will be re-confirmed after model, prompt, data, retrieval, or workflow changes, particularly where dataset shift, concept drift, or technical debt may affect production performance [18], [19], [21]–[23]. For higher-risk

workflows, a comparable test or parallel environment can help assess material updates before production release [28].

Monitoring should be treated as operational infrastructure, not as an audit document. The monitoring plan should define ownership, review frequency, thresholds, reassessment triggers, and required actions when performance changes. For executive readers, drift should be understood in three practical forms: data drift, when production inputs no longer resemble the data used during development; concept drift, when the correct answer, decision logic, or operating context changes over time; and model drift, when the underlying model or model behavior changes through vendor updates, retraining, prompt changes, configuration changes, or retrieval changes. For higher-risk workflows, monitoring should be resourced, executed, and periodically tested as part of normal operations, consistent with lifecycle governance, risk-based assurance, and accountability expectations [1], [2], [18], [19], [21]–[23].

The scientific and operational point is that AI governance cannot end at launch. Performance at launch is a baseline, not a guarantee. A defensible lifecycle standard allows an organization to scale AI while maintaining evidence that the AI-enabled workflow remains appropriate for its intended use.



**Figure 2. AI Lifecycle Governance Model for Regulated Life Sciences.** Adapted from the OECD AI system lifecycle model, this figure maps regulated life sciences AI governance gates onto the lifecycle from application context and data processing through model development, verification and validation, deployment, operation, monitoring, and potential retirement. The seven gates translate the lifecycle into executive decision points for assessing context, evidence, ownership, control, and continued fitness for use. Patients, people, and society are placed at the center to emphasize purpose, safety, trust, and accountability across the lifecycle. **Source:** Adapted from the OECD AI system lifecycle model, with life sciences governance gates developed by the author based on FDA/EMA Good AI Practice principles, EMA lifecycle guidance, NIST AI RMF, and OECD accountability guidance [1], [2], [9], [37].

Figure 2 should be read as a governance map, not as a linear project plan. The OECD lifecycle provides the structure, while the seven gates translate that lifecycle into executive decision points. Each gate asks whether the organization has enough context, evidence, ownership, and control to allow the AI-enabled workflow to proceed, change, scale, or remain in operation. The gates are intentionally recursive. A workflow may return to an earlier gate when the model changes, the data changes, the prompt is modified, the intended use expands, the risk profile increases, or monitoring shows performance degradation. In this sense, the gates are not a waterfall sequence. They are recurring accountability checkpoints that help leadership preserve purpose, safety, trust, and defensibility across the AI lifecycle. Table 1 translates these lifecycle checkpoints into seven executive governance gates, each framed as a decision question and a corresponding governance focus.

These standards should be understood not only as technical safeguards, but also as governance mechanisms that improve the quality of executive decision-making [1,3,4,13,14]. In that sense, the purpose of AI oversight in life sciences is not merely to improve model performance. It is to improve the reliability, defensibility, and trustworthiness of the decisions organizations make with and about AI systems [1,3,4].

**Table 1. AI Lifecycle Governance Gates for Regulated Life Sciences.** This table translates the AI lifecycle model into seven executive governance gates. Each gate defines a practical decision point for assessing whether an AI-enabled workflow has sufficient context, evidence, ownership, and control to proceed, change, scale, remain in operation, or be retired. **Source:** Author's synthesis based on the OECD AI system lifecycle model, FDA/EMA Good AI Practice principles, EMA lifecycle guidance, NIST AI RMF, and OECD accountability guidance [1], [2], [9], [37].

