

Validation Strategies for Agentic AI Platforms in GxP-Regulated Clinical Data Environments

RAJESH HAGALWADI, Maxis AI, Edison, NJ USA

ABSTRACT

Agentic AI platforms, composed of collaborating digital agents, are increasingly used to orchestrate complex clinical data workflows ranging from data ingestion and transformation to quality checks and documentation. In GxP-regulated environments, these capabilities must be validated with the same rigor as other computerized systems, while addressing new challenges introduced by learning components, multi-agent interactions, and continuous change. This paper presents a problem-first, risk-based framework for validating agentic AI platforms used in clinical data programming and statistical computing.

Since late 2025, regulators and industry bodies have sharpened their expectations for AI in drug development, emphasizing intended use and context of use, risk-based credibility across the lifecycle, human–AI interaction, and transparent information about performance and limitations. Updated GAMP 5 AI/ML validation guidance reinforces risk-based planning, data strategy, and continuous monitoring for AI components in GxP systems, including treatment of electronic records and audit trails. Life-science organizations are also reporting real-world agentic AI use cases that traverse data silos and automate multistep processes, underscoring the need for robust platform-level validation and agent-level lifecycle controls that keep pace with these deployments.

This paper presents two complementary contributions for GxP clinical data environments. First, it outlines an AI Platform Validation Framework that structures validation across Concept, Requirements, Development and Design, and Acceptance and Release phases, with ongoing operation and monitoring aligned to risk-based computer software assurance and 21 CFR Part 11 expectations. Second, it defines AI Agent Lifecycle Phases for validation, safety, and governance: Concept, Requirements, Design, Verification and Validation, Acceptance and Release, and Operations and Continuous Oversight that can be applied to individual agents running on the platform.

This paper also describes how to scope AI uses, segment risks, design validation plans that cover both platform and agents, and implement lifecycle monitoring—including data drift and model updates—within GxP expectations. Practical patterns for documenting agent behavior, guardrails, and human-in-the-loop oversight are outlined, along with illustrative scenarios. The result is a pragmatic, regulator-aligned approach to validating agentic AI platforms, enabling innovation in clinical data science while preserving patient safety, data integrity, and audit readiness.

INTRODUCTION

Agentic AI platforms promise to streamline and augment clinical data workflows by coordinating multiple specialized AI agents to perform tasks such as data ingestion, quality review, mapping and documentation. In this context, an AI agent can be viewed as a bounded component often tool-using and context-aware with a defined intended use and autonomy level, such as a data mapping agent or an anomaly-detection agent. The agentic AI platform provides the orchestration layer that plans, sequences, and supervises these agents under policy constraints, creating system-level behavior that must be validated as a GxP-relevant computerized system.

In clinical data environments governed by GxP, these platform-level capabilities must be introduced carefully. The platform and its agents need to be demonstrably fit for their intended use and context of use, controlled, and documented in a way that satisfies regulatory expectations for computerized systems and AI-enabled functions. Historically, computer system validation (CSV) in clinical research focused on relatively static systems whose behavior changed infrequently and in predictable ways. Agentic AI systems, by contrast, may incorporate AI/ML components that are sensitive to data, orchestrate many interdependent agents, and interact directly with business-critical workflows, demanding explicit lifecycle controls, risk-based credibility assessment, and human oversight.

Recent regulatory developments—especially joint EMA–FDA guiding principles on AI in drug development—stress lifecycle management, risk-based performance assessment, consideration of human–AI interaction, and transparency about AI purpose, performance, and limitations. In parallel, GAMP 5 Appendix D11 and related ISPE materials

emphasize risk-based computer software assurance, AI maturity, and supporting processes such as risk management, data governance, and change management throughout the AI subsystem life cycle.

