

A FAIR Data Foundation for AI-Supported Regulatory Decision-Making in Drug and Biological Products

Colin Wood



Author



Colin Wood

Strategy & Solutions Leader,
Life Sciences

A seasoned Life Sciences data and strategy leader with 34+ years of experience across R&D and enterprise information management. He has led large-scale initiatives in Master Data, Metadata, Data Quality, and FAIR data principles, building scalable foundations for innovation.

A recognized thought leader in pharmaceutical data, Colin has deep expertise in standards such as IDMP, HL7 FHIR, and USDM. He has helped organizations modernize data ecosystems and accelerate R&D transformation.



Content

1	Introduction	03
2	Disclaimer	04
3	AI Governance & Data Governance Must Work Together	05
4	Implications for AI Agents	06
5	New opportunities offered by AI Agents	07
6	AI agents also come with challenges	08
7	Ensuring persistence through GUPRIs	09
8	Structuring and Governing Decision Data	11
9	Appendix	13
10	Glossary	15



1. Introduction

The FDA's 2025 draft **Industry Guidance for the Use of Artificial Intelligence to Support Regulatory Decision-Making for Drug and Biological Products** marks a significant shift in how AI will contribute to regulatory processes. The guidance emphasizes transparency, traceability and governance, including the use of **Credibility Plans** and clear documentation of **Context of Use**. The potential scope of AI across the drug product lifecycle is vast. AI may generate synthetic control arm data, predict toxicology outcomes, estimate stability profiles, or support decisions in study design, safety surveillance, and many other regulated decisions. As a result, AI has the potential to influence nearly every facet of pharmaceutical development. For Chief Data Officers, Chief AI Officers, Heads of Regulatory Affairs, and Quality leaders, this guidance implies a shift from AI-as-a-tool to AI-as-a-regulated component of the product lifecycle.

Other recently released documents relating to the use of AI within the industry include:

- **Guiding Principles of Good AI Practice in Drug Development (2026)**, FDA/EMA
- **EU GMP Annex 22** (Jul 2025), specific to GMP manufacturing

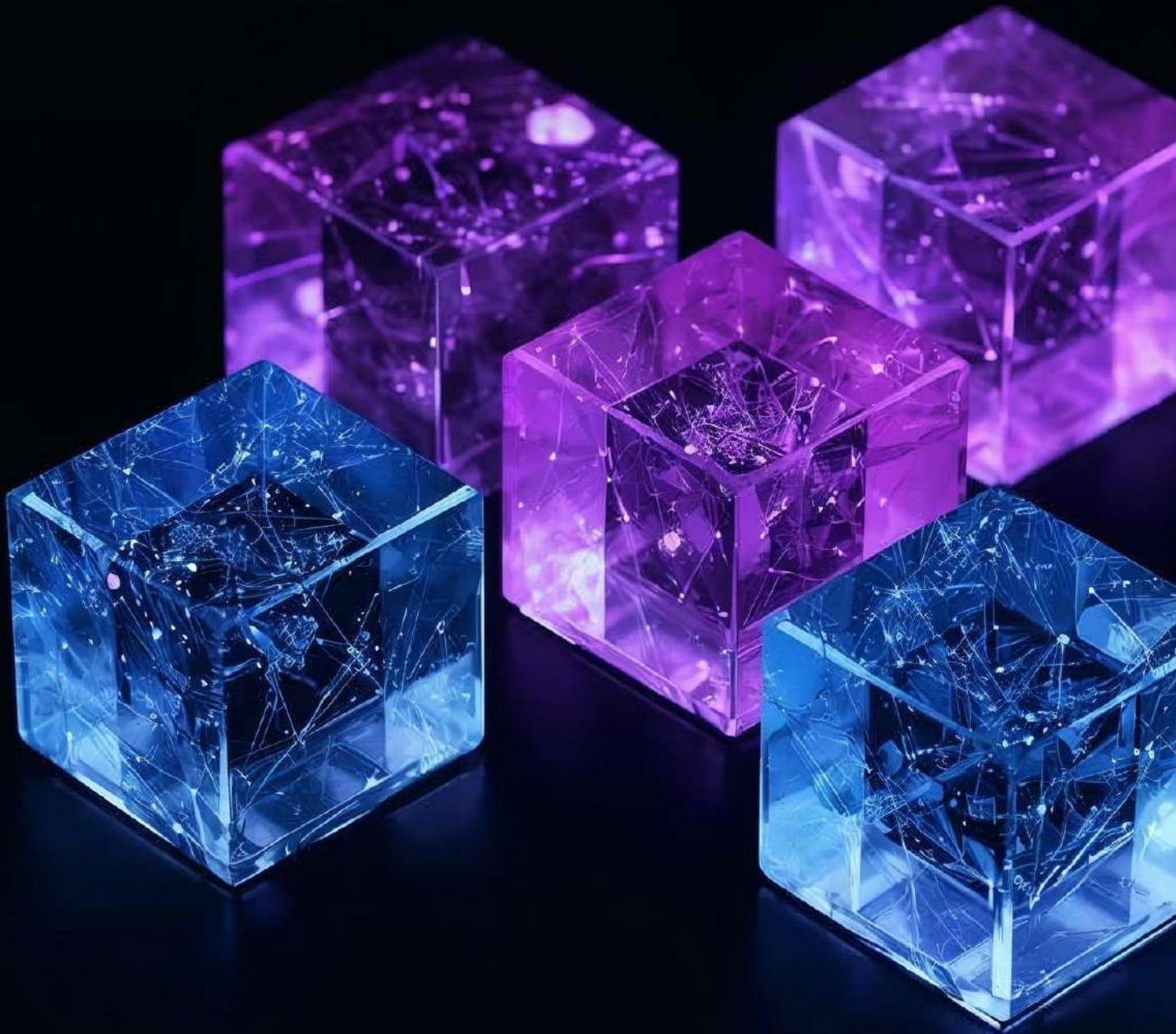
These give a consistent picture of the need to carefully govern the use of AI.