| Gate                                                     | Executive Question                                                                                                                       | Governance Focus                                                                                                           |
|----------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------|
| <b>Gate 1: Use-Case Definition</b>                       | What problem, goal, or outcome is the AI-enabled workflow intended to support, and who will use or be affected by it?                    | Context of use, intended users, value hypothesis, patient or regulatory relevance, success metrics                         |
| <b>Gate 2: Risk Classification</b>                       | What level of risk does the use case create?                                                                                             | Decision impact, patient relevance, regulatory exposure, data sensitivity, automation level, accountability requirements   |
| <b>Gate 3: Data Suitability</b>                          | Are the data sources, lineage, quality, permissions, orchestration, and representativeness fit for the intended use?                     | Data provenance, access rights, data quality, bias considerations, pipeline reliability, fitness for purpose               |
| <b>Gate 4: Model, Prompt, and Architecture Selection</b> | Are the model or model set, prompts, retrieval logic, configurations, and integration patterns appropriate for the risk profile?         | Model choice, prompt design, retrieval architecture, agent permissions, configuration controls, portability considerations |
| <b>Gate 5: Evaluation and Operational Validation</b>     | Does the workflow meet intended-use metrics in the operating context where it will actually be used?                                     | Acceptance criteria, subgroup performance, human-AI interaction, robustness, explainability, security, failure modes       |
| <b>Gate 6: Deployment Approval</b>                       | Has the workflow been approved for controlled use?                                                                                       | Release approval, trained users, access controls, SOPs, change baseline, monitoring obligations, named accountability      |
| <b>Gate 7: Monitoring and Retirement</b>                 | Does the workflow continue to perform as intended, and are there clear triggers for reassessment, redesign, revalidation, or retirement? | Drift monitoring, usage, cost, access, incidents, model updates, reassessment triggers, retirement criteria                |

### 3.5. Treat Human Oversight as a Designed Control

Human oversight should be treated as a designed control, not as a general assurance statement. In regulated life sciences, the presence of a human reviewer does not automatically reduce risk unless the reviewer has the expertise, authority, time, context, and evidence required to challenge the AI output. A nominal “human in the loop” can become a weak control if the workflow encourages passive acceptance, if the basis for the AI recommendation is opaque, or if the reviewer lacks a clear escalation path.

A more defensible approach is to define human oversight in relation to the AI system’s context of use. For low-risk productivity tasks, oversight may focus on user awareness, appropriate-use boundaries, and confirmation that the output is not used as regulated evidence. For higher-risk workflows, such as clinical, regulatory, quality, safety, manufacturing, commercial compliance, or patient-facing decisions, oversight should include explicit review criteria, documented accountability, escalation procedures, and evidence that the reviewer can meaningfully accept, reject, modify, or challenge the AI-generated output.

This distinction is consistent with the FDA and EMA emphasis on assessing the complete AI system, including human-AI interactions, and on managing AI performance across the lifecycle in a risk-based manner [1]. It also aligns with the broader regulatory principle that the organization remains accountable for the evidence it generates, the decisions it supports, and the risks created by AI-enabled processes [2], [4]. Meaningful oversight should therefore be treated as part of the accountability ecosystem for AI, not merely as a final human review step [37].

In practice, meaningful oversight requires more than assigning a reviewer at the end of the workflow. It requires designing the workflow so that uncertainty, limitations, source evidence, confidence signals, and known failure modes are visible to the human decision-maker. It also requires training, decision rights, auditability, and a clear definition of when AI output may be used, when it must be independently verified, and when it must not be used at all.

For life sciences CIOs, this is an important governance design point. Human oversight should not be positioned as a symbolic safeguard added after deployment. It should be embedded into the operating model as a control that connects technical design, risk classification, business ownership, and patient impact.

### 3.6. Govern AI Cost as a Design Constraint

AI cost governance should be treated as an architectural and operating concern, not merely as a finance afterthought. Token usage, embedding refresh frequency, retrieval strategy, model tier, inference volume, logging, storage, evaluation runs, and user behavior can all materially affect the economics of AI-enabled workflows. Without intentional design, organizations may scale pilots that are technically impressive but economically unsustainable.

Cost management should therefore be connected to architecture. Not every task requires the most advanced or expensive model. Some workflows may be better served by smaller models, deterministic automation, business rules, retrieval tools, or traditional analytics. Other workflows may justify more expensive models when the value of reasoning, synthesis, or decision support is materially higher. This aligns with cloud resource optimization literature, which emphasizes cost visibility, informed decision-making, and tailored optimization strategies across heterogeneous cloud environments [41].