This paper focuses on:

- Applying an AI Platform Validation Framework to agentic AI platforms in clinical data environments, organizing validation across concept, requirements, development and design, and acceptance and release phases, with risk-based testing (including IQ/OQ where appropriate), traceability, and summary reporting.
- Applying AI Agent Lifecycle Phases for validation, safety, and governance to individual agents on the platform, including risk assessment (clinical context, Model influence, decision consequence), guardrails design, verification and validation against measurable performance criteria, and continuous oversight

The target readers are Clinical Data Programmers, Statisticians, Data scientists, Forward Deployed Engineers (FDE), IT, CSV and quality-compliance professionals responsible for introducing and maintaining agentic AI within GxP frameworks.

INDUSTRY CONTEXT: 2025–2026 AI & VALIDATION TRENDS

Regulatory principles for AI in drug development

In early 2026, EMA and FDA issued aligned guiding principles for AI use in drug development, building on earlier EMA AI reflection papers and agency-specific guidance. These principles stress

- **Lifecycle management** – AI systems must be managed through their lifecycle, including continuous monitoring, periodic reevaluation, and change control to address data drift, evolving use contexts, and model or configuration changes
- **Risk-based performance assessment** – Validation should consider the complete system, including AI components and human–AI interaction, with evidence proportional to patient, product, and data integrity risk
- **Transparency and information** – Users should receive clear information about AI purpose, intended and context of use, performance, limitations, and interpretability so they can appropriately rely on and challenge outputs.

These principles apply directly to agentic AI platforms whose outputs influence clinical data quality and analysis, even if they do not make direct clinical decisions. They shape both the AI Platform Validation Framework (by defining expectations for platform-level lifecycle) and the AI Agent Lifecycle (by requiring clear intended uses, risk assessments, and oversight for individual agents).

GxP and GAMP 5 AI/ML validation guidance

Updated guidance on AI/ML validation in GxP, such as GAMP 5 Second Edition and Appendix D11, reinforces a shift toward risk-based computer software assurance and lifecycle validation for AI components. Key elements include

- Emphasizing fitness for intended use, patient safety, and data integrity as overarching goals rather than exhaustive testing of all possible behaviors.
- Recognizing that AI introduces opaque model behavior and continuous change, shifting much of the risk from static code to data quality, model configuration, and operating context.
- Advocating proportional validation, where higher-risk uses (for example, those affecting participant safety or critical data integrity) receive more intensive validation and monitoring than low-risk, supportive functions.
- Highlighting supporting processes—Risk management, data governance, and change management—that operate across concept, project, and operation phases of the AI subsystem life cycle.

These positions underpin the AI Platform Validation Framework by guiding when and how to perform platform-level risk assessment, requirements specification, development and design controls, and acceptance activities, and they inform the AI Agent Lifecycle by defining expectations for agent-level risk stratification, guardrails, and continuous oversight

Segmenting uses by risk

Agentic AI platforms often support multiple use cases within the same environment. Building on risk assessment approaches such as FMEA and AI-specific risk matrices, functions can be grouped into tiers such as

- **Low-risk, operational support** (for example, formatting reports, pre-populating non-critical documentation), where unscripted or exploratory testing aligned with intended use may be sufficient, combined with general controls.
- **Medium-risk, data quality and oversight** (for example, suggesting data queries, flagging anomalies, generating draft code), where more structured testing, guardrails, and human-in-the-loop oversight are required.
- **High-risk, direct impact on data integrity or participant safety** (for example, inclusion/exclusion decision support, safety signal triage), where robust, fully scripted testing, stringent monitoring, and close regulatory alignment are needed.

At the platform level, an initial high-level risk assessment informs the depth of validation and the rigor of IQ/OQ and other tests. At the agent level, a risk matrix based on clinical context, Model influence and decision consequence supports consistent classification and determines whether human-in-the-loop or human-on-the-loop oversight is appropriate

ARCHITECTURE OF AGENTIC AI PLATFORMS IN CLINICAL DATA ENVIRONMENTS

Typical components and data flows

A generic agentic AI platform for clinical data environments includes:

- Orchestration layer managing triggers, sequences, and dependencies among agents.
- Agents dedicated to tasks such as data ingestion, mapping, validation, reconciliation, code generation, or documentation.
- Data layer (e.g., CDR/Data Lakehouse etc) providing structured inputs and receiving outputs.
- Observation and logging layer capturing events, metrics, and audit trails.