Operating in a regulated environment introduces additional obligations around data integrity (ALCOA+), long-term retention of models and datasets, and compliance with 21 CFR Part 11. Full exploitation of the potential of AI also requires organisations to move away from governance-by-documentation and instead build FAIR (Findable, Accessible, Interoperable, Reusable) data infrastructures that enable automated introspection and long-term reproducibility.



2. Disclaimer

The concepts presented here are exploratory and should not be interpreted as definitive regulatory guidance. Organisations should carefully evaluate regulatory requirements and AI governance frameworks before implementing any solutions.





3. AI Governance and Data Governance Must Work Together

Most organisations now recognise the need for AI governance, including compliance with the EU AI Act and ethical AI principles. The FDA guidance aligns and extends this requiring the identification of AI Systems, assessment of AI Model Risks, and the use of Credibility Plans supported by AI Experiments. Failure to coordinate AI and data governance risks long-term untraceability, increased inspection burden, and re-work during regulatory submissions or audits. Integrated AI and data governance enable faster, auditable, and repeatable regulatory submissions by ensuring that every AI-influenced decision can be reconstructed on demand.

However, effective regulatory use of AI also demands robust data governance. Long term traceability requires persistent identifiers for AI Systems, AI Models and versions, Datasets and Dataset versions, AI Experiments and AI Agents.





4. Implications for AI Agents

The FDA guidance predates the rapid rise of LLM-driven AI agents capable of performing tasks, interacting with other agents, and making or recommending decisions. Whilst not explicitly addressed, several principles can be inferred.

- Each AI agent should be supported by a Credibility Plan demonstrating trustworthiness and fitness for purpose.
- The AI System must define a clear Context of Use for each agent, including required human oversight.
- AI Model Risk frameworks should be extended to the agent level, enabling risk attribution to individual agents.
- Keep Human-in-the Loop for critical decisions that may impact patient safety or product quality.

The level of scrutiny should scale with regulatory risk, especially for decisions impacting patient safety or product quality.