AI-specific cost governance also requires model routing, usage thresholds, consumption dashboards, alerts, and clear ownership of budget and value. Large language model inference can create substantial computational and memory demands, and efficiency must be addressed across data-level, model-level, and system-level techniques rather than through



reporting alone [42]. Where multiple models are available, routing strategies can help balance cost, latency, scalability, and performance instead of sending every request to the most capable or most expensive model [43].

An AI initiative should not be scaled only because it is technically possible. It should be scaled because the expected value justifies the cost, risk, and operating burden. Cost visibility helps leaders compare initiatives, stop low-value experimentation, and direct investment toward the parts of the value chain where AI can create meaningful impact.

### 3.7. Manage AI as a Portfolio, Not a List of Pilots

AI use cases should be managed as a portfolio rather than as a disconnected collection of pilots. In many organizations, enthusiasm produces long lists of possible use cases, but not a clear method for deciding which ones should be funded, scaled, redesigned, or stopped. This leads to experimentation without enterprise learning.

**Table 2. From AI Risk to Enterprise Control Evidence.** Common AI risks in regulated life sciences can be translated into practical controls, evidence artifacts, and regulatory or standards anchors that support auditability, accountability, lifecycle governance, and responsible scale. **Source:** Author's synthesis based on the regulatory, standards, and scholarly sources cited in the "Regulatory / Standards Anchor" column.

| Risk Domain                                  | Failure Mode                                                                                  | Enterprise Control                                                                                                                    | Evidence Artifact                                                                                                             | Regulatory / Standards                                                                             |
|----------------------------------------------|-----------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------|
| <b>Context of Use &amp; Value</b>            | AI is used without a clear business problem, desired outcome, or intended purpose             | Use-case intake, context-of-use definition, value hypothesis, and risk classification                                                 | Approved use-case record, context-of-use statement, value hypothesis, and risk-tier assessment                                | FDA/EMA Good AI Practice; FDA AI Regulatory Decision-Making Guidance [1], [4]                      |
| <b>Accountability</b>                        | Shared input becomes unclear ownership or no ownership                                        | Named accountable owner for outcome, risk acceptance, and ongoing fitness for use                                                     | Accountability record, approval log, governance council decision, and escalation path                                         | NIST AI RMF; ISO/IEC 42005; OECD accountability guidance [9], [34], [37]                           |
| <b>Data Trust</b>                            | AI uses data that is incomplete, biased, stale, poorly orchestrated, or not fit for purpose   | Certified data products, data suitability review, lineage, access controls, and pipeline monitoring                                   | Data product card, lineage record, quality report, pipeline run log, and access review                                        | FDA/EMA Good AI Practice; EMA AI lifecycle paper; ICH E6(R3); FAIR principles [1], [2], [16], [17] |
| <b>Prompt &amp; Configuration Governance</b> | Prompts, system instructions, retrieval settings, or agent permissions change without control | Prompt ownership, version control, testing, access protection, and change review                                                      | Prompt register, version history, test evidence, and approval log                                                             | Prompt-engineering literature; NIST Generative AI Profile [5], [6]                                 |
| <b>Model, Model Set &amp; Architecture</b>   | Model choice, model routing, or platform dependency creates unmanaged risk                    | Model/model-set evaluation, architecture review, abstraction-layer assessment, portability requirements, and exit planning            | Architecture decision record, model evaluation record, vendor/platform risk assessment, portability assessment, and exit plan | NIST AI RMF; NIST Generative AI Profile [6], [9]                                                   |
| <b>Access &amp; Retrieval</b>                | AI exposes over-permissioned, confidential, or context-blind content                          | Dynamic access controls, retrieval-layer permissions, DLP, indexed-corpus review, and exclusion of unmanaged repositories when needed | Access review, retrieval source inventory, DLP report, indexed-corpus approval, and audit log                                 | NIST Zero Trust Architecture; NIST SP 800-53 [38], [39]                                            |

A mature AI portfolio should distinguish between broad productivity enablement and strategic value-chain transformation. Productivity use cases, such as document summarization, meeting synthesis, internal knowledge search, draft generation, and workflow automation, can build literacy and momentum when they are governed appropriately. Strategic use cases require stronger scrutiny because they are tied to core outcomes such as accelerating clinical development, improving patient identification, strengthening regulatory readiness, enhancing pharmacovigilance triage, reducing manufacturing investigation time, or improving commercial execution.