For validation, this architecture must be understood and documented so that each component's role, interfaces, and failure modes can be assessed.

A generic agentic AI platform for clinical data environments includes:

- **Agents** – Task-focused agents designed to support activities such as trial startup and initiation documentation, structured and unstructured data ingestion, data quality management, aggregation, standardization, validation, code generation, and quality checks or verification. Each agent operates with a clearly defined intended use, level of autonomy, risk category, and lifecycle phase.
- **Orchestration layer** - managing triggers, sequences, planning, and dependencies among agents, implementing the “agentic” behavior of decomposing tasks and coordinating multiple agents within defined policies and guardrails.
- **Data Layer** –This includes systems like clinical data repositories or lakehouse platforms, vector databases, and knowledge graph databases that store both structured and unstructured data. These systems act as the backbone for Data analytics, statistical computing environments (SCEs), and AI or agentic workflows by securely holding the data they use and generate. Depending on how the data is used, especially in regulated contexts, parts of this layer may fall under GxP requirements and 21 CFR Part 11 controls.
- **Observation and logging layer** – capturing events, prompts, inputs, intermediate reasoning steps, outputs, metrics, and audit trails to support traceability, incident investigation, and regulatory inspection.

For validation, this architecture must be understood and documented so that each component's role, interfaces, and potential failure modes can be assessed in relation to intended use and context of use. Clearly documented configurations, architectural decisions, and version-controlled artifacts are also essential to support both platform-level and agent-level lifecycle phases.

Where validation needs to focus

From a GxP standpoint, validation should concentrate on:

- High-impact agent tasks that can alter data content, metadata, or interpretation used in regulated outputs, including automated transformations and anomaly detection that feed clinical analysis or reporting.

- Interfaces between agents and GxP data stores for example, how agents read from and write to the clinical data repository or secured repositories supporting Statistical computing environments and what controls exist at these boundaries.
- Controls and guardrails, such as thresholds, constraints, access restrictions (Role based access control, Study/Program/Project, Data Privacy), policy-based response filtering, and escalation triggers that limit agent behavior to validated scopes.
- Human interaction points, including review, override, and approval steps, which are central to risk mitigation and must be reflected in both platform validation and agent lifecycle design.

Mapping these elements early supports a targeted AI Platform Validation Framework that leverages existing CSV/CSA practices while extending them to multi-agent orchestration and AI-specific hazards, and it provides a concrete basis for the AI Agent Lifecycle Phases

VALIDATION STRATEGY: LIFECYCLE AND DELIVERABLES

Lifecycle stages for AI Platform Validation

Following EMA–FDA principles, GAMP 5 Appendix D11, and the AI Platform Validation Framework, a lifecycle validation strategy for an agentic AI platform typically include

1. Concept Phase

- Intended use of the system** – Document the specific purpose of the AI platform, including functional scope, users, and business impact in clinical data environments.
- Platform high-level risk assessment** – Identify potential risks associated with the platform as a whole and their potential impact on patient safety, data integrity, and regulatory compliance.
- 21 CFR Part 11 applicability** – Evaluate whether the platform falls under electronic records and signatures regulations and determine the need for controls such as audit trails, secure access, and electronic signature compliance.

2. Requirements Phase

- Risk assessment** – Apply a Failure Modes and Effects Analysis (FMEA) or similar approach to prioritize risks associated with the platform's functions: high, medium, or low risk priority, guiding the depth and rigor of validation and testing activities.
- Functional and non-functional requirements** – Specify requirements for performance, security, logging, traceability, and interoperability, ensuring that intended use and context of use are fully captured.
- Validation planning** – Define the validation strategy, including IQ/OQ where applicable, risk-based test coverage, and deliverables such as test plans and summary reports.

