

A Life-Cycle Approach to AI Governance in Medical Device Development

Miguel Azevedo, PhD

Qity BV, Belgium
miguel@qity.be

Abstract—Medical devices increasingly incorporate AI-enabled functions whose performance depends on development and evaluation data, statistical modelling, deployment conditions, human interaction and continuing control of change. These dependencies create a governance problem that cannot be resolved by treating an AI Model as an ordinary software component or by adding a one-time AI assessment to the product file. At any point in the life cycle, a manufacturer should be able to determine what system is being governed; which uses are authorised; which model and data baselines support each use; which assumptions, risks and limitations were accepted; what evidence supported approval; and which changes or operating signals require those decisions to be revisited. This paper proposes a vendor-neutral life-cycle architecture for an Artificial Intelligence Management System in medical device development. The architecture extends from initial concept and product-realisation planning through design, development, release, deployment, post-market surveillance, modification, retirement, and the retention and continued retrievability of historical records. Its objectives are to establish accountable decision rights; align AI development with the intended medical purpose and organisational objectives; maintain technical and evidential traceability; identify and treat wider AI and medical-device safety risks; control the reuse and modification of models and data; monitor continuing performance and impact; and enable timely restriction, correction, rollback or withdrawal. The architecture represents the AI System, AI Use Case, AI Model, Dataset, Dataset-Model Association and AI Risk as distinct but related controlled entities. Semantic relationships and governance gates connect these entities to the established quality system without duplicating its authoritative records. Operational records for monitoring Signals, AI Reviews and Planned Modifications act upon that configuration. A Controlled Evolution capability defines, before change is attempted, the objectives, permitted boundaries, methods, evidence, reassessment triggers and decision authority through which the system may evolve. In the United States, those records can be assembled into, and used to execute, an authorised Predetermined Change Control Plan without allowing the regulatory document to become detached from the system and evidence it describes. The paper’s central contribution is a mechanism for decision continuity: every approval remains connected to the configuration, context, assumptions and evidence on which it was based. This allows changes and post-market information to be traced to the decisions they may invalidate. A coordinated risk representation also connects broader AI effects to device-safety analysis without treating the two as equivalent. The resulting model is intended to make AI-enabled medical-device governance more reviewable, responsive and proportionate, and can be adapted to other regulated sectors by replacing the medical-device control layer with the relevant sector authority.

Keywords—artificial intelligence; AI-enabled medical devices; artificial intelligence management system; life-cycle governance; technical traceability; quality management; risk management; controlled

evolution; predetermined change control plan; post-market monitoring; machine learning

I. INTRODUCTION

Medical-device manufacturers are increasingly incorporating AI-enabled functions into screening, diagnosis, prognosis, prioritisation, treatment planning, workflow support and patient monitoring. In some products, the intended medical functionality depends materially on AI rather than merely being supplemented by it. Such products are sometimes described commercially as *AI-native*. The expression is useful as shorthand, but it is not used here as a regulatory classification. *AI-enabled* is the more precise general term because the governance obligations arise from the function, intended purpose and risk of the actual system, not from how the product is branded.

AI-enabled functions introduce dependencies that are not fully described by source code or software version. Behaviour may depend on the provenance and composition of training data, the labelling process, the selected model and hyperparameters, the separation of development and evaluation data, the input distribution at a deployment site, the design of human oversight and the way an output is interpreted in a clinical workflow. A model can remain technically unchanged while its performance or consequences change because the patient population, acquisition equipment, clinical pathway or user behaviour has changed. Conversely, a retrained model can preserve the same user interface while altering the product’s statistical behaviour. Research on clinical translation, dataset shift, socio-technical implementation and hidden machine-learning dependencies repeatedly identifies these conditions as barriers to safe and effective adoption [21,24,25,29,30,40-42].

The resulting governance problem extends across the complete device life cycle. It begins when a proposed AI-enabled function is framed during concept and product-realisation planning, because the intended medical purpose, affected persons, expected benefit and acceptable uncertainty determine what data, development controls and evidence will be needed. It continues through design and development planning, data acquisition, model development, verification, validation, transfer, release and deployment. After release, the same governance context is needed to interpret complaints, incidents, performance trends, drift, changes in clinical practice, supplier changes and proposed retraining. Retirement ends authorised use, but it does not end accountability: the approved configuration, decisions, risk

analyses and supporting evidence must remain retrievable for safety review, field action, regulatory enquiry and the applicable retention period.

An AI management approach for medical devices should therefore achieve more than policy compliance. It should establish accountability for each AI-enabled system and authorised use; translate organisational and regulatory objectives into controlled technical decisions; maintain traceability among intended purpose, models, data, risks, controls and evidence; ensure that AI-related effects are assessed at the appropriate organisational, societal and product-safety levels; control suppliers, reuse, configuration and change; verify that approved assumptions remain valid in operation; and provide defined routes for correction, restriction, suspension, rollback, resumption and retirement. These objectives correspond closely with the life-cycle, quality-integration and continuing-oversight needs identified across academic, health-system and regulatory literature [21-30,43-50].

An established medical-device quality management system already provides the authoritative processes for design control, software life-cycle evidence, supplier control, clinical and performance evaluation, release, complaint handling, corrective and preventive action, and post-market activities. The practical gap is one of AI-specific granularity and continuity. Conventional product records may not preserve the distinction between the complete AI System, a particular AI Model, the Dataset versions used for different purposes, the clinical context in which outputs acquire consequence, and the risks created by relationships among those elements. This is not an argument that the quality system is deficient. It is an argument that AI-specific governance must make those objects and dependencies explicit while continuing to use the quality system as the authority for regulated product evidence.

The converse failure occurs when an Artificial Intelligence Management System (AIMS) is implemented as a separate compliance programme. Policies, principles, impact assessments and inventories may exist, yet remain detached from design inputs, the medical-device risk-management file, release evidence and post-market processes. The organisation then maintains two representations of the same product: one for AI governance and another for regulated development. This duplication makes inconsistency likely and obscures causality when a model, dataset, deployment context or monitoring signal changes. Prior proposals to integrate health-AI governance into established quality infrastructure directly support avoiding this parallel-system failure [21,22].

No single instrument is designed to resolve the whole problem. ISO/IEC 42001:2023 establishes an organisational AIMS; ISO/IEC 23894:2023 supplies a more developed discipline for AI-related risk; ISO/IEC 23053:2022 provides a technical system model and shared machine-learning vocabulary; EN 18286:2026 translates EU AI regulatory quality obligations into product-focused processes; and ISO/TS 24971-2:2026 applies established medical-device risk management to machine-learning-specific sources of harm [1-7]. ISO 13485, ISO 14971 and IEC 62304 remain the sector authorities for medical-device quality,

safety risk and software life-cycle evidence [7,13,14]. The value lies in allocating these contributions to one operational architecture rather than treating the number of adopted standards as an end in itself.

This paper asks how those contributions can be combined so that AI governance remains connected to the technical and medical-device decisions it is intended to control. It proposes a vendor-neutral entity-and-relationship architecture that can be implemented through controlled documents, application life-cycle management systems, specialised governance platforms or combinations of them. The implementation technology is secondary. What defines the architecture is the identity of the governed entities, the semantics of their relationships, the authority of each system of record and the conditions under which decisions must be made or revisited.

Four propositions are examined. First, the organisational governance requirements of ISO/IEC 42001 become operationally stronger when their risk process is implemented using the deeper AI-specific guidance of ISO/IEC 23894. Secondly, AI Risk and medical-device safety risk should remain distinct but explicitly mapped, with ISO/TS 24971-2 providing the bridge to ISO 14971. Thirdly, controlled entities and semantic relationships provide stronger decision continuity than documents alone because they preserve the dependencies needed to govern reuse, change and post-market feedback. Fourthly, a PCCP becomes operationally reliable when it is treated as the authorised regulatory projection of a wider Controlled Evolution capability, rather than as a static plan disconnected from the quality processes that must execute each modification.

II. METHOD

This work is a conceptual standards synthesis and structured narrative review rather than an experimental study or a claim of validated best practice. The analytical method began with the governance objectives defined in the Introduction, rather than with a list of available standards. Five functions required normative support: organisational governance; AI-specific risk management; technical decomposition and traceability of machine-learning systems; product-focused regulatory quality management; and integration with medical-device quality, software and safety processes. An instrument was treated as core only where it supplied a distinct function that could not be assigned more appropriately to another selected source.

A structured narrative review was performed in parallel to test whether the design propositions were supported, contradicted or left unresolved by prior work. The review covered peer-reviewed literature indexed through biomedical and multi-disciplinary scholarly search services, direct publisher repositories, standards catalogues, regulatory repositories and publicly available professional guidance. Search themes combined artificial intelligence or machine learning with life cycle, governance, quality management, risk management, medical device, clinical implementation, model documentation, dataset documentation, monitoring, drift, change control and post-market surveillance. Backward citation chasing was performed from relevant systematic and scoping reviews. Priority was given to sources that

described an end-to-end framework, reported implementation experience, synthesised multiple governance frameworks or defined a reusable assurance artefact. The search was updated through 1 August 2026.

This was not a systematic review and no claim of exhaustive retrieval is made. Its purpose was analytical triangulation: to identify the most relevant academic, health-system, regulatory and professional approaches, compare their principal conclusions with the proposed architecture and avoid presenting convergent ideas as novel. Public professional and industry evidence was included where it disclosed operational expectations or assessment practice; confidential organisational experience was not treated as citable evidence.

The selected instruments were analysed at the level of stated scope, process model, principal controlled subjects, life-cycle expectations, decision authority and relationship with other management systems. Requirements and guidance were allocated to governance objectives before any record structure was designed. Where a source requires an outcome but does not prescribe an artefact, the model defines the minimum controlled entity or semantic relationship needed to maintain that outcome over time. These design-derived elements are identified as such and are not presented as records explicitly mandated by a named instrument.