Portfolio governance should measure both value and learning. Some initiatives will fail technically; others will work technically but not justify the operating cost or risk. For this reason, each initiative should have a value hypothesis, accountable sponsor, risk classification, monitoring plan, and explicit stop criteria. A practical stop trigger may be reached when the cost of maintaining performance, correcting drift, securing data access, or sustaining human review exceeds the operational value generated by the AI-enabled workflow.

The discipline of stopping matters. Without stop criteria, organizations accumulate AI debt: pilots that remain alive, tools that are not adopted, workflows that are not monitored, and costs that persist without clear value. A portfolio approach makes AI investment more transparent and connects innovation to enterprise strategy.

#### 4. Conclusion

The central conclusion of this whitepaper is that AI readiness in regulated life sciences is a form of enterprise readiness. The performance of an individual model is important, but it is not sufficient to determine whether an AI-enabled workflow can be trusted, scaled, or defended. The more consequential question is whether the organization can connect models or model sets to governed data, controlled architecture, prompt and configuration governance, documented context of use, named human accountability, lifecycle monitoring, and measurable business value.

This conclusion is increasingly aligned with the direction of regulatory and standards-based guidance. FDA and EMA principles emphasize that AI in drug development should be managed through human-centric design, risk-based performance assessment, data governance, lifecycle management, and clear information for intended audiences [1]. EMA's lifecycle reflection paper further reinforces that AI risk depends on the interaction among the technology, the data, the intended use, the decision context, and the stage of the medicinal product lifecycle in which the system is applied [2]. NIST's AI Risk Management Framework similarly frames AI risk management as a governance, mapping, measurement, and management discipline rather than as a narrow model-performance exercise [9]. ISO/IEC 42001 and ISO/IEC 23894 reinforce the need for organizational management systems and risk-management practices that can be embedded into institutional processes [10], [11].

The cost of inadequate AI governance is therefore not limited to inaccurate model output. In life sciences, weak governance can produce invalid evidence, biased patient identification, poor safety signal handling, uncontrolled software, privacy leakage, over-permissioned knowledge access, unmanaged prompt or configuration changes, untraceable regulatory work, escalating cost, vendor or platform dependency, and loss of executive trust. These risks can translate into inspection exposure, delayed decisions, operational waste, reputational damage, and, in the most consequential cases, patient harm. The machine learning literature on hidden technical debt reinforces that deployed AI systems can accumulate risk through

unmanaged data dependencies, model dependencies, feedback loops, and changing operating conditions [21].

The value of stronger governance is equally material. Effective AI governance creates the right control at the right decision points. Trusted data products reduce repeated discovery work and improve reuse. Prompt and configuration governance makes AI behavior more traceable and consistent. Lifecycle controls help leaders understand whether an AI-enabled workflow remains fit for purpose after deployment, including whether data drift, concept drift, or model drift requires reassessment. Designed human oversight clarifies when AI is assisting judgment, when it is exceeding its intended role, and when a professional or institution must intervene. This view is consistent with OECD work linking AI accountability to lifecycle risk management and institutional governance [37].

For life sciences CIOs, the leadership implication is clear. AI should be led neither as a purely technical program nor as an unconstrained innovation agenda. It should be led as a sociotechnical enterprise capability: technical because it depends on models, prompts, data, platforms, cybersecurity, validation, architecture, and monitoring; social because it changes work, decision-making, accountability, culture, trust, and patient-facing consequences. Human-centered AI governance depends not only on technical controls, but also on diverse judgment, multidisciplinary review, named accountability, and institutional responsibility [26].

The organizations most likely to create durable AI value will not simply be those with access to the most advanced models. They will be those that can define which decisions they intend to improve, which data can be trusted, which prompts and configurations are governed, which controls are proportionate, which risks are acceptable, which human responsibilities remain non-delegable, and which value-chain outcomes matter most. In regulated life sciences, that is the difference between adopting AI and governing AI well.