3. Development and Design Phase

- Development** – Develop and configure the AI platform following standard software development practices, ensuring documented code, version control, and secure data handling.
- Design and configurations** – Configure the platform to meet functional requirements, document architecture, algorithm selection, data flows, and configurable parameters, and ensure that design reflects risk assessment outcomes.
- Testing (IQ/OQ)** – Conduct Installation Qualification (IQ) to verify that the platform is correctly installed, and Operational Qualification (OQ) to confirm that it performs according to specifications, tailoring scripted vs. exploratory testing to risk priority.

4. Acceptance and Release Phase

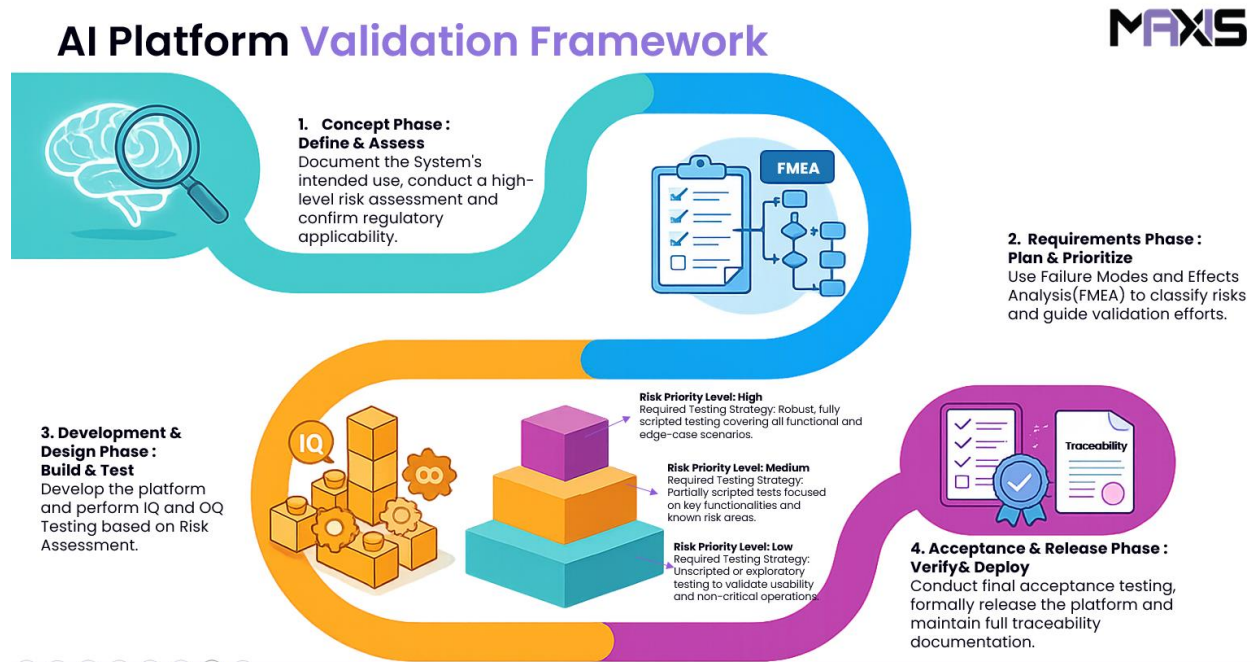
- Acceptance testing** – Verify that the platform meets all specified requirements and functions as intended in the production environment, with traceability back to requirements (intended use) and risk assessments.
- Release** – Formally approve the AI platform for operational use under a quality management system, documenting release notes, configuration settings, and risk mitigation measures.
- Traceability and Validation Summary report** – Maintain traceability between requirements, design, testing, and risk assessment outcomes and prepare a validation summary report that consolidates evidence of fitness for intended use.

5. Operation, monitoring, and change control

- Continuous operation** – Operate the platform under defined procedures, maintaining audit trails for prompts, inputs, intermediate reasoning steps, and outputs.