The synthesis followed five design rules. First, no AI-specific record was created where an established QMS record already provides the authoritative evidence. Secondly, technically different objects were not collapsed for administrative convenience: an AI Model is not synonymous with an AI System, and a Dataset is not merely an attachment to a model. Thirdly, a many-to-many relationship was represented as a controlled association when it carried properties material to approval, risk or reproducibility. Fourthly, risk information could exist at more than one semantic level only where the decision scope and authority of each record remained explicit. Fifthly, every entity, relationship, field and gate had to support at least one practical control purpose: accountability, approval, impact analysis, monitoring, change control or evidence retrieval.

The resulting model was tested conceptually against an end-to-end scenario involving an AI-assisted diagnostic function. The test asked whether a reviewer could navigate from intended clinical use to the deployed model version; identify the exact datasets and transformations used for training and independent evaluation; locate a population-specific performance concern and the assumptions that made it material; trace the resulting AI Risk to the corresponding medical-device safety analysis, controls and effectiveness evidence; reconstruct the release decision; and determine how a subsequent monitoring signal would reopen affected decisions. A second test examined whether the same query remained possible after a model, dataset, supplier or deployment-site change.

The architecture was also instantiated in a working prototype to test whether the proposed semantics could support navigation, deterministic completeness findings, independent review, life-cycle state control and execution of a PCCP. The Controlled Evolution scenario required the prototype to define a system-level change boundary, maintain specific Planned Modifications,

link each modification to its protocol and impact assessment, distinguish planned from implemented model and data versions, hold change authorisation when source records were incomplete, and assemble a PCCP from the controlled records. The prototype is an implementation probe, not empirical evidence of regulatory conformity or governance effectiveness. Portability was assessed by replacing the medical-device QMS and safety-risk layer with another sector's authoritative control framework while retaining the AI-specific entities and relationships.

III. RESULTS: A LAYERED STANDARDS ARCHITECTURE

A. *Selection of the core standards*

The selection is based on functional complementarity. The proposed approach needs an organisational mechanism for setting policy, assigning responsibility and evaluating effectiveness; a sufficiently developed method for identifying and treating AI-related risk; a technical vocabulary that prevents systems, models and data from being conflated; product-focused processes responsive to European AI regulation; and a controlled interface with existing medical-device quality, software and safety processes. Each selected instrument fills one or more of these functions, but none is allowed to displace an instrument with more specific authority.

ISO/IEC 42001 governs the organisation's capability to direct and control AI. ISO/IEC 23894 guides the discipline through which AI-related uncertainty and its effects on objectives are identified, analysed, evaluated, treated, communicated and monitored. ISO/IEC 23053 supplies a shared technical language for the components and functions of machine-learning systems. The EU AI Act establishes binding provider obligations, and EN 18286 translates its quality-management expectations into product-focused processes. ISO 14971 remains the authority for medical-device safety risk, while ISO/TS 24971-2 directs that process towards machine-learning-specific causal factors. ISO 13485 and IEC 62304 preserve the regulated quality and software processes through which controls and evidence are implemented.

The architecture is deliberately selective rather than exhaustive. ISO/IEC 5338:2023 can add a more detailed horizontal AI life-cycle process model, ISO/IEC 42005:2025 can deepen AI system impact assessment, and the ISO/IEC 5259 series can provide specialised data-quality measures and processes [8-10]. Cybersecurity and privacy should be governed through the applicable ISMS, PIMS and product-security standards. These are meaningful extensions where the system context requires them, but adding them to the core set would not change the six-entity relationship model. Their controls and evidence can be linked through the same AI System, AI Model, Dataset and AI Risk context.

B. *Why ISO/IEC 42001 requires a stronger risk method*

ISO/IEC 42001 is a management-system requirements standard. Its primary contribution is to turn AI governance into a repeatable organisational capability: context and objectives are established, responsibilities assigned, resources and competence provided, operations controlled, performance evaluated and defi-

Table 1. Selected instruments and their function in the reference architecture

Instrument	Primary contribution	Why it is included	Boundary in the proposed model
ISO/IEC 42001:2023	Organisational AIMS requirements and continual improvement	Establishes governance, accountability, policy, objectives, operational planning, performance evaluation and improvement	Governs the AI portfolio and management system; it does not replace product design control or prescribe the complete technical risk method
ISO/IEC 23894:2023	AI-specific risk-management guidance	Extends risk thinking beyond device safety to affected persons, groups, society, organisational objectives and AI-specific uncertainty	Informs the method represented by the AI Risk entity; risks that concern medical-device harm are linked to, but do not replace, the ISO 14971 analysis
ISO/IEC 23053:2022	Conceptual framework for AI systems using ML	Distinguishes systems, models, data and ML processes and supplies shared technical language	Informs the entity model; it is descriptive and does not itself establish approval or regulatory controls
EN 18286:2026	QMS for EU AI Act regulatory purposes	Converts Article 17 and associated life-cycle duties into auditable product-focused processes	Adds EU regulatory quality expectations; it can extend an existing QMS but does not make the AIMS or ISO 13485 QMS redundant
Regulation (EU) 2024/1689, as amended	Binding EU AI requirements	Establishes risk classification and obligations for risk management, data governance, records, transparency, human oversight, performance and monitoring	Supplies legal obligations and timing; conformity remains product- and role-specific
ISO/TS 24971-2:2026	Application of ISO 14971 to ML-enabled medical devices	Addresses ML-specific hazards, data and model behaviour, bias, drift, retraining, monitoring and rollback	Provides the bridge to the medical-device risk file; it excludes LLM- and generative-AI-enabled devices
ISO 14971:2019	Medical-device safety risk management	Preserves the accepted hazard-to-harm logic and life-cycle risk-control process	Remains authoritative for medical-device safety risk and benefit-risk decisions
ISO 13485:2016	Medical-device QMS	Controls design and development, suppliers, records, nonconformity, CAPA, change and release	Remains the authoritative management system for regulated medical-device processes and evidence
IEC 62304:2006+A1:2015	Medical-device software life-cycle processes	Connects AI components to controlled software planning, implementation, maintenance and problem resolution	Remains authoritative for software life-cycle records; AI governance supplies additional model, data and monitoring context

ciencies improved. Its risk requirements and control framework place AI risk, impact, data, system life cycle, interested parties, third-party relationships and responsible use within that management cycle [1]. This makes it the appropriate backbone for accountability and assurance.

That backbone does not, and should not, prescribe a universal technical risk method. The same requirements must be usable by an organisation procuring a low-impact administrative tool, a provider of a general-purpose service and a manufacturer developing a safety-critical medical device. ISO/IEC 42001 consequently requires risk assessment and treatment to be established and operated, but leaves the organisation to define the method, criteria and level of detail appropriate to its context. It does not supply a medical-device hazard-to-harm method, nor does it define the object and relationship granularity needed to trace model, data and deployment dependencies.

The practical failure mode is subtle. An organisation can implement the management-system clauses convincingly while using an AI risk register whose entries are too broad to direct technical work. A statement such as “biased outputs may affect patients” identifies a concern, but does not identify the affected use, population, dataset role, model baseline, causal mechanism, control owner, acceptance criterion or monitoring signal. The

management system may then prove that reviews occurred without proving that the relevant source of risk was understood or controlled.

ISO/IEC 23894 is the appropriate enrichment because it is horizontal, AI-specific and expressly concerned with integrating AI risk management into organisational activities [2]. It broadens the framing of consequence beyond patient injury to objectives and effects involving individuals, groups, organisations and society. It directs analysis towards data, models, system behaviour, human interaction, deployment context, uncertainty and changing operating conditions. It also supports communication, monitoring and review across the life cycle. This scope is necessary because discriminatory performance, loss of human agency, opacity, privacy, security, environmental effect, operational dependence or interference with fundamental rights may be material even when the concern has not yet formed part of a medical-device hazard-to-harm sequence.

The relationship between the two standards is therefore procedural and methodological. ISO/IEC 42001 establishes the accountable system within which risk work must occur, how it is resourced, how results influence objectives and controls, and how effectiveness is audited and reviewed. ISO/IEC 23894 informs how individual AI Risks are framed, analysed, evaluated,