## 5. Call to Action

Life sciences CIOs should make AI readiness an executive agenda item, not only a technology initiative. The immediate call to action is clear: CIOs should convene their executive teams, boards, and cross-functional leaders to assess whether the organization is truly prepared to scale AI safely, repeatedly, and defensibly. The discussion should move beyond model selection and vendor comparison. The real question is whether the enterprise has the governed data, architecture, prompt and configuration controls, human oversight, lifecycle monitoring, accountability model, and value discipline required to use AI responsibly in regulated life sciences.

This conversation matters because AI is no longer confined to experimentation. It is becoming part of how evidence is generated, how work is performed, how decisions are supported, and how risk is distributed across the enterprise. As a result, AI must be treated as a sociotechnical capability with real human consequences, potential patient impact, and material implications for value creation across the life sciences value chain.

The CIO's role is not simply to provide access to AI tools. It is to help the organization define where AI should be used, where it should not be used, what risks are acceptable, what controls are required, and what value the organization expects to create. Every AI use case should have a clearly defined context of use, business owner, named accountable owner, risk classification, human accountability model, prompt and configuration controls, lifecycle monitoring requirements, and value hypothesis.

The practical next step is to establish an enterprise AI readiness review. This review should identify the highest-value AI opportunities, assess whether underlying data, prompts, model sets, access controls, and architecture are fit for purpose, define human oversight expectations, confirm monitoring requirements for data drift, concept drift, and model drift, and determine which pilots should be scaled, redesigned, or stopped. This should not be treated as a one-time governance exercise, but as the foundation for an operating model that connects AI adoption to patient impact, regulatory defensibility, and measurable business value.

Most importantly, CIOs should help leadership teams ask the right question. The issue is no longer only, “Which AI model should we use?” It is, “Which decisions are we prepared to augment, which risks are we prepared to accept, which prompts, data, models, and workflows must be governed, which parts of the value chain are we trying to improve, and which human responsibilities must remain non-delegable?”

AI will not create trust or value by itself. Trust will come from disciplined leadership, named accountability, defensible decisions, transparent evidence, and operating models that recognize the human and patient impact of AI-enabled systems. Value will come from applying AI to the workflows, decisions, and value-chain constraints that matter most. For life sciences CIOs, the opportunity is to lead this shift now: from AI experimentation to accountable enterprise scale.

**Acknowledgements:** The authors gratefully acknowledge the participants of the BoAF SIM Boston Life Sciences CIO Roundtable, held in Boston, Massachusetts, on April 29, 2026, chaired by Fernando Bardella and Jim Bilotta. The session was hosted and sponsored by PwC and provided the practical executive context that informed this whitepaper’s discussion of artificial intelligence, data foundations, governance, and responsible scaling in regulated life sciences.

The authors especially thank the CIOs and technology leaders who participated in the discussion and contributed perspectives that helped shape the themes explored in this whitepaper: Adi Shah, Amit Mathur, Andreas Bauer, Anthony Martin, Brian Blood, Darren Tong, Jagan Balasubramanian, Jason Manfredi, Jeff Russo, John Larkin, John Masso, John Robbins, Kenneth Matta, Kevin Durfee, Kevin Lind, Mark Yunger, Matt Nerney, Michael Belmarsh, Michael Tobin, Mike Fowler, Orhan Karsligil, Ragy Emile, Raj Padmanabhan, Ryan White, Sean Spingler, Srini Parthasarathy, Stacey Fernandes, Steve Aschettino, Steven Swan, Subhashish Dhar, Suhas Mulki, Surya Wiryo, Dana J. Seehale, and Tim Fox. \*

\* Participation in the discussion does not imply endorsement of all views or conclusions presented in this whitepaper.

The Society for Information Management (SIM) and the SIM Boston Life Sciences SIG are volunteer-led, not-for-profit professional communities open to technology leaders and practitioners.