- b) **Monitoring and drift detection** – Monitor performance and usage metrics and detect data or model drift, with predefined thresholds that trigger investigation or revalidation.
- c) **Change control** – Manage changes to platform configurations, models, and data sources under structured change control within the quality management system, with risk-based planning and documented approvals.

Within this lifecycle, the AI Platform Validation Framework provides the structure for platform-level assurance, while the AI Agent Lifecycle Phases manage the validation, safety, and governance of individual agents running on the platform.



AI Agent Lifecycle Phases for validation, safety, and governance

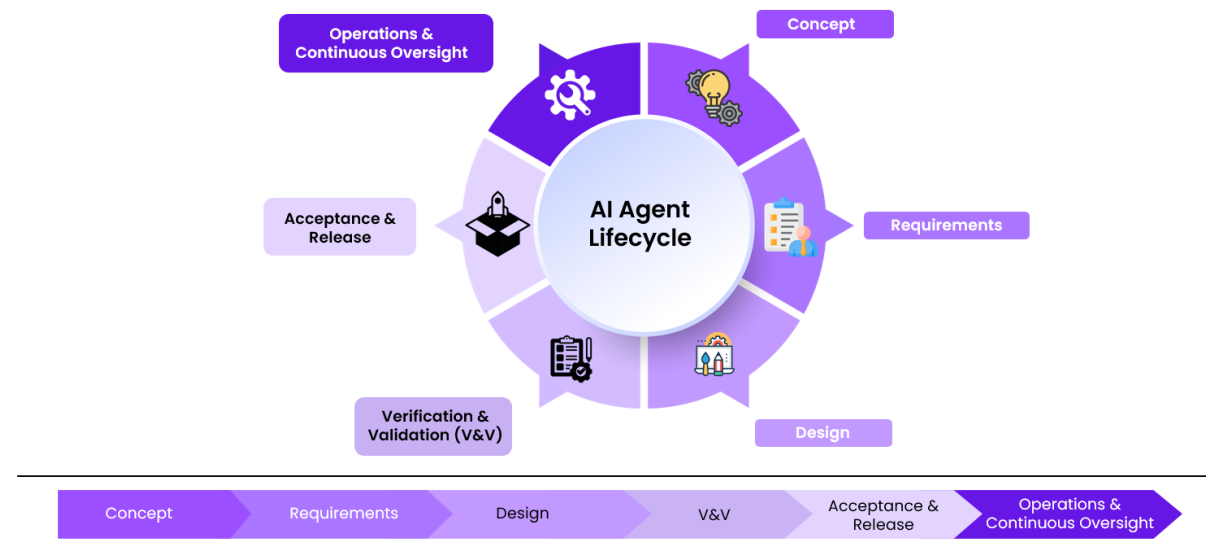
AI Agent Lifecycle Phases can be applied to individual agents whose outputs influence GxP activities:

1. **Concept Phase**
 - a) Define the agent's intended use, clinical context, and GxP impact, including whether it supports quality checks, documentation, or other clinical data processes.
 - b) Assess AI system risk through Model influence, decision consequence using a matrix that combines these dimensions to derive an agent risk category (low, medium, high).
2. **Requirements Phase**
 - a) Define explicit safety, functional, and monitoring requirements, including comprehensive logging to ensure traceability of inputs, outputs, intermediate reasoning, and decision pathways.
 - b) Establish organizational accountability via a RACI matrix or equivalent, clarifying responsibilities for safety, performance monitoring, incident management, and regulatory compliance.
 - c) Specify performance targets (for example, accuracy, latency, robustness) and embed transparency, security, privacy, and guardrail requirements.
3. **Design Phase**
 - a) Translate requirements into architectural and procedural controls, including configuration of guardrails such as input validation, output constraints, policy-based response filtering, and escalation triggers.
 - b) Design monitoring dashboards to surface key performance indicators, drift signals, fairness metrics, security events, and usage patterns, and define incident management workflows.
 - c) Create version-controlled artifacts capturing agent configurations, architectural decisions, model updates, rationale for changes, and sign-offs to support traceability and audits.
 - d) Calibrate human oversight according to risk: human-in-the-loop for high- to medium-risk agents, human-on-the-loop for low-risk agents.
4. **Verification and Validation Phase**