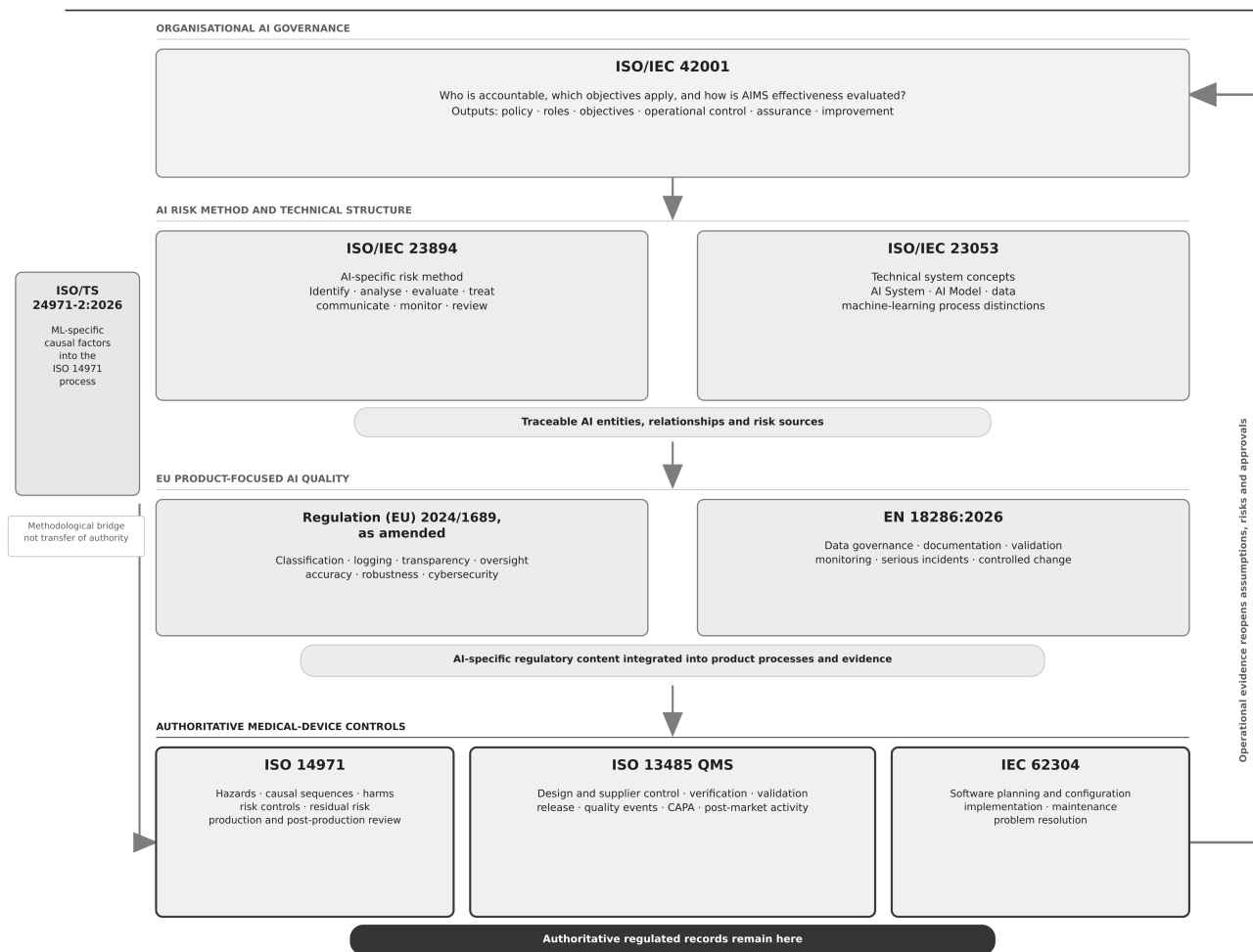


Fig. 1. The selected instruments answer different control questions. Their value arises from explicit interfaces and preserved authority: the AIMS governs AI-specific context and decisions, while established medical-device processes remain authoritative for regulated evidence and safety conclusions.

treated, communicated and monitored. The proposed AI Risk entity is a design response to this combination; it is not a record format mandated by either standard. Its purpose is to make the chosen method operational by connecting each material risk to the technical or contextual source, accountable decision, treatment evidence and monitoring condition.

This enrichment is still not sufficient for medical-device safety. The wider AI-risk method may identify an effect on fairness, trust, regulatory compliance or clinical performance, but it does not replace the causal safety analysis, control hierarchy, residual-risk evaluation and production or post-production feedback required for a medical device. The architecture therefore uses ISO/IEC 23894 to deepen AI risk management and a separate, explicit bridge to ISO 14971 for safety-relevant concerns. Section 5 explains that bridge and the reasons the two risk representations should remain coordinated rather than collapsed.

C. Why the EU AI Act belongs in the model before full high-risk application

The EU AI Act is already law, but its application is phased. Regulation (EU) 2024/1689 entered into force on 1 August

2024. Prohibited practices and AI-literacy provisions began to apply in February 2025, and governance and general-purpose AI obligations began to apply in August 2025. As at 1 August 2026, further general provisions and transparency obligations apply from 2 August 2026 [5]. Regulation (EU) 2026/1744, the Digital Omnibus on AI, extended the application of high-risk rules for Annex III uses to 2 December 2027 and for AI embedded in regulated products, including qualifying medical devices, to 2 August 2028 [11].

The extended date is not a reason to exclude the Act from current governance. Medical-device evidence is cumulative. Intended purpose, data provenance, dataset separation, model selection, performance acceptance criteria, human oversight, logging, supplier controls and post-market monitoring cannot be reconstructed reliably at the end of development. A model entering development in 2026 may be placed on the market, materially modified or still supported when the high-risk provisions apply. Designing the evidence architecture now avoids retrofitting legal concepts onto technical decisions that are no longer reproducible.

The Act also changes the scope of what must be governed. MDR and IVDR requirements address safety and performance [15,16], but the AI Act adds explicit AI-system concerns, including data governance, functional logging, transparency, human oversight, accuracy, robustness, cybersecurity and risks to fundamental rights. MDCG 2025-6 describes the regimes as simultaneous and complementary and encourages manufacturers to integrate the required testing, reporting and documentation into existing MDR or IVDR procedures [12]. The proposed architecture follows that direction: it creates AI-specific context and traceability while routing authoritative product evidence through the established QMS.

EN 18286:2026 makes this anticipatory approach more practical. Published on 31 July 2026, it defines a cross-sector QMS for organisations providing AI systems and addresses regulatory strategy, management responsibility, product realisation, risk treatment, data management, supply chains, verification and validation, deployment, monitoring, change, serious incidents and nonconformity [4]. Publication alone should not be confused with citation as a harmonised standard in the Official Journal. Once cited, a harmonised standard can support a presumption of conformity for the requirements it covers. Even before that step, EN 18286 provides the most direct published quality-management interpretation of Article 17 and a valuable benchmark for current design.

D. Convergence with academic and professional literature

The literature review found strong convergence on the need for life-cycle governance, although the proposed mechanisms and institutional vantage points differ. Overgaard and colleagues provide the closest quality-management analogue to this paper. They argue that health AI should be governed through a risk-based QMS across development, deployment and use and that this capability should be integrated with existing organisational quality infrastructure to reduce redundancy and close the translation gap [21]. Bartels and colleagues reach a comparable conclusion from hospital practice: manufacturer controls are necessary but insufficient unless the deploying healthcare organisation also embeds model operation, change, monitoring and responsibility in a life-cycle QMS [22]. Bedoya and colleagues report an operational governance framework for local deployment and oversight, demonstrating that portfolio-level governance and continuing model management are practical rather than merely aspirational [23].

Process studies reinforce the same conclusion. Boag and colleagues' Algorithm Journey Map shows that an implemented clinical AI solution is a socio-technical sequence of problem selection, development, integration and continuing life-cycle management involving many stakeholders and decisions [24]. Khan and colleagues' systematic review identified a large body of proposed clinical AI frameworks but found an imbalance: planning activities were substantially better represented than implementation, monitoring and evaluation [25]. Crossnohere and colleagues similarly found that available guidance gives less attention to engagement and translational surveillance than to earlier-stage technical activity [26]. Those findings support ex-

plicit operational-acceptance, monitoring, change and retirement gates rather than treating governance as a pre-release assessment.

Consensus and legal scholarship support a broader conception of the governed object. FUTURE-AI organises trustworthy healthcare AI around fairness, universality, traceability, usability, robustness and explainability, with practices spanning design, development, validation, regulation, deployment and monitoring [27]. Solaiman and colleagues independently propose a "True Lifecycle Approach" that integrates healthcare law, ethics and patient protections across research and development, approval and post-implementation governance [28]. McCradden and colleagues describe healthcare AI as a sociotechnical system whose normative assessment cannot be reduced to model performance [42]. These sources converge with the proposed AI System and AI Use Case entities: the model must be governed in relation to people, workflow, institutional responsibilities, legal duties and the conditions in which outputs acquire consequence.

The translational literature also supports the separation of model evidence from system approval. Wiens and colleagues call for an engaged, systematic process from problem formulation through deployment [29], while Kelly and colleagues identify dataset shift, external validation, workflow integration, human factors and post-deployment performance as barriers to clinical impact [30]. The TEHAI framework assesses capability, utility and adoption rather than accuracy alone [31], and DECIDE-AI requires early clinical evaluation to report the actual human-AI interaction, workflow and errors [32]. CONSORT-AI and SPIRIT-AI extend trial reporting so that AI interventions, inputs, outputs, failure handling and analysis are sufficiently transparent to evaluate [33,34]. These sources are complementary to the present architecture: reporting guidance defines what a study should disclose, whereas the entity model preserves the governed baseline and its evidence before, during and after any one study.

Documentation research supplies direct support for versioned Model and Dataset entities. A scoping review by Brereton and colleagues found that model documentation supports transparency and translation but that formats and intended users vary, requiring documentation to be tied to concrete decisions [35]. Model Facts labels [36], Model Cards [39] and early algorithm registration [37] make intended use, limitations, performance and development history visible to different audiences. Datasheets for Datasets make data provenance, composition, collection, recommended uses and limitations explicit [38]. These artefacts are valuable evidence, but none is a substitute for configuration control: a model card or datasheet describes an object, while the proposed architecture gives that object identity, version, status, relationships and decision rights.

Software-engineering and field research explain why the Dataset-Model Association is necessary. Sculley and colleagues show that data dependencies, configuration, feedback loops and changes in the external world generate technical debt that conventional code-centric controls do not expose [40]. Sambasivan and colleagues show how neglected data work produces "data cascades" whose downstream effects become difficult to locate in high-stakes systems [41]. A controlled association between an exact dataset version, its role, transformation and a model

baseline makes those dependencies reviewable. It also supplies the transitive impact analysis needed when a label set, source population, preprocessing rule or independence claim changes.

Institutional and regulatory guidance reaches similar operational conclusions. WHO guidance treats ethics, accountability and governance as continuous responsibilities and its regulatory considerations emphasise documentation, intended use, validation, data quality, human intervention, stakeholder collaboration, cybersecurity and post-release learning [43,44]. The IMDRF Good Machine Learning Practice principles, NIST AI Risk Management Framework and FDA total-product-life-cycle guidance similarly emphasise multidisciplinary governance, representative data, independent test evidence, human-AI performance, transparency, monitoring and controlled change [45-48]. The non-mandatory Team-NB/IG-NB questionnaire makes these topics visible from notified-body assessment practice [49], while the Coalition for Health AI public-comment draft organises practical health-system considerations across a defined AI life cycle [50]. These professional sources do not all carry the same legal or normative authority, but their convergence is material evidence of operational expectations.

The synthesis therefore supports the architecture while limiting its originality claim. Life-cycle governance, quality integration, socio-technical assessment, documentation and monitoring are well established themes. The contribution of this paper is their operational combination: a vendor-neutral relationship model, a deliberate enrichment of ISO/IEC 42001 with ISO/IEC 23894, a controlled Dataset-Model Association, and a two-layer risk architecture in which ISO/TS 24971-2 links AI-specific causal factors to an ISO 14971 medical-device safety analysis without erasing broader AI risk.