More information about joining SIM, participating in CIO roundtables and other events can be found at <https://www.simboston.org> and <https://www.simboston.org/boston/boston-programs/life-sciences>

## Reference List

- [1] U.S. Food and Drug Administration and European Medicines Agency, *Guiding Principles of Good AI Practice in Drug Development*, Jan. 2026. [Online]. Available: <https://www.fda.gov/media/189581/download>. Accessed: Jun. 8, 2026.
- [2] European Medicines Agency, *Reflection Paper on the Use of Artificial Intelligence (AI) in the Medicinal Product Lifecycle*, EMA/CHMP/CVMP/83833/2023, Sep. 2024. [Online]. Available: [https://www.ema.europa.eu/en/documents/scientific-guide-line/reflection-paper-use-artificial-intelligence-ai-medicinal-product-lifecycle\\_en.pdf](https://www.ema.europa.eu/en/documents/scientific-guide-line/reflection-paper-use-artificial-intelligence-ai-medicinal-product-lifecycle_en.pdf). Accessed: Jun. 8, 2026.
- [3] R. Schwartz, A. Vassilev, K. Greene, L. Perine, A. Burt, and P. Hall, *Towards a Standard for Identifying and Managing Bias in Artificial Intelligence*, NIST Special Publication 1270. Gaithersburg, MD, USA: National Institute of Standards and Technology, Mar. 2022. doi: 10.6028/NIST.SP.1270.
- [4] 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*, Jan. 2025. [Online]. Available: <https://www.fda.gov/media/184830/download>. Accessed: Jun. 8, 2026.
- [5] B. Chen, Z. Zhang, N. Langrené, and S. Zhu, “Unleashing the potential of prompt engineering for large language models,” *Patterns*, vol. 6, no. 6, Art. no. 101260, 2025. doi: 10.1016/j.patter.2025.101260.
- [6] National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1. Gaithersburg, MD, USA: National Institute of Standards and Technology, Jul. 2024. [Online]. Available: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>. Accessed: Jun. 8, 2026.
- [7] International Medical Device Regulators Forum, *Good Machine Learning Practice for Medical Device Development: Guiding Principles*, IMDRF/AIML WG/N88 FINAL:2025, Jan. 2025. [Online]. Available: [https://www.imdrf.org/sites/default/files/2025-02/IMDRF\\_AIML%20WG\\_GMLP\\_N88%20Final.pdf](https://www.imdrf.org/sites/default/files/2025-02/IMDRF_AIML%20WG_GMLP_N88%20Final.pdf). Accessed: Jun. 8, 2026.
- [8] International Organization for Standardization, *ISO/IEC 23053:2022, Framework for Artificial Intelligence (AI) Systems Using Machine Learning (ML)*. Geneva, Switzerland: ISO, 2022. [Online]. Available: <https://www.iso.org/standard/74438.html>. Accessed: Jun. 8, 2026.
- [9] National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1. Gaithersburg, MD, USA: National Institute of Standards and Technology, Jan. 2023. [Online]. Available: <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>. Accessed: Jun. 8, 2026.
- [10] International Organization for Standardization, *ISO/IEC 42001:2023, Information Technology—Artificial Intelligence—Management System*. Geneva, Switzerland: ISO, 2023. [Online]. Available: <https://www.iso.org/standard/81230.html>. Accessed: Jun. 8, 2026.
- [11] International Organization for Standardization, *ISO/IEC 23894:2023, Information Technology—Artificial Intelligence—Guidance on Risk Management*. Geneva, Switzerland: ISO, 2023. [Online]. Available: <https://www.iso.org/standard/77304.html>. Accessed: Jun. 8, 2026.
- [12] Amazon Web Services, “What is Amazon Bedrock?,” *Amazon Bedrock User Guide*. [Online]. Available: <https://docs.aws.amazon.com/bedrock/latest/userguide/what-is-bedrock.html>. Accessed: Jun. 8, 2026.
- [13] Microsoft, “Foundry Models from partners and community,” *Microsoft Learn*. [Online]. Available: <https://learn.microsoft.com/en-us/azure/ai-foundry/foundry-models/concepts/models-from-partners>. Accessed: Jun. 8, 2026.
- [14] Google Cloud, “Model Garden on Gemini Enterprise Agent Platform.” [Online]. Available: <https://cloud.google.com/model-garden>. Accessed: Jun. 8, 2026.
- [15] Model Context Protocol, “Introduction.” [Online]. Available: <https://modelcontextprotocol.io/docs/getting-started/intro>. Accessed: Jun. 8, 2026.