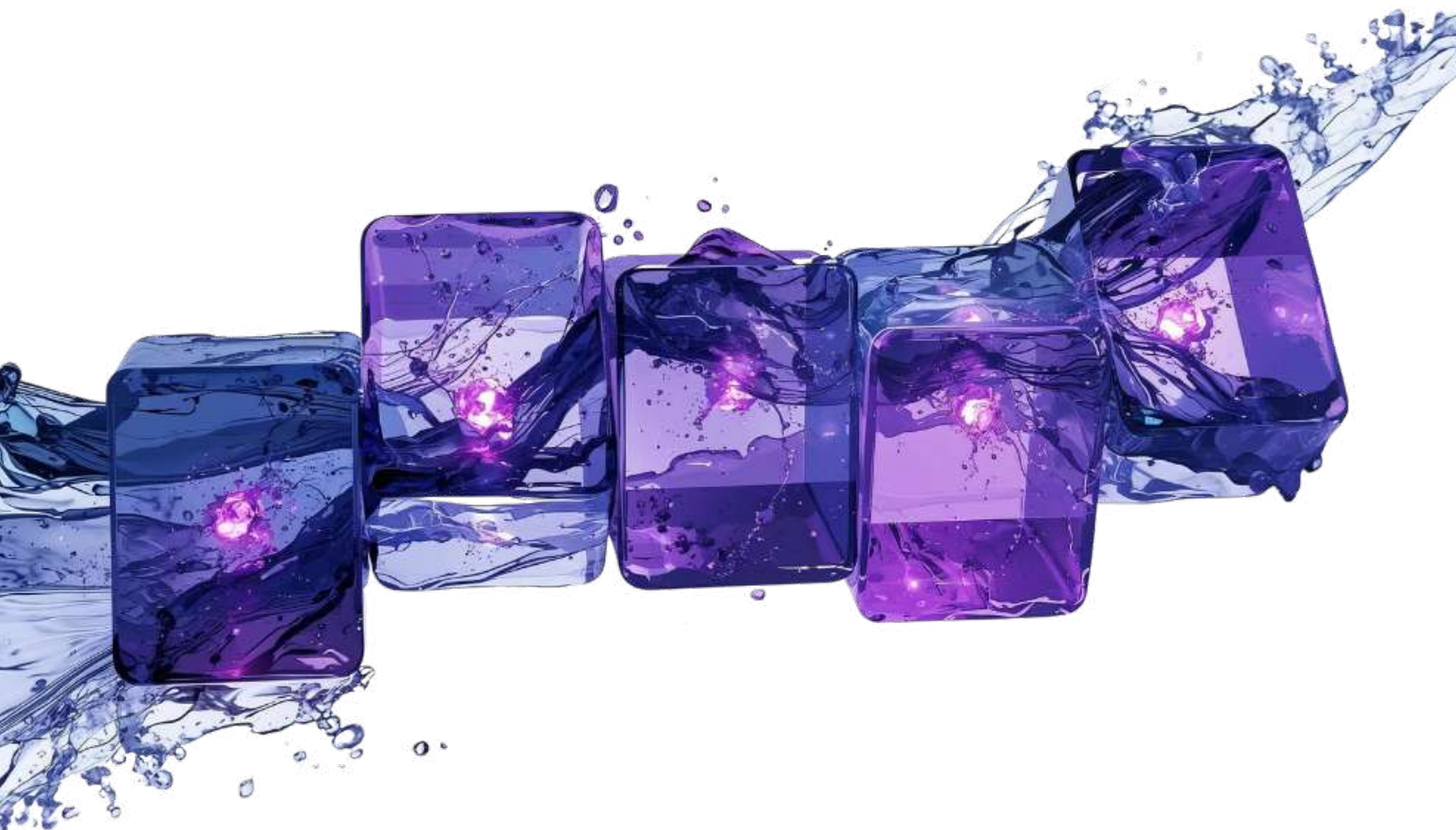
The EU GMP Annex 22 goes further and states that generative AI and Large Language Models should **not** be used in critical GMP applications. While this doesn't directly impact GCP or GLP applications, it does illustrate the potential risk relating to the use of generative AI. This restrictive stance may evolve as regulators gain experience with auditable, explainable, and FAIR grounded AI systems. Therefore, early investment in traceability and governance can position companies to meet future expectations.



5. New opportunities offered by AI Agents

Conversational AI agents introduce capabilities that go beyond process automation. Because agents can be defined with persistent knowledge bases, they enable unprecedented transparency. This could transform regulatory inspections by enabling inspectors to interrogate the exact agent involved in a decision, even years later.