IV. THE PROPOSED ENTITY AND RELATIONSHIP MODEL

A. *Why governance requires entities rather than documents alone*

A document-centred system is effective when a governed object and its evidence have a stable one-to-one relationship. Machine-learning development rarely has that property. One model may support several clinical uses; one dataset may contribute to several models; one system may contain or depend on several models; a dataset may be suitable for training in one configuration but unsuitable as independent test evidence in another; and one post-market signal may challenge a model, deployment context, risk control and release assumption simultaneously. Reports can describe these facts, but a repository of reports does not, by itself, preserve them as queryable dependencies.

The control problem becomes most visible after change. If a labelling error is discovered in one Dataset version, the manufacturer must identify which data partitions were affected, how each partition was used, which models and evaluations relied on it, which authorised use cases depend on those models, which risks and acceptance decisions assumed the data were valid, and which released configurations may require action. Finding the relevant reports is insufficient if their conclusions cannot be connected reliably to the exact configuration under review.

The architecture therefore uses six core entities. They are not six document templates and are not all named artefacts in the source instruments. Each is a controlled representation with a stable identity, defined owner, state, scope, material relationships and supporting evidence. Version control is applied where a change can alter technical behaviour or evidential validity; review state is used where the underlying object remains stable but its authorisation must be reconsidered. Documents continue to provide approved analysis and objective evidence, while the entities preserve what that evidence describes and which decision it supports.

These six entities describe the governed configuration; they are not the complete set of records needed to operate the governance system. Monitoring Signals, AI Reviews, Planned Modifications, findings and gate decisions are transactional control records that act upon one or more core entities. The distinction is consequential. A Planned Modification is not part of the deployed technical architecture, and an AI Review is not evidence merely because a meeting occurred. Their control value lies in identifying the affected baseline, the decision being requested, the evidence considered, the accountable participants and the resulting change in authorisation.

B. *The AI System as governance anchor*

The AI System is the primary governance anchor because organisational, legal and product decisions attach to a system placed into a defined context, not to a mathematical model in isolation. Its boundary should identify the functions and components inside the governed system; the data flows, infrastructure, external services and human activities on which it depends; and the outputs through which it can influence the clinical or operational environment. The relationship to the regulated medical device must be explicit. Depending on the product architecture, one AI System may coincide with the device software, form one function within a broader device, or be an externally supplied component on which the device depends.

The AI System record should not reproduce the technical documentation. It should identify the current approved governance context: intended purpose, boundary, roles, jurisdictions, AI Use Cases, approved AI Models, relevant data dependencies, unresolved AI Risks, gate decisions, monitoring status and links to controlled product evidence. A reviewer should be able to use it to establish which configuration is authorised and then navigate to the evidence held in the competent system of record.

C. *The AI Use Case as the unit of context*

An AI Use Case translates a system-level intended purpose into a concrete interaction that can be assessed and authorised. It identifies who uses the system, whose interests may be affected, what information enters, what output is produced, how that output is interpreted, what decision or action may follow, which human oversight is available and what happens if the output is incorrect, unavailable or used outside validated conditions. The entity should also identify the applicable environment, population, workflow dependencies and operational owner. Separating use cases prevents evidence for a model from being treated as

Table 2. Convergence of the literature with the proposed architecture

Evidence cluster	Representative sources	Principal conclusion in prior work	Response in the proposed architecture	Remaining distinction
Life-cycle QMS and organisational oversight	Overgaard et al. [21]; Bartels et al. [22]; Bedoya et al. [23]	Safe translation requires accountable, risk-based governance across development, deployment, operation and change, integrated with existing quality structures	AIMS governance gates are linked to the ISO 13485 QMS, rather than operated as a parallel compliance programme	The present model allocates authority at entity and relationship level and adds an explicit ISO/IEC 42001 and ISO/IEC 23894 pairing
End-to-end implementation and monitoring	Boag et al. [24]; Khan et al. [25]; Crossnohere et al. [26]	Frameworks often describe planning better than operational integration, surveillance, maintenance and decommissioning	Operational Acceptance, Monitoring Review, Change, Restriction, Rollback, Resumption and Retirement are controlled decisions	Decision scope is tied to a versioned system, use case, model, dataset and risk baseline
Trustworthy, patient-centred and socio-technical governance	FUTURE-AI [27]; Solaiman et al. [28]; McCradden et al. [42]	Governance must span the full life cycle and include people, workflow, law, ethics, equity and patient rights rather than technical performance alone	AI System and AI Use Case records make affected parties, context, oversight, misuse and consequences governable	The reference architecture focuses on auditable operational semantics and can link, but does not replace, legal or ethical analysis
Clinical translation and evaluation	Wiens et al. [29]; Kelly et al. [30]; Reddy et al. [31]; DECIDE-AI [32]; CONSORT-AI and SPIRIT-AI [33,34]	Evidence must address the clinical problem, external validity, workflow, users, failure modes and real-world effects	Study and validation evidence is linked to the exact approved use, model and data configuration	Reporting a study is distinguished from controlling the system baseline and continuing authorisation
Model and data transparency	Brereton et al. [35]; Model Facts [36]; algorithm registration [37]; Datasheets [38]; Model Cards [39]	Reusable documentation should expose provenance, purpose, performance, limitations and intended users	Model and Dataset are versioned entities; documentation becomes governed evidence supporting them	The Dataset-Model Association records the material properties of each use of data rather than relying on two static documents
Dependency, drift and data quality	Sculley et al. [40]; Sambasivan et al. [41]	ML risk arises from hidden dependencies, configuration, feedback loops, neglected data work and external change	Relationship-level lineage, change impact analysis, monitoring triggers and rollback preserve causal traceability	Risk can attach directly to a material association and propagate to affected systems and safety analyses
Public professional and regulatory practice	WHO [43,44]; IMDRF [45]; NIST [46]; FDA [47,48]; Team-NB [49]; CHAI [50]	Multidisciplinary governance, representative data, independent evaluation, transparency, monitoring and controlled modification are recurring expectations	These expectations are allocated to auditable entities, Planned Modifications, governance gates and authoritative QMS records; PCCP content remains traceable to the source records used to execute it	Their different jurisdictions and authority levels remain visible; convergence is not misrepresented as equivalence or conformity

proof of acceptability in every situation in which the model could technically be used.

This distinction is particularly important for reused or foundation models. A model may provide acceptable performance as a prioritisation aid under specialist oversight but be unacceptable as an autonomous diagnostic decision. The model is unchanged, but the exposure, ability to intervene, severity of consequences and transparency needs are different. The AI Use Case is therefore the natural location for context-specific impact assessment, foreseeable misuse and human-oversight requirements.

D. Models, datasets and the association between them

The AI Model is a versioned technical object. It should identify the model architecture and configuration sufficiently to distinguish one reproducible baseline from another. Its approval state should reflect defined evidence, not merely successful training. A validated model is not necessarily approved for deployment, and a model approved for one AI Use Case is not automatically approved for another.

The Dataset is independently governed because its provenance, population, labels, legal constraints and quality characteristics can change or be reinterpreted without changing the model record. Dataset governance should distinguish raw sources from curated versions and should preserve the transformations and decisions that produced an approved data baseline. Training, validation and test sets should be identifiable as distinct roles even when they originate from a common source.

The Dataset-Model Association is the most consequential design-derived entity in the model. A simple statement that a model “uses” a dataset is insufficient because it omits the properties that determine evidential validity. The association should identify the exact versions and partitions used; the role assigned to the data; preprocessing, augmentation and feature-generation steps; the controls used to prevent leakage; the basis for any claim of independence; the represented population and acquisition conditions; predefined acceptance criteria; and the evidence produced. The same Dataset can be appropriate for training while being unsuitable as independent test evidence. Elevating the association to a controlled object makes these differences reviewable and permits an AI Risk to attach to the relationship in which the concern actually arises. This operationalises the dependency and data-work problems described in the technical-debt and data-cascade literature [40,41], while allowing Datasheets or equivalent records to remain the descriptive evidence for the underlying data asset [38].

E. Relationship semantics

Relationship names must express meaning in both directions. A generic “related to” link creates navigability but not traceability because it does not tell a reviewer why the connection exists or what must be reconsidered when either side changes. Table 4 defines the minimum vocabulary, inverse semantics, preferred cardinality and control consequence. Implementations may extend this vocabulary, but should not replace a precise relationship with a weaker generic link.

Table 3. Core entities, classifications and derivation

Entity	Classification and scope	Minimum governed content	Practical control value	Derivation and authority
AI System	Top-level governance and system-context entity	Intended purpose; system boundary; provider, manufacturer and deployer roles; applicable classification; constituent components and external dependencies; life-cycle state; accountable owner; overall governance decisions	Defines the socio-technical object to which accountability attaches and provides the navigation point for the approved governance context	Corresponds to the governed AI system in ISO/IEC 42001, ISO/IEC 23053, EN 18286 and the AI Act; it may map to, form part of, or depend upon a medical-device product configuration controlled in the QMS
AI Use Case	Design-derived context-of-use and impact entity	Specific objective; target users and affected persons; workflow and environment; input and output meaning; reliance and human oversight; validated conditions; foreseeable misuse; applicable performance and impact criteria	Prevents evidence for one application from being treated as approval for another; localises context-specific benefit, consequence, oversight and transparency decisions	Derived from intended-purpose, foreseeable-use, interested-party, impact and deployment-context concepts across ISO/IEC 42001, ISO/IEC 23894, EN 18286 and the AI Act; not a mandated record name
AI Model	Versioned technical configuration item	Model family and architecture; training configuration and material hyperparameters; code and dependency baseline; version; performance and limitations; opacity characteristics and explanation methods; approval, deployment and deprecation state	Distinguishes the statistical artefact from the complete system, supports reproducibility and allows reuse without transferring approval automatically between uses	Informed by the model and ML-process distinctions in ISO/IEC 23053 and by model-specific controls in ISO/TS 24971-2, EN 18286 and AI technical-documentation requirements
Dataset	Versioned data asset and evidence entity	Provenance; population and acquisition context; inclusion and exclusion; labels and ground truth; quality characteristics; legal and ethical constraints; known limitations; version and approval state	Makes fitness, representativeness, bias, provenance and change impact reviewable independently of any one model	Informed by data concepts and governance expectations in ISO/IEC 23053, ISO/IEC 23894, ISO/TS 24971-2, EN 18286 and Article 10 of the AI Act
Dataset-Model Association	Design-derived associative configuration and evidence entity	Exact Dataset and AI Model versions; role such as training, tuning, validation, test or monitoring; partition; preprocessing and feature generation; leakage and independence controls; acceptance criteria; resulting evidence	Controls material properties of a many-to-many relationship that belong neither to the Dataset nor the AI Model alone; enables transitive impact analysis and reproducibility	Derived from the separation of data and ML processes in ISO/IEC 23053 and from lineage, evaluation and risk-control needs in ISO/TS 24971-2, EN 18286 and the AI Act; not a mandated record name
AI Risk	Design-derived governance and risk-decision entity	Risk source or changed condition; affected objectives, parties and entities; plausible consequences; uncertainty and assumptions; analysis and evaluation; owner; treatment; residual conclusion; monitoring and links to competent risk authorities	Preserves the wider AI-risk context, attaches risk to its technical or contextual source and supports explicit mapping to device safety, privacy, security or other authoritative risk processes	Operational representation of the ISO/IEC 42001 risk process as enriched by ISO/IEC 23894 and applicable AI Act obligations; the ISO 14971 risk-management file remains authoritative for medical-device safety