- [16] International Council for Harmonisation of Technical Requirements for Pharmaceuticals for Human Use, *ICH E6(R3) Guideline for Good Clinical Practice*, Jan. 2025. [Online]. Available: [https://database.ich.org/sites/default/files/ICH\\_E6%28R3%29\\_Step4\\_Final-Guideline\\_2025\\_0106.pdf](https://database.ich.org/sites/default/files/ICH_E6%28R3%29_Step4_Final-Guideline_2025_0106.pdf). Accessed: Jun. 8, 2026.
- [17] M. D. Wilkinson *et al.*, “The FAIR guiding principles for scientific data management and stewardship,” *Scientific Data*, vol. 3, Art. no. 160018, Mar. 2016. doi: 10.1038/sdata.2016.18.
- [18] J. Quiñonero-Candela, M. Sugiyama, A. Schwaighofer, and N. D. Lawrence, Eds., *Dataset Shift in Machine Learning*. Cambridge, MA, USA: MIT Press, 2009.
- [19] J. Gama, I. Žliobaitė, A. Bifet, M. Pechenizkiy, and A. Bouchachia, “A survey on concept drift adaptation,” *ACM Computing Surveys*, vol. 46, no. 4, pp. 1–37, Mar. 2014. doi: 10.1145/2523813.
- [20] C. Amrit and A. Kolar Narayanappa, “An analysis of the challenges in the adoption of MLOps,” *Journal of Innovation & Knowledge*, vol. 10, no. 1, Art. no. 100637, 2025. doi: 10.1016/j.jik.2024.100637.
- [21] D. Sculley *et al.*, “Hidden technical debt in machine learning systems,” in *Proc. 28th International Conference on Neural Information Processing Systems*, Montreal, QC, Canada, 2015, pp. 2503–2511.
- [22] E. Breck, S. Cai, E. Nielsen, M. Salib, and D. Sculley, “The ML test score: A rubric for ML production readiness and technical debt reduction,” in *Proc. IEEE International Conference on Big Data*, Boston, MA, USA, 2017, pp. 1123–1132. doi: 10.1109/Big-Data.2017.8258038.
- [23] D. Kreuzberger, N. Kühn, and S. Hirschl, “Machine Learning Operations (MLOps): Overview, definition, and architecture,” *IEEE Access*, vol. 11, pp. 31866–31879, 2023. doi: 10.1109/ACCESS.2023.3262138.
- [24] R. Parasuraman and V. Riley, “Humans and automation: Use, misuse, disuse, abuse,” *Human Factors*, vol. 39, no. 2, pp. 230–253, Jun. 1997. doi: 10.1518/001872097778543886.
- [25] S. Amershi *et al.*, “Guidelines for human-AI interaction,” in *Proc. CHI Conference on Human Factors in Computing Systems*, Glasgow, Scotland, U.K., 2019, Paper no. 3, pp. 1–13. doi: 10.1145/3290605.3300233.
- [26] F. Bardella, *Why Diverse Thinking Builds Better AI: Human-Centered Governance for Global Life Sciences*. NucleAI, 2026. doi: 10.5281/zenodo.19682427.
- [27] Electronic Code of Federal Regulations, “Title 21, Chapter I, Subchapter A, Part 11 — Electronic Records; Electronic Signatures.” [Online]. Available: <https://www.ecfr.gov/current/title-21/chapter-I/subchapter-A/part-11>. Accessed: Jun. 8, 2026.
- [28] U.S. Food and Drug Administration, *Computer Software Assurance for Production and Quality System Software*, Sep. 2025. [Online]. Available: <https://www.fda.gov/media/188844/download>. Accessed: Jun. 8, 2026.
- [29] M. Souppaya, K. Scarfone, and D. Dodson, *Secure Software Development Framework (SSDF) Version 1.1: Recommendations for Mitigating the Risk of Software Vulnerabilities*, NIST Special Publication 800-218. Gaithersburg, MD, USA: National Institute of Standards and Technology, Feb. 2022. doi: 10.6028/NIST.SP.800-218.
- [30] European Parliament and Council of the European Union, “Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence,” *Official Journal of the European Union*, Jun. 13, 2024. [Online]. Available: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>. Accessed: Jun. 8, 2026.
- [31] I. D. Raji *et al.*, “Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing,” in *Proc. 2020 Conference on Fairness, Accountability, and Transparency*, Barcelona, Spain, 2020, pp. 33–44. doi: 10.1145/3351095.3372873.
- [32] M. Mitchell, S. Wu, A. Zaldivar, P. Barnes, L. Vasserman, B. Hutchinson, E. Spitzer, I. D. Raji, and T. Gebru, “Model cards for model reporting,” in *Proc. Conference on Fairness, Accountability, and Transparency*, Atlanta, GA, USA, 2019, pp. 220–229. doi: 10.1145/3287560.3287596.