- a) Perform configuration verification, functional testing, and statistical verification to demonstrate that the agent meets its specified requirements and is fit for its intended use.
 - b) Map evaluation metrics such as task completion rate, accuracy, precision, recall, F1 score, and AUC-ROC to risk categories; higher performance corresponds to lower risk and may justify less intensive oversight.
 - c) Conduct independent quality review of traceability matrices, deviation logs, and validation reports, ensuring segregation of duties.
 - d) Perform Explainability, bias, fairness, robustness, and security testing, including stress tests under noisy inputs, edge cases, partial failures, or adversarial prompts, and verify guardrail effectiveness.
5. **Acceptance and Release Phase**
- a) Implement human oversight mechanisms at critical decision points so that high-risk actions remain subject to human review or intervention.
 - b) Govern release through structured change control and authorization within the quality management system, with documented approval, risk acceptance, and readiness assessment.
 - c) Define rollback procedures, escalation paths, and contingency plans to support rapid recovery in the event of post-release issues.
6. **Operations and Continuous Oversight Phase**
- a) Use continuous monitoring mechanisms—including automated logging, alerts, and anomaly detection—to identify performance degradation, drift, security incidents, or ethical concerns.
 - b) Maintain operational controls such as guardrail enforcement, audit trails, and incident management processes, and perform periodic reviews for ongoing compliance and effectiveness.
 - c) Monitor for privacy breaches and security vulnerabilities to ensure that the agent remains resilient to evolving threats and regulatory expectations.

MAXIS Governance Framework for the AI Agent Lifecycle

End to End Strategies for Validation, Safety and GxP Compliance



Validation documentation

Key validation deliverables tailored to Agentic AI include:

- Platform-level validation plan, risk assessment, requirements specifications, IQ/OQ protocols and reports, traceability matrices, and validation summary report.
- Agent-level risk assessments (including the model influence vs decision consequence), requirements, design specifications, V&V plans and reports, metrics and thresholds, and oversight procedures.
- Logs and audit trails capturing prompts, inputs, intermediate reasoning, outputs, and human decisions, treated as records for inspection and incident investigation.
- Governance artifacts such as RACI matrices, change control records, and periodic review minutes documenting lifecycle oversight.

These align traditional CSV artefacts with the added complexity of AI while retaining a risk-based, pragmatic stance

DATA STRATEGY, DRIFT, AND CONTINUOUS MONITORING

Data strategy for AI validation

AI/ML validation guidance emphasizes that data quality, representativeness, and alignment with context of use are central to validation. For agentic AI platforms in clinical data environments, this translates to

- Documenting training, reference, and evaluation datasets for AI components and agents, including sources, selection criteria, and relevant characteristics.
- Defining intended operating ranges such as trial types, therapeutic areas, data volume profiles, and coding systems, and ensuring that validation data reflect these ranges.
- Characterizing known limitations in performance on rare edge cases or under-represented populations and communicating these limitations to users and oversight bodies.

These data-centric aspects must be reflected in both the AI Platform Validation Framework (for example, in requirements and testing phases) and the AI Agent Lifecycle (for example, in design, V&V, and continuous oversight).

Monitoring for drift and performance degradation.

Because AI components can be sensitive to changes in input data distribution or operational context, continuous monitoring is vital. For the platform and its agents, monitoring might include

- Performance metrics on key tasks (for example, anomaly detection precision, false positive and false negative rates, documentation accuracy, task completion rates) tracked over time.
- Volume and nature of human overrides and escalations, especially sudden increases that may indicate degraded behavior or insufficient guardrails.
- Changes in source data characteristics, such as new coding standards or data structures, that could affect assumptions and require retraining or revalidation.