The practical benefit is impact analysis. If a Dataset version is withdrawn because labels were found to be unreliable, the organisation can identify every Dataset-Model Association that used it, the affected model baselines, the AI Use Cases that rely on those models, the open or closed AI Risks whose assumptions changed, and the released systems for which corrective action may be necessary. A document repository can contain all of this information, but only a controlled relationship model makes the dependency query reliable.

V. INTEGRATED AI AND MEDICAL-DEVICE RISK MANAGEMENT

A. Two risk perspectives that should not be collapsed

ISO/IEC 23894 and ISO 14971 both organise risk work across a life cycle, but they answer different decision questions. ISO/IEC 23894 adopts the broader organisational risk concept of uncertainty affecting objectives and considers effects involving organisations, individuals, groups and society [2]. ISO 14971 is directed towards the safety of a medical device. Its analysis traces hazards, foreseeable sequences of events, hazardous situations and harms, then estimates and evaluates risk, implements controls, evaluates residual risk and uses production and post-production information to maintain the analysis [7].

These perspectives overlap whenever an AI characteristic can contribute to harm, but neither contains the other. A biased model may create a medical-device safety risk if reduced sensitivity in a subgroup delays diagnosis. It may also create discrimination, inequitable access, loss of trust or regulatory exposure that should be governed even where the ISO 14971 definition of harm is not met. Conversely, a device may have electrical, mechanical or biocompatibility hazards that are not AI-specific and do not belong in the AI Risk register.

The architecture therefore coordinates two representations with different authority. The AI Risk entity captures the AI-specific source or changed condition, affected context, parties, objectives, uncertainty, consequence domains, treatment and monitoring. The ISO 14971 risk-management file remains the authority for medical-device safety analysis, control and residual-risk decisions. Where the same factual concern contributes to both, an explicit mapping connects the records. The mapping is semantic, not numerical: a rating assigned under an AI-risk method should not be copied into the device risk file unless the estimation criteria are genuinely the same.

The mapping can be many-to-many. One AI Risk, such as non-representative development data, may contribute to several hazardous situations across different clinical uses. One device

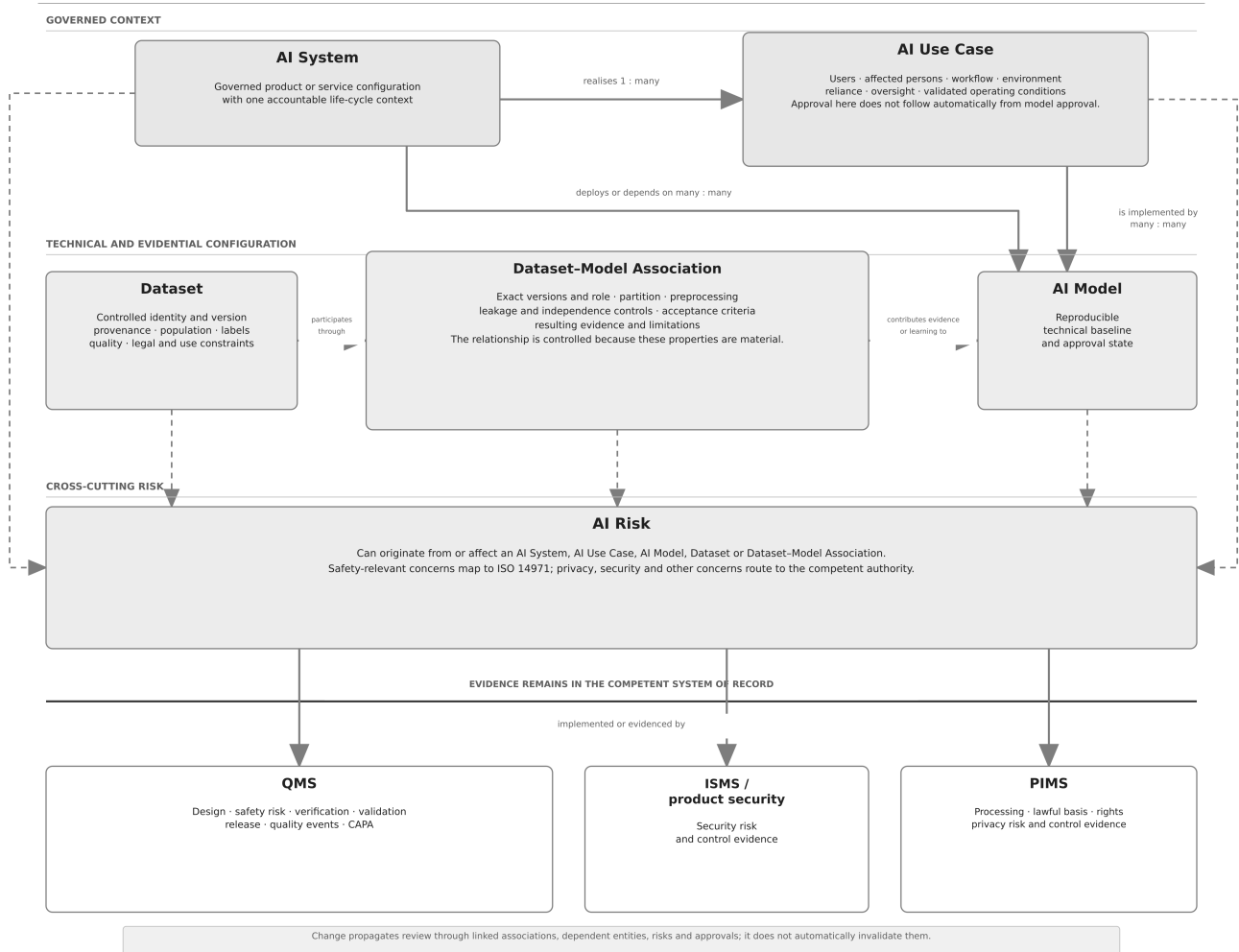


Fig. 2. The entity topology preserves the distinction between authorised context, technical baseline, data lineage and risk. AI Risk can attach to any material source or relationship, while controls and objective evidence remain in the management system competent to authorise them.

safety-risk item may also have several AI-specific sources, such as data shift, model brittleness, inadequate human oversight and a supplier change. Preserving this cardinality is necessary for causal review. A one-to-one cross-reference would conceal which source changed and which controls or use contexts must be reconsidered.

A practical screening rule is required. When an AI Risk is created or materially revised, the owner should determine whether the source or consequence can contribute to harm involving a patient, user or other person. If it can, the applicable ISO 14971 analysis is created or updated and linked. If it cannot, the AI Risk remains governed through the appropriate organisational, legal, ethical, security, privacy or operational process. This screening conclusion is itself reviewable when new evidence changes the plausible consequence.

B. The role of ISO/TS 24971-2:2026

ISO/TS 24971-2 is the critical bridge because it is intended for use with ISO 14971 and the general application guidance in ISO/TR 24971, while directing them towards technologies whose behaviour depends on data, statistical learning

and deployment context [6,7,17]. It covers machine-learning-specific hazards and hazardous situations, data quality, non-representative training data, overfitting and underfitting, model design, bias, transparency, explainability, usability, autonomy, human oversight, model drift, retraining, updating, rollback and post-production monitoring [6]. These topics convert general safety obligations into questions that can be answered at the Dataset, Dataset-Model Association, AI Model and AI Use Case levels.

Its importance is methodological. ISO 14971 remains fully applicable, but an analysis can describe “incorrect output” as a software failure without examining why statistical performance varies across populations, acquisition devices, sites or time. Such a description is too coarse to select or monitor the most effective control. ISO/TS 24971-2 directs attention towards the conditions that produce those errors and towards controls that can be implemented in data selection, labelling, model design, performance evaluation, user information, human oversight, deployment restriction, change control and monitoring. It enriches the causal chain without creating a competing definition of medical-device risk.