- [33] T. Gebru, J. Morgenstern, B. Vecchione, J. W. Vaughan, H. Wallach, H. Daumé III, and K. Crawford, "Datasheets for datasets," *Communications of the ACM*, vol. 64, no. 12, pp. 86–92, Dec. 2021. doi: 10.1145/3458723.
- [34] International Organization for Standardization, *ISO/IEC 42005:2025, Information Technology—Artificial Intelligence—AI System Impact Assessment*. Geneva, Switzerland: ISO, 2025.
- [35] E. Papagiannidis, P. Mikalef, and K. Conboy, "Responsible artificial intelligence governance: A review and research framework," *The Journal of Strategic Information Systems*, vol. 34, no. 2, Art. no. 101885, 2025. doi: 10.1016/j.jsis.2024.101885.
- [36] National Institute of Standards and Technology, *NIST AI Risk Management Framework Playbook*. [Online]. Available: <https://airc.nist.gov/airmf-resources/playbook/>. Accessed: Jun. 8, 2026.
- [37] Organisation for Economic Co-operation and Development, *Advancing Accountability in AI: Governing and Managing Risks Throughout the Lifecycle for Trustworthy AI*, OECD Digital Economy Papers, no. 349. Paris, France: OECD Publishing, 2023. doi: 10.1787/2448f04b-en.
- [38] S. Rose, O. Borchert, S. Mitchell, and S. Connelly, *Zero Trust Architecture*, NIST Special Publication 800-207. Gaithersburg, MD, USA: National Institute of Standards and Technology, Aug. 2020. doi: 10.6028/NIST.SP.800-207.
- [39] Joint Task Force, *Security and Privacy Controls for Information Systems and Organizations*, NIST Special Publication 800-53 Rev. 5. Gaithersburg, MD, USA: National Institute of Standards and Technology, Sep. 2020. doi: 10.6028/NIST.SP.800-53r5.
- [40] I. Blohm, F. Wortmann, C. Legner, and F. Köbler, "Data products, data mesh, and data fabric," *Business & Information Systems Engineering*, vol. 66, pp. 643–652, 2024. doi: 10.1007/s12599-024-00876-5.
- [41] P. Nawrocki, M. Cybulski, and J. Sniezynski, "A survey of cloud resource consumption optimization methods," *Journal of Grid Computing*, vol. 23, 2025. doi: 10.1007/s10723-024-09792-0.
- [42] Z. Zhou *et al.*, "A survey on efficient inference for large language models," arXiv:2404.14294, 2024. [Online]. Available: <https://arxiv.org/abs/2404.14294>. Accessed: Jun. 8, 2026.
- [43] C. Varangot-Reille, S. Casola, N. Kaushik, R. L. Ogundepo, M. Gallé, D. Filippova, and D. Alistarh, "Doing more with less: A survey on routing strategies for resource optimisation in large language model-based systems," arXiv:2502.00409, 2025. [Online]. Available: <https://arxiv.org/abs/2502.00409>. Accessed: Jun. 8, 2026.

**Disclaimer/Publisher's Note:** This article is self-published by its author(s) through the NucleAI platform. All statements, opinions, and data are solely those of the author(s) and contributor(s) and do not necessarily represent the views of NucleAI, its editors, or affiliated organizations. NucleAI disclaims responsibility for any harm, loss, or damage resulting from the use of any ideas, methods, instructions, technologies, or products described in this publication.