AI agents also enable introspection. When new information emerges (e.g., unexpected study results), agents could identify which past decisions might be affected and summarise the implications. This creates a dynamic, self-updating regulatory knowledge environment.

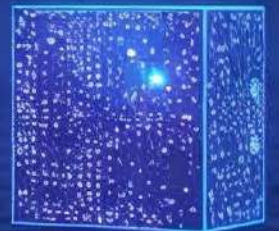
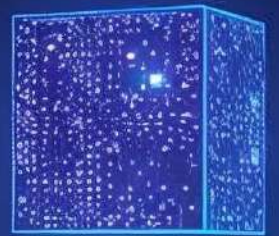
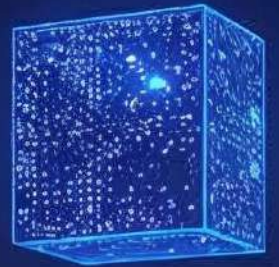




6. AI agents also come with challenges

The rapid evolution of AI platforms means organisations may cycle through multiple generations of models and agents within a few years. This creates a risk that historical decision records will reference agent identifiers that no longer resolve to an operational system. While documentation may exist, regulators increasingly expect automated traceability and reproducibility. Platform dependent identifiers that cannot be resolved years later may force organisations to manually reconstruct AI-related decision histories at audit or inspection, significantly increasing cost and risk. This is also a core concern for the full exploitation of the potential of AI.

This makes persistent, resolvable identifiers essential.





7. Ensuring persistence through GUPRIs

FAIR practices teach us that identifiers are foundational and should be designed to be Globally Unique, Persistent, and Resolvable Identifiers (GUPRI). Many of the practices for this were outlined in a 2017 paper [Identifiers for the 21st Century](#).

Any AI developer will rightly claim that existing AI Platforms are already able to generate unique identity for the range of data entities that support AI activities, including AI Models, AI Model Versions, AI Agents and AI Experiments. Such identifiers are locally generated and use a variety of platform-specific techniques, which could include the use of UUIDs, sha256 hash codes, sequences (e.g. agent_1234) or simple version codes (e.g. “Model name: Model Version number”).

This gives a huge variety of potential identifier formats, which can of course change after migration to a new AI platform. Furthermore, not all the identifiers are immutable. The identifier formats that use a Model Name, for example, can change following a rename of a model. Indeed, an analysis of common AI registries shows, these identifiers are often:

- Not globally unique
- Not persistent across migrations
- Not resolvable via standard protocols
- Not accessible once archived or deleted

In other words, they fail the FAIR test. Some AI platforms may be better than others, but they all fail on at least the last two of these criteria. The good news is that organisations do not need to replace existing identifiers. They can **wrap** them in a persistent, resolvable namespace. For example, if the original model identifier is 12345, it can be wrapped to form a GUPRI of the form:

<https://pid.company.com/ai-model-namespace/12345>

(Don't click on the link, it doesn't exist)



This identifier isn't configured yet. However, with the right infrastructure this can be configured to act as a GUPRI. For an example of a public domain GUPRI see: http://purl.obolibrary.org/obo/NCIT_C2985. Clicking on this identifier will rewrite the URL and will resolve to the NCI Thesaurus.

Before: Datasets and content contained static identifiers that could not be automatically linked to metadata or provenance.

After: Identifiers automatically link to metadata, provenance and even archived content, giving full traceability and enabling introspection by AI agents. The links can be secured to ensure individuals and AI agents can only access data appropriate to their role.

This approach:

- Preserves existing platform identifiers
- Adds global uniqueness
- Applies the Registry Style of MDM, allowing all identifiers to be governed and traced
- Enables resolution to metadata or archived content
- Enables resolution to changed platform identifiers following migration
- Supports compliance with FAIR principle A2 (metadata accessible even when data is gone)

This is the most practical and scalable path to FAIR maturity. The approach can be supported using commercial solutions ([Accurids FAIR Data Registry](#)) and may be built using open-source tooling (see [FAIR cookbook](#)).

With this in place, organisations can design governed behaviours for retired agents, such as redirecting to archives, replacement agents, or deprecation notices.