Table 4. Minimum relationship vocabulary and control semantics

Source	Outward relationship	Target	Inverse meaning	Preferred cardinality	Relationship class	Control consequence and principal rationale
AI System	realises	AI Use Case	is governed within	One-to-many within one governed system	Context and authorisation	Defines the applications included in the system boundary; an intended-purpose, workflow, user, population or jurisdiction change reopens the use-context assessment
AI System	deploys or depends on	AI Model	contributes to behaviour of	Many-to-many	Configuration and dependency	Identifies the model baseline on which system behaviour depends; replacement, retraining, supplier or runtime changes trigger transitive impact review
AI Use Case	is implemented by	AI Model	supports	Many-to-many	Context-to-technology allocation	Prevents model approval from being transferred automatically between uses with different reliance, oversight, consequence or performance criteria
Dataset	participates through	Dataset-Model Association	uses this Dataset version	One Dataset version to many Associations; each Association references one Dataset version	Lineage and evidence configuration	Preserves the role, partition, transformations, independence and acceptance evidence for each use of the data
Dataset-Model Association	contributes evidence or learning to	AI Model	is developed, evaluated or monitored through	One AI Model baseline to many Associations; each Association references one model baseline	Lineage and evidence configuration	Makes a many-to-many dependency reproducible and permits the relationship itself to carry risk and approval state
AI Risk	originates from or affects	Any core entity or Association	is source or affected context of	Many-to-many	Risk causality and propagation	Locates the source, affected contexts and assumptions; new evidence, drift, incidents, complaints or change can reopen the linked assessment
AI Risk	maps to	ISO 14971 risk-analysis item	receives an AI-specific causal contribution from	Many-to-many	Cross-authority risk mapping	Connects an AI source or effect to a hazard, event sequence, hazardous situation, harm, control or production/post-production conclusion without equating the records
AI Risk treatment	is implemented or evidenced by	Authoritative QMS, ISMS or PIMS record	implements or demonstrates treatment of	Many-to-many	Control and evidence	Preserves one competent system of record for design outputs, tests, labelling, procedures, supplier controls, CAPA, privacy or security evidence
Monitoring signal	challenges	Entity, Association, AI Risk or safety control	has an assumption tested by	Many-to-many	Feedback and review trigger	Routes operational evidence to every approval assumption it can invalidate, rather than recording the signal only in a monitoring report
Planned Modification	proposes a bounded change to	Entity, Association, approved baseline or control	is affected by	Many-to-many	Change intent and impact	Identifies the precise proposed change, its rationale, initiating signal or objective, affected configuration, boundary assessment, expected evidence and regulatory route before implementation begins
Planned Modification	produces or revises	AI Model version, Dataset version, Dataset-Model Association, use condition or controlled evidence	results from	One modification may produce several controlled outputs; each output identifies its originating change	Change execution and lineage	Separates authorisation to perform a change from the versioned outputs produced by it and preserves the route from change intent to released configuration
AI Review	decides upon	Planned Modification, entity, risk position or identified baseline	is approved, conditionally approved or rejected by	One review may decide several explicitly scoped records; each record may be reviewed repeatedly over time	Independent decision and signature	Records subject, evidence, participants, independence, signatures, conditions, decision, effective date and supersession without treating record completeness as approval
Governance decision	approves, restricts, suspends, resumes or retires	Identified entity and effective baseline	is subject to	One decision may cover several explicitly listed objects; each object may have many decisions over time	Authority and state transition	Makes decision scope, accountable approver, effective configuration, conditions and superseded decision unambiguous

The Technical Specification is particularly valuable at four interfaces. During planning, it helps translate the intended purpose and foreseeable use into ML-specific risk questions and evidence needs. During development, it connects data and model choices to hazards, event sequences and risk controls. At release and deployment, it tests whether the evaluated configuration and operating context support the safety conclusions. During production and post-production activity, it directs monitoring and change review towards drift, retraining, site differences, feedback loops and rollback. Those interfaces align directly with the entity relationships and governance gates proposed in this paper.

The Technical Specification also clarifies a boundary that must be respected: it does not apply to machine-learning-enabled medical devices employing large language models or generative AI [6]. The entity architecture remains usable for those systems, but the medical-device risk method requires additional controls and emerging guidance. A governance claim should not imply that adoption of ISO/TS 24971-2 alone covers generative-AI behaviour.

C. Comparative allocation of risk concepts

Table 5 distinguishes the decision scope, unit of analysis, control evidence and feedback mechanisms of the two risk perspectives, and defines the rule by which they remain connected.

D. A worked traceability example

Consider an AI-assisted histopathology classifier intended to highlight image regions for specialist review. A Dataset contains images from several laboratories, but one scanner type used in a target hospital is under-represented. The defect is not merely a property of the Dataset: materiality depends on the exact Dataset version, preprocessing pipeline, model version and target AI Use Case. The initial AI Risk should therefore be attached to the

Dataset-Model Association and propagate to the model and use context.

The AI Risk can be stated as follows: under-representation of images acquired using scanner type B may cause reduced sensitivity in the target deployment environment, with disproportionate performance effects for the patients served by that site. The broad assessment considers performance, fairness, clinical workflow, regulatory compliance and stakeholder impact. The linked ISO 14971 analysis traces how a false-negative output, combined with reliance on the highlighted regions, could contribute to a missed finding, delayed diagnosis and patient harm.

Risk controls may include acquisition of representative data, predefined subgroup acceptance criteria, site-specific independent testing, input-quality checks, user information, mandatory specialist review, deployment restriction and monitoring stratified by scanner type. The QMS remains authoritative for the corresponding Design Inputs, Design Outputs, verification, validation, usability evidence, labelling, release approval and post-market plan. The AIMS links those records to the AI Risk and preserves why the controls were selected.

If monitoring later detects a performance shift, the signal challenges the same assumptions. It can trigger reassessment of the Dataset, Dataset-Model Association, AI Model, AI Use Case, AI Risk and ISO 14971 risk-control effectiveness. Depending on severity and uncertainty, the governance decision may maintain use, impose a restriction, suspend the affected configuration, roll back the model, initiate CAPA or require retraining and revalidation. The trace remains continuous because development and post-market evidence refer to the same governed entities.

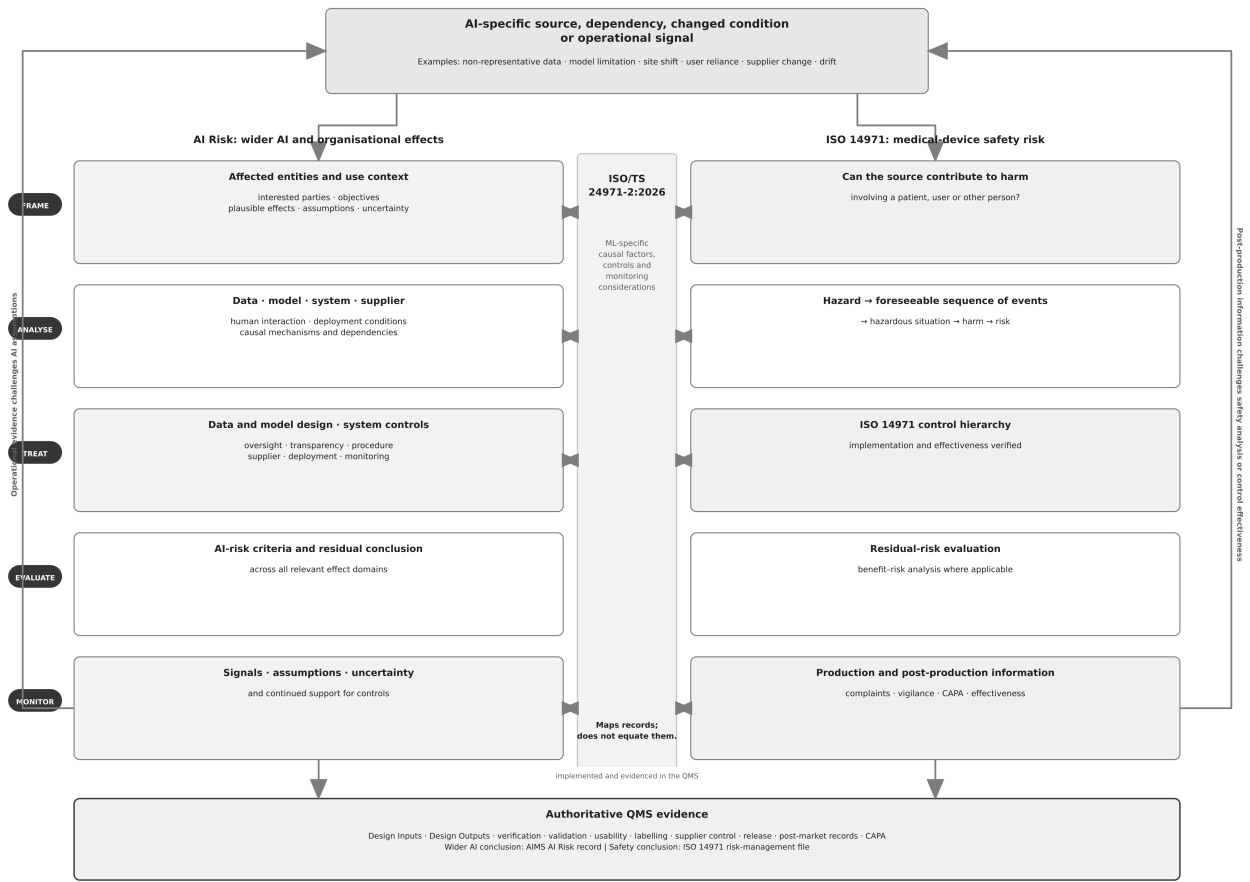


Fig. 3. ISO/IEC 23894 informs the wider AI-risk method, while ISO/TS 24971-2 directs ML-specific causal factors into the ISO 14971 process. The architecture coordinates sources, treatments, evidence and feedback without merging scope, criteria, records or decision authority.