Monitoring results should feed into risk-based decisions about when to adjust thresholds, retrain models, refine prompts or agent policies, or perform partial revalidation under change control, supporting the Operations and Continuous Oversight Phase of the agent lifecycle and the operation/monitoring phase of the platform framework

HUMAN-AI INTERACTION, GUARDRAILS, AND OVERSIGHT

Designing guardrails

Guardrails constrain AI behavior within acceptable bounds and are a central design element in the AI Agent Lifecycle. For clinical data environments, examples include

- Configurable thresholds for flagging anomalies or suggesting changes, tuned according to risk category and validated performance.
- Rules that prevent agents from directly modifying critical datasets without human confirmation, ensuring that high-impact actions remain under human control.
- Scope limitations that keep agents within defined tasks and contexts unless expanded through deliberate, validated changes in intended use and risk assessment.

These guardrails reduce risk and make behavior more predictable, supporting both platform-level validation and ongoing operation in line with GxP expectations.

Human oversight models and Governance

Regulatory principles call for explicit consideration of human-AI interaction and oversight. In practice, this is implemented through the AI Agent Lifecycle and platform governance by

- Defining oversight models aligned to risk: human-in-the-loop mechanisms for high- and medium-risk agents and human-on-the-loop mechanisms for low-risk agents.
- Assigning responsibilities through RACI or similar models, clarifying who is responsible and accountable for reviewing outputs, monitoring performance, managing incidents, and approving changes.
- Recording decisions and rationales when humans accept, modify, or override AI outputs, and maintaining audit trails suitable for inspection.

Validation should include testing these interaction flows and verifying that oversight remains effective as data, models, and operational contexts evolve. Governance processes such as structured change control, periodic reviews, and AI governance committee oversight help ensure that both the platform and its agents remain within acceptable risk and compliance thresholds

ILLUSTRATIVE SCENARIOS AND METRICS

Scenario 1: Agentic AI for data anomaly detection

An agentic platform monitors integrated clinical datasets, flagging potential data anomalies such as out-of-range values, missing visits, or inconsistent dates etc. Within the AI Platform Validation Framework, this function is scoped and risk-assessed at the platform level, and within the AI Agent Lifecycle it is treated as a medium-risk agent whose intended use, Model influence and decision consequence are explicitly documented

Validation focuses on:

- Testing anomaly rules and AI detection performance on representative datasets.
- Measuring false positive and false negative rates.
- Verifying that flagged anomalies appear correctly in review tools and that resolution paths are auditable.

Monitoring in the Operations and Continuous Oversight Phase tracks performance metrics and override rates and triggers change control if performance falls below defined thresholds, ensuring ongoing safety and data integrity

Scenario 2: Agentic Study Startup and Initiation

An agentic AI platform accelerates study start-up by coordinating protocol development through Regulatory submission and study preparation using a hierarchical Supervisor–Worker agent pattern. Supervisor Agent interprets high-level study objectives in natural language, decomposes them into tasks, and orchestrates specialized agents for protocol drafting, document distribution, Regulatory package assembly and initial data-integration setup.

Each Worker Agent has a defined intended use and risk category, is aligned with ICH-GCP and 21 CFR Part 11 expectations, and is augmented with sponsor-specific protocols, SOPs, and templates via retrieval-augmented generation. Independent Evaluator Agents review key outputs for accuracy, completeness, and regulatory alignment before release, so only artifacts meeting predefined criteria reach human reviewers. All prompts, intermediate steps, outputs, evaluations, and human decisions are logged, enabling risk-based validation of the platform orchestration and agent lifecycles, while providing an inspection-ready audit trail that supports faster, more consistent Regulatory submissions.

PRACTICAL RECOMMENDATIONS AND LESSONS LEARNED

Drawing from current guidance and industry narratives, several recommendations emerge for applying the AI Platform Validation Framework and AI Agent Lifecycle Phases in GxP clinical data environments.