While this may not be seen as a pressing issue since companies are still experimenting with AI or are working from a single AI stack, this may lead to long-term issues tracing and interacting with AI supported decisions over a period of years. The time to resolve this issue is **before** the identifiers are deeply embedded in regulatory data, which means this should be seen as a priority **now**.

8. Structuring and Governing Decision Data



A decision should be formally recorded when it affects product quality, patient safety, data integrity, regulatory compliance, or validated processes. AI will increasingly influence such decisions, whether through predictive modelling, automated review, or LLM-assisted SOP development.

To make decision records FAIR and suitable for retrospective analysis, organisations should:

- Assign each decision a persistent, resolvable identifier
- Define a canonical decision schema with minimum metadata requirements
- Use provenance standards (e.g., PROV-O) to capture human and AI contributors, inputs, and dependencies
- Provide templates for decisions captured in unstructured formats such as meeting minutes

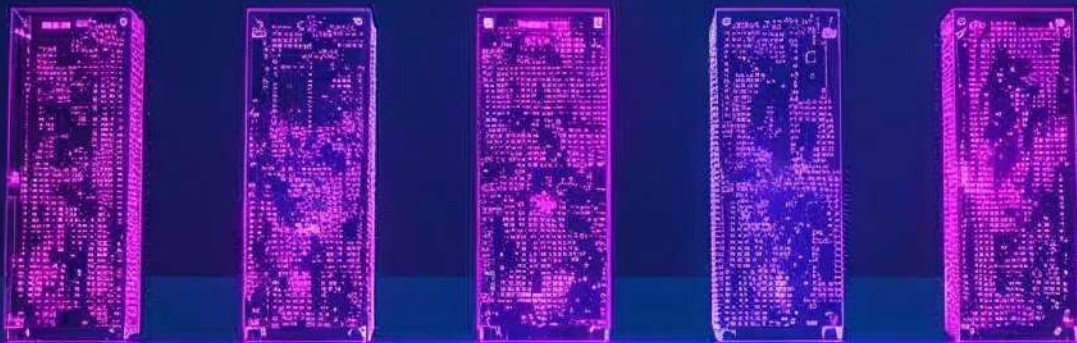
This foundation enables the development of Context Graphs, which are connected representations of decisions, events, entities, and evidence. While the technology is new, the underlying need for consistent data standards remains unchanged.



Next Steps

AI-supported regulatory decision-making presents both opportunity and complexity. Realising its potential requires coordinated investment in AI governance, data governance, and persistent identification strategies such as GUPRIs. These foundations will enable automated introspection, long-term traceability, and context-aware AI systems that improve drug development and regulatory compliance.

Therefore, organisations should treat FAIR-compliant AI infrastructures as a strategic programme, not a side project to reduce long-term regulatory risk, accelerate cycle times, and position themselves as a leader in AI-enabled drug development.