Table 5. Allocation of risk concepts across the integrated architecture

Question	AI Risk under ISO/IEC 23894 and the AIMS	Medical-device risk under ISO 14971	Integration rule
What is affected?	Organisational objectives, system performance, persons, groups, society, rights, security, privacy and other relevant domains	Safety of patients, users and other persons through device-related harm	Keep both scopes explicit; do not discard non-safety AI effects
What is the unit of analysis?	AI System, AI Use Case, AI Model, Dataset, association, supplier or governance process	Hazard, sequence of events, hazardous situation, harm and associated risk	Map causal AI sources to the applicable ISO 14971 item rather than treating the records as identical
How is context represented?	Intended and foreseeable use, stakeholders, environment, data distribution, human interaction and organisational role	Intended use, reasonably foreseeable misuse, device characteristics and use-related safety	Reuse the approved intended-purpose and use-context evidence; preserve the different decision questions
Where are controls implemented?	Data, model, system, human oversight, supplier, procedure, information, monitoring or governance decision	Inherent safety by design, protective measures and information for safety, supported by verification	Link each AI treatment to the authoritative QMS control and effectiveness evidence
How are criteria applied?	Organisation-defined AI-risk criteria appropriate to affected objectives and consequence domains	Approved medical-device safety-risk criteria and benefit-risk provisions	Do not translate scores mechanically; preserve each evaluation and connect the factual basis
How is change handled?	Reassess affected entities, assumptions, impacts and relationships	Review risk analysis, control effectiveness, residual risk and benefit-risk as applicable	One change event can trigger both reviews; conclusions remain in their authoritative records
What closes the loop?	Monitoring, impact review, incidents, complaints, audit and management review	Production and post-production information, PMS, vigilance, CAPA and risk-management review	Route shared signals once, then update all affected records through controlled links

VI. INTEGRATION WITH THE ISO 13485 QUALITY MANAGEMENT SYSTEM

A. The AIMS as an integrated governance layer

The AIMS should be integrated with an ISO 13485-compliant QMS, not deployed as a competing quality system. ISO 13485

remains authoritative for controlled documents and records, competence, supplier control, design and development, product realisation, nonconforming outputs, complaints, CAPA and change [13]. IEC 62304 remains authoritative for software development and maintenance processes [14]. ISO 14971 remains authoritative

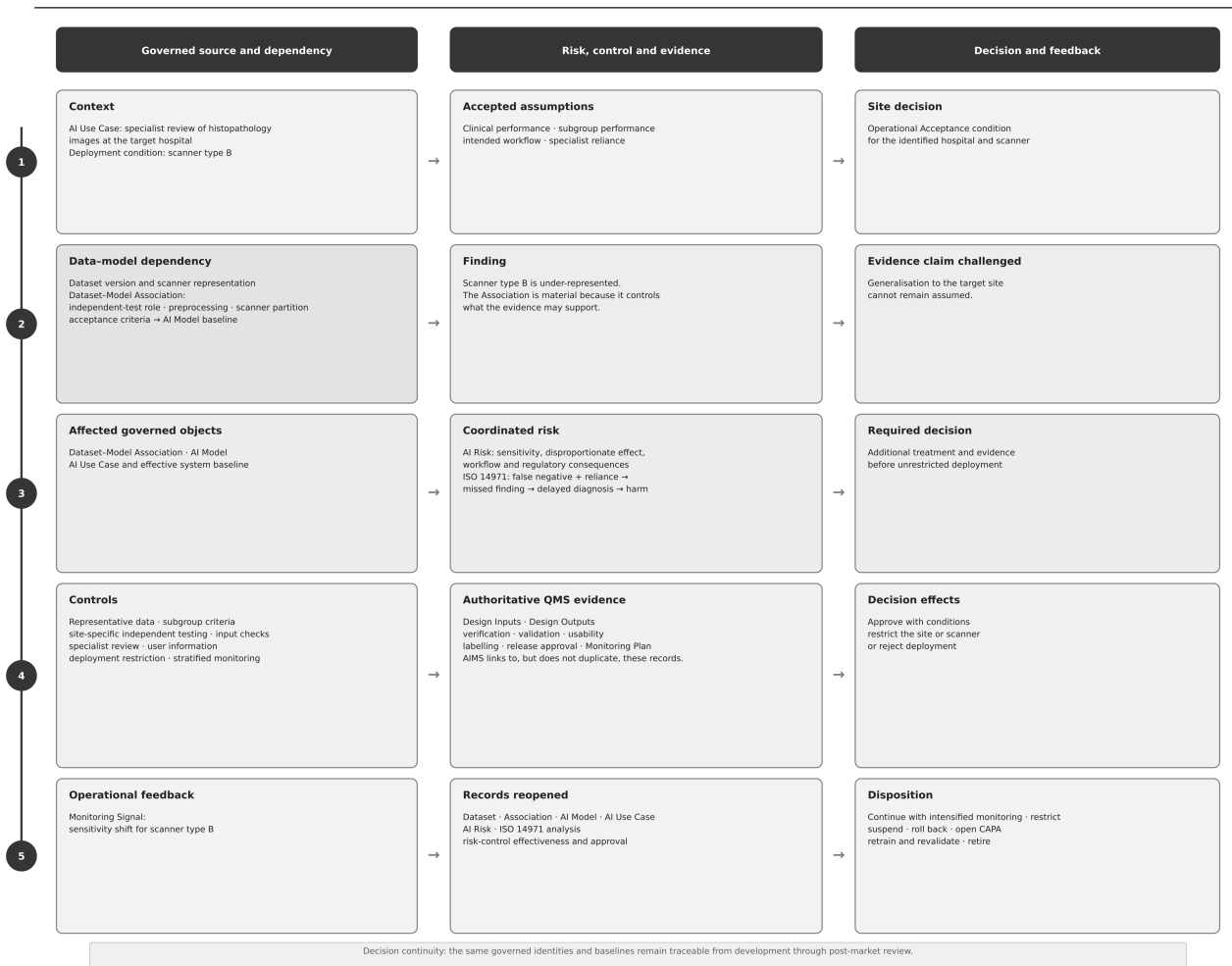


Fig. 4. The worked example preserves one trace from a site-specific data limitation to the model and use context, coordinated AI and medical-device risk conclusions, authoritative control evidence and the operational signal that may reopen approval.

tive for the device risk-management file [7]. The AIMS defines how AI-specific objects, assumptions, decisions and risks invoke, consume and remain traceable to those processes. This allocation is consistent with prior proposals that identify integration into established quality infrastructure as a practical means of reducing redundancy while supporting end-to-end oversight [21,22].

The same allocation now has a direct United States regulatory connection. FDA's Quality Management System Regulation became effective on 2 February 2026 and incorporates ISO 13485:2016 by reference, subject to the definitions and supplemental requirements retained in 21 CFR Part 820 [53]. A manufacturer implementing an authorised PCCP therefore does so through the same design, change, record and risk-based quality system that controls the device baseline. The PCCP does not create a separate execution system.

This allocation is consistent with the European regulatory direction. MDCG 2025-6 states that AI Act documentation and procedures may be integrated with those already established under the MDR or IVDR, and it emphasises a single set of technical documentation for high-risk medical-device AI [12].

Integration does not mean that every AI obligation is already satisfied by ISO 13485. Data governance, AI-system logging, model-specific traceability, fundamental-rights concerns and certain transparency and human-oversight duties require additional content. The control mechanism should be added to the existing process wherever that process is the correct authority.

For example, a model-development procedure may be governed by an AI-specific SOP, but an approved model baseline that forms part of the device design is controlled as a Design Output and its verification remains under design control. Dataset acquisition may invoke supplier qualification, data-protection controls and data-quality acceptance. A drift threshold may be defined in an AI Monitoring Plan, while an excursion enters the QMS through feedback, complaint handling, nonconformity, CAPA, regulatory reporting or change control according to its nature. The AIMS identifies why the event matters to the AI configuration and which governance decision is affected; it does not create an alternative route for handling the quality event.

B. System-of-record allocation

A practical implementation should define the system of record for each evidence class. The AIMS should own the AI inventory, entity identity, relationship graph, governance status, AI Risk records, impact decisions, AI-specific monitoring triggers and gate outcomes. The QMS should own approved controlled documents, design records, safety risk files, quality events, CAPA, suppliers, release records and training evidence. The ISMS should own information-security risk and organisational security controls; the PIMS should own personal-data processing, lawful basis, data-subject rights and privacy risk; and the applicable product-security process should remain authoritative for secure product-development evidence [18-20]. Links should be bidirectional where operational follow-up is required.

This allocation prevents the AIMS from becoming an uncontrolled index. A link is effective only when it identifies the current approved record, preserves its meaning and is reviewed when either side changes. Governance procedures should therefore define ownership of broken links, superseded evidence and cross-system change notifications. The principle is simple: evidence remains where it is authoritatively controlled, while the AIMS maintains the AI-specific context needed to interpret that evidence.

C. Life-cycle allocation and record continuity

Integration should be defined across the life cycle rather than only at the point of release. Table 6 shows the minimum allocation. The purpose is not to create an AI copy of each QMS phase. It is to ensure that each quality decision can identify the AI configuration and assumptions to which it applies, and that each AI governance decision invokes the established process competent to create or change regulated evidence.

VII. LIFE-CYCLE GOVERNANCE AND DECISION GATES

The entity model becomes operational through governance gates. A gate is a controlled decision based on defined criteria, accountable authority and current evidence. It is not necessarily a meeting and should not be confused with a software-development phase. Different entities may progress asynchronously, and the outcome of a gate may route work to design control, safety risk management, supplier control, privacy, security, CAPA or another competent process.

AI System Readiness is deliberately earlier than formal model approval. Exploratory work may precede it, but regulated development should not depend on an undefined system purpose or an unidentified use context. The gate creates a point at which the organisation accepts responsibility for the proposed AI-enabled function and defines the processes and evidence needed to govern it.

Product release and Operational Acceptance are separated because evidence that supports a centrally approved baseline may not establish suitability for every deployment. Local population, acquisition equipment, workflow, integration, user competence and oversight can change exposure and control effectiveness without changing the model. Operational Acceptance records

the conditions under which the released baseline may actually be used.

Monitoring Review is both scheduled and event-driven. A fixed annual review cannot be the sole control where a material drift indicator, complaint, incident, supplier notification or data change becomes known earlier. The monitoring design should therefore identify which assumptions are being tested, the threshold or qualitative evidence that challenges them, the accountable recipient and the action expected when the signal is confirmed.

Change classification follows impact analysis, not the other way round. Retraining creates a new model baseline even if architecture and interface are unchanged; relabelling or repartitioning a Dataset can invalidate evaluation evidence; and a workflow change can alter human reliance without changing software. The relationship model makes the affected path visible before the organisation decides how much verification, validation, safety review or regulatory action is required.

Restriction, Suspension, Rollback, Resumption and Retirement remain distinct decisions even though they share one disposition gate. Restriction narrows authorised use; Suspension temporarily prevents it; Rollback activates a previously approved baseline; Resumption confirms that explicit recovery criteria have been met; and Retirement ends authorisation permanently. Retirement must also identify deployed configurations, required communications, residual post-market or field obligations, record ownership and the evidence that must remain retrievable after active support ends.