- Anchor platform validation in intended use and risk – Clearly define platform scope, users, and 21 CFR Part 11 applicability and perform high-level risk assessment to scale validation effort proportionally.
- Use AI-specific risk matrices for agents – Combine clinical context, model influence and decision consequence into a risk matrix to classify agents and determine the appropriate degree of guardrails, testing, and oversight.
- Integrate data strategy and monitoring into lifecycle phases – Treat data quality, representativeness, and drift as first-class concerns in requirements, design, V&V, and continuous oversight for both platform and agents.
- Prioritize explainability, guardrails, and human control – Design agents and workflows so that reviewers can understand and challenge outputs, enforce guardrails, and maintain effective human-in/on-the-loop oversight.
- Embed governance into the quality system – Align platform and agent lifecycle activities with quality management processes, including change control, periodic review, and AI governance oversight, to ensure that innovation is balanced with compliance.

These recommendations help organizations modernize validation practices without overstating the maturity of agentic AI in GxP and provide a concrete path to managing both the platform and agent lifecycles over time.

CONCLUSION

AI Platform Validation Framework presented in this paper offers a lifecycle structure—spanning concept, requirements, development and design, acceptance and release, and ongoing operation and monitoring—that enables organizations to validate an entire agentic AI platform as a GxP-relevant computerized system. By grounding activities in intended use, risk assessment, and proportional testing, the framework helps ensure platform-level fitness for purpose, data integrity, and inspection readiness.

Complementing this, the AI Agent Lifecycle Phases for Validation, Safety, and Governance provide a focused view on individual agents operating within the platform. Through explicit phases covering concept, requirements, design, verification and validation, acceptance and release, and operations with continuous oversight, the lifecycle links risk stratification, guardrails, human-in/on-the-loop oversight, and ongoing monitoring to each agent's role and impact.

Together, these two contributions offer a coherent, regulator-aligned foundation for adopting agentic AI in clinical data environments. They enable organizations to introduce hierarchical, multi-agent architectures in a way that balances innovation with robust validation, clear accountability, and sustainable lifecycle governance.

REFERENCES

- 1 FDA. *Considerations for the Use of Artificial Intelligence To Support Regulatory Decision-Making for Drug and Biological Products*. <https://www.fda.gov/media/184830/download>
- 2 European Biotechnology Magazine. *EMA and FDA join forces: Aligned principles on AI use in drug development*. <https://european-biotechnology.com/latest-news/ema-and-fda-join-forces-aligned-principles-on-ai-use-in-drug-development/>
- 3 PharmTech. *EMA and FDA Collaborate on Framework for AI Use in Drug Development*. <https://www.pharmtech.com/view/ema-and-fda-collaborate-on-framework-for-ai-use-in-drug-development>
- 4 IntuitionLabs. *AI/ML Validation in GxP: A Guide to GAMP 5 Appendix D11*. <https://intuitionlabs.ai/pdfs/ai-ml-validation-in-gxp-a-guide-to-gamp-5-appendix-d11.pdf>
- 5 arXiv. *AI Agents vs. Agentic AI: A Conceptual Taxonomy, Applications and Challenges* (arXiv:2505.10468). <https://arxiv.org/pdf/2505.10468>
- 6 MMS Holdings. *The FDA's AI Guidance and its Seven Steps into the Future of Drug Development*. <https://mmsholdings.com/regulatory/the-fdas-ai-guidance-and-its-seven-steps-into-the-future-of-drug-development/>
- 7 <https://globalforum.diaglobal.org/issue/november-2025/the-next-ai-revolution-computer-system-validation/>

RECOMMENDED READING

Recommended reading lists go after your acknowledgments. This section is not required.

CONTACT INFORMATION

Contact the author at:

Author Name: RAJESH HAGALWADI

Company: Maxis AI

Address: 510 Thornall Street

Work Phone: 732 494 2005

Email: rajesh@maxisit.com

Website: www.maxisit.com

Brand and product names are trademarks of their respective companies.