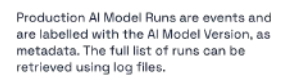


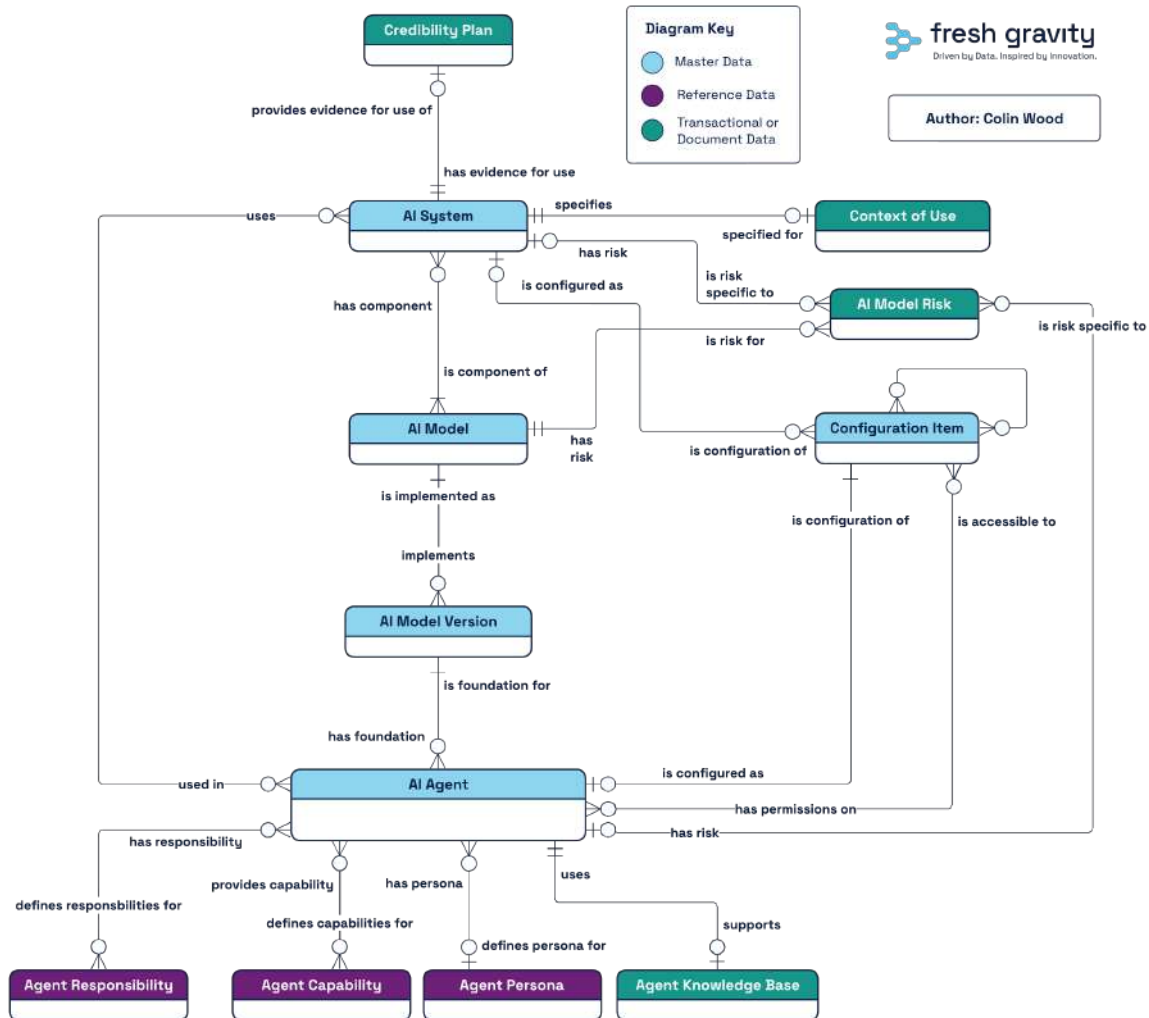


When planning this document, I searched for a relevant domain model or ontology that would define the concepts within scope of AI Models and AI agents and supporting governance activities. As I couldn't find anything that represented the full domain, I built a simple conceptual data model, shown below, using ER notation and defined a basic glossary.

The Credibility plan will reference content from a series of experiments used to build and test the AI model. There may also be experiments used to give specific evidence supporting the use of the AI system within the defined context of use.

An AI Model Experiment may use a single model or a series of models in a pipeline. Configuration of an AI Model, as a new version, would also be defined in a series of experiments.





The colour coding on the models is intended to help understand the impact on FAIR data.

- **Sky Blue - Master Data Entities**

These require persistent identifiers and a strategy for maintaining a single source of identity. Examples include AI Model, AI Model Version, Dataset, and Dataset Version. Although datasets are not traditionally treated as master data, in this context they are critical because they may be reused across multiple models and must be precisely referenced.

- **Purple - Reference Data**

These are candidates for controlled vocabularies, taxonomies, or ontologies. Examples include AI Model Type, Experiment Type.

- **Green - Transactional or Document Data**

These represent documented content such as Credibility Plans, Context of Use.

In addition to these, this type of model can also be used to plan data standards and the creation or use of ontologies. That's a potential topic for a future discussion.



10. Glossary

The following table gives a glossary supporting the definition of each of the entities within this model. These have largely been sourced from a series of external sources, including ontologies, that relate to the definition and management of AI models. While these form a basic glossary, they should not be regarded as definitive.

Entity	Definition
Agent Capability	An ability or an action that an AI agent can perform to solve problems, interact with systems, or complete tasks.
Agent Knowledge Base	Information repository, specific to an AI agent, that an AI agent consults to provide accurate answers and perform its task. This may also include long-term memory, storing conversations or outputs to ensure consistent responses.
Agent Persona	The defined character, role, behaviour style, and communication pattern an AI agent adopts when interacting with users.
Agent Responsibility	A task or obligation that an AI agent must handle to perform its role. This describes what the agent is responsible for achieving or maintaining, rather than just the individual actions it can perform.
AI Agent	An AI Agent is an intelligent system that interprets inputs, makes decisions, and performs actions to accomplish defined goals. AI agents utilise LLM (Large Language Models) to support reasoning and understanding, but add additional tools, memory and workflows to perform tasks and take actions.
AI Deployment	The installation and activation of a specific AI Model Version on a designated compute environment for operational use. Deployments are typically recorded in a Configuration Management Database (CMDB) to ensure traceability, reproducibility, and lifecycle control.
AI Model	A mathematical or computational construct that generates predictions, inferences, classifications, or other outputs based on input data (AIRO). It may use statistical, machine learning, or deep learning techniques.
AI Model Experiment	A structured, repeatable investigation designed to test hypotheses about data, model architectures, parameters, or configurations, with the goal of evaluating their impact on model performance, behaviour, or risk.
AI Model Type	A classification describing the underlying architecture or methodological approach of an AI Model, such as a Convolutional Neural Network, Gradient Boosted Tree, or Transformer.
AI Model Risk	The potential for an AI Model's outputs to contribute to incorrect, unsafe, or biased decisions that may result in adverse outcomes for patients, products, or regulatory processes.



AI Model Version	A controlled, immutable release of an AI Model with defined training datasets, hyperparameters, and executable artefacts. Each version must be uniquely identifiable and reproducible.
AI Model Run	A discrete execution event of a deployed AI Model Version that generates predictive outputs. Each run is logged with identifiers for the model, version, input data, and compute environment to support traceability and auditability.
AI Platform	An integrated software environment that supports the end-to-end lifecycle of AI models, including data preparation, experimentation, training, validation, deployment, monitoring, and governance.
AI System	A machine-based system capable of operating with varying levels of autonomy, adapting after deployment, and generating outputs such as predictions or decisions based on received inputs (EU AI Act). An AI System may incorporate one or more AI Models as components.
Configuration Item	A component of an IT environment that must be tracked and managed to deliver a service. In the AI context, this typically includes physical or virtual compute environments where AI Model Versions are deployed and executed.
Context of Use	A description of the specific role, scope, and conditions under which an AI Model is applied to address a Question of Interest. It includes how outputs will be used, what additional information or human review is required, and any constraints or assumptions relevant to decision-making.
Credibility Plan	A structured, risk-based framework that defines the evidence, validation activities, and documentation required to demonstrate that an AI Model is trustworthy, reliable, and appropriate for its intended use.
Dataset	A curated collection of data used for training, validating, or testing AI Models. A dataset may include raw data, processed data, annotations, or engineered features.
Dataset Version	A specific, immutable snapshot of a dataset at a defined point in time, preserved to ensure reproducibility of model training, evaluation, and regulatory review.
Experiment Type	A classification describing the nature of an experiment, such as an AI Model Training Experiment or Hyperparameter Tuning Experiment. In an enterprise context, this should align with a broader taxonomy or ontology to support consistent access to experimental data.
Feature Type	A measurable property, variable, or characteristic derived from a dataset and used as an input to an AI Model.
Parameter Type	A variable that defines or influences the behaviour of an AI Model. This includes trainable parameters (e.g., weights, biases) and nontrainable parameters (e.g., hyperparameters) that shape model performance.
Question of Interest	A clearly defined question, decision, or problem that an AI Model is intended to address. It establishes the purpose and scope of model use within a regulated or scientific context.



About Fresh Gravity

Founded in 2015, Fresh Gravity is an innovation-driven Data and Analytics consulting firm that helps businesses make data-driven decisions. We are driven by data and its potential as an asset to drive business growth and efficiency. We are passionate innovators who solve clients' business problems by applying best-in-class data and analytics solutions. For more information, contact us at info@freshgravity.com.