VIII. CONTROLLED EVOLUTION AND PCCP OPERATIONALISATION

A. Controlled Evolution as a governance capability

Change is not an exceptional event in an AI-enabled medical device. New data become available, populations and acquisition equipment change, labels are corrected, monitoring identifies drift, suppliers revise dependencies, clinical practice evolves and performance objectives are refined. Preventing all post-authorisation change would cause a system to become progressively less representative of its operating environment. Permitting change merely because the manufacturer expected the system to evolve would create the opposite failure: the approved device could become an increasingly weak reference for what is actually deployed.

This paper uses *Controlled Evolution* to describe the governance capability that resolves that tension. The term is a design construct, not a defined regulatory term. It means that the organisation determines in advance why an AI System may change, which kinds of change are contemplated, which characteristics must remain invariant, which methods and evidence will govern implementation, which conditions require reassessment, and who may authorise a revised baseline. The capability begins during design and regulatory planning, remains active during operation and continues until the system is retired. It is therefore wider than retraining and wider than a single change-control procedure.

A PCCP is one important regulatory expression of Controlled Evolution, but the terms are not synonymous. FDA describes a

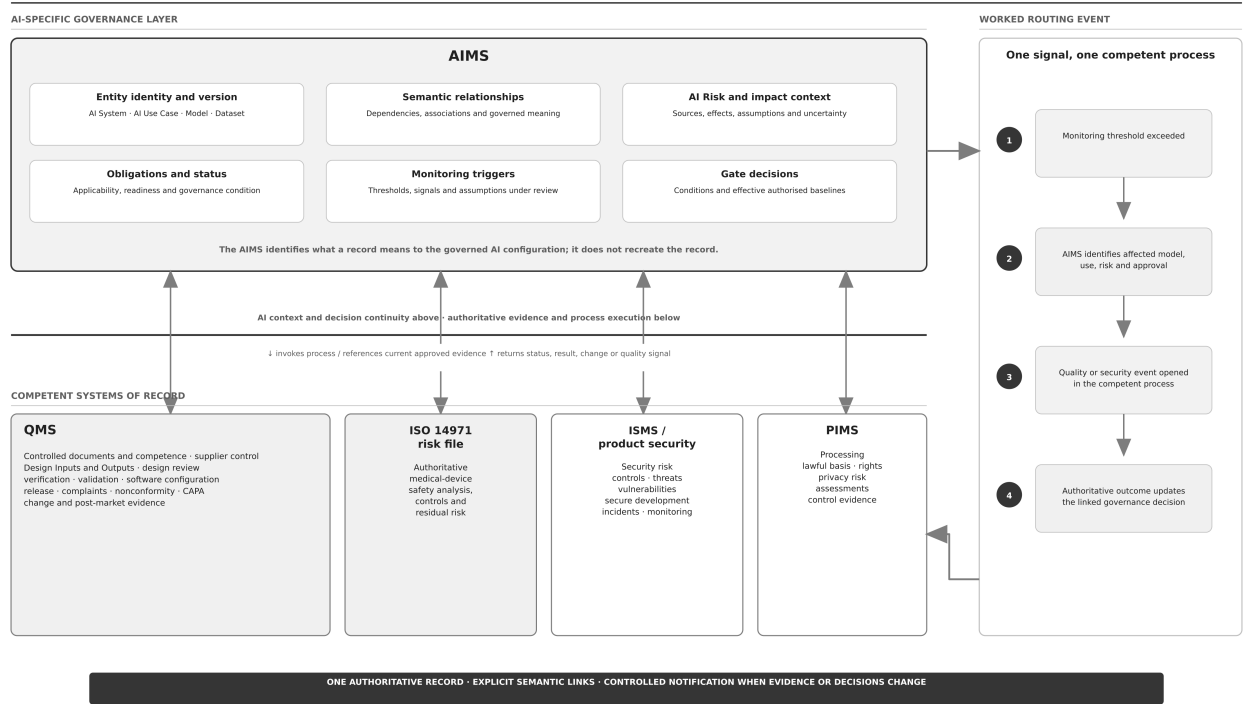


Fig. 5. The AIMS owns AI-specific identity, relationships, risk context and governance decisions. Established management systems remain authoritative for product, safety, security and privacy evidence; bidirectional links preserve meaning, status and change propagation without duplication.

PCCP through three interrelated components: the Description of Modifications, the Modification Protocol and the Impact Assessment. An authorised PCCP allows specified modifications to be implemented without a new marketing submission when each modification is both described in the PCCP and implemented in accordance with its authorised protocol [48]. Controlled Evolution supplies the enduring organisational structure needed to create that plan, maintain its relationship to the approved system and execute it through the QMS. It also governs changes that are outside the PCCP, changes in markets where the PCCP has no authorising effect, and decisions not to proceed with a change.

The distinction prevents two category errors. First, authorisation of a PCCP is not advance approval of every implementation produced under it. Each modification still has to be developed, verified, validated, risk assessed, reviewed, documented and released through the manufacturer’s quality system [48]. Secondly, an internally acceptable evolution boundary is not a regulatory authorisation. The same proposed change may be covered by an authorised PCCP in the United States, require notified-body assessment or another regulatory route in the European Union, and be treated differently in a third market. Controlled Evolution must therefore preserve a common technical change record while maintaining market-specific authorisation overlays.

B. Defining the evolution boundary

The evolution boundary is approved before individual modifications are executed. It should state the objectives served by future change, such as recovery from population drift, reduc-

tion of subgroup performance disparity, recalibration within an approved operating range or improvement of robustness to an identified acquisition condition. An objective such as “improve performance” is too broad to control development. The intended improvement should identify the affected performance characteristic or use condition and the constraints that cannot be traded away to obtain it.

The boundary should then distinguish permitted, conditionally permitted and excluded change. Permitted change might include retraining the same model architecture on a released Dataset version drawn from defined sources, or adjusting a decision threshold within a validated numerical range. Conditions may include use of a frozen training procedure, independent evaluation data, specified subgroup analyses, unchanged human oversight and predefined acceptance criteria. Exclusions should identify changes that alter the governance basis itself, such as a new intended purpose, clinical claim, input modality, model architecture, autonomy level, user group or market. Those changes return to the earlier readiness, design and regulatory-planning decisions appropriate to their effect.

A robust boundary is multidimensional. A numerical threshold alone does not define it. The boundary comprises the proposed change, the eligible source data, the development method, the characteristics required to remain invariant, the performance and safety acceptance criteria, the affected population and environment, the risk controls, the update and communication process, and the monitoring and rollback arrangements. A candidate change is within the boundary only when the complete combina-

Table 6. AIMS integration with the medical-device quality life cycle

Life-cycle activity	Authoritative QMS process and evidence	AI-specific governance contribution	Decision-continuity requirement
Concept and planning of product realisation	Quality planning; regulatory strategy; initial product and project planning	Define the AI System boundary, proposed AI Use Cases, accountable roles, jurisdictions, preliminary classification, affected persons, objectives and initial AI Risks	Preserve the origin of the intended purpose and the assumptions used to decide whether regulated development may proceed
Design and development planning and inputs	Design and Development Plan; Design Inputs; risk-management planning; clinical or performance-evaluation planning; software planning	Translate each authorised AI Use Case into data, model, performance, fairness, transparency, human-oversight, monitoring and change requirements	Link every AI-specific requirement to its use context, risk rationale and intended verification or validation method
Data acquisition and model development	Supplier control; purchasing information; Design Outputs; software configuration; development records; data-protection and security processes	Control AI Model and Dataset identity, provenance and status; create Dataset-Model Associations; record experiment-to-baseline decisions; identify new AI Risks	Distinguish exploratory work from controlled development and preserve which exact data and model configuration produced each claimed result
Design review, verification and validation	Design reviews; software verification; system validation; usability engineering; clinical or performance evaluation; risk-control verification	Confirm that evidence addresses the approved use, population, operating conditions, model baseline and Dataset-Model Associations; maintain limitations and unresolved uncertainty	Prevent evidence generated for one configuration, population or workflow from being transferred to another without assessment
Transfer and release	Design transfer; release criteria; approved device and software configuration; labelling; risk-management report	Conduct Model and System Release review against current AI Risks, approved entities, monitoring readiness and applicable governance conditions	Record the exact released baseline, the evidence set relied upon and any conditions or limitations attached to approval
Deployment and operational acceptance	Installation and servicing controls; configuration records; training; distribution and deployment procedures	Confirm local population, equipment, workflow, user competence, oversight, infrastructure, input controls, monitoring and rollback arrangements	Separate product release from authorisation for a specific operating environment and retain site-specific acceptance conditions
Production and post-production activity	Feedback; complaint handling; regulatory reporting; analysis of data; nonconformity; CAPA; post-market surveillance and vigilance under applicable law	Route performance, drift, subgroup, misuse, supplier and contextual signals to the affected entities, Associations, AI Risks and gate decisions	Make each signal capable of challenging the original approval assumption and triggering both AI-governance and safety-risk review where applicable
Change, retraining and corrective action	Design change; software maintenance; supplier change; nonconformity; CAPA; verification, validation and release as required	Check the Controlled Evolution boundary and any market-specific authorised PCCP; perform relationship-based impact analysis before classifying the change; create new entity versions or Associations; determine which gates and monitoring baselines must be repeated	Do not classify significance from a version label or PCCP heading alone; base the decision on affected uses, assumptions, risks, controls, evidence, protocol conformance and jurisdiction-specific authorisation
Restriction, suspension, rollback, resumption and retirement	Field action; communication; distribution, installation and servicing records; regulatory reporting; discontinuation and support controls	Define the scope, effective configuration, deployed population, conditions, accountable decision maker and follow-up obligations for each disposition	Preserve which baseline may or may not be used, where it is deployed, and which evidence supports any later resumption or replacement
Post-retirement records	Document and record control; technical-documentation and regulatory-retention arrangements	Maintain an index of the retired AI System, use contexts, final model and data lineage, risk and monitoring history, dispositions and links to archived evidence	Retirement ends authorisation, not traceability; records must remain retrievable for safety investigation, field action, legal accountability and the applicable retention period

tion remains within those conditions. This prevents a nominally familiar change, such as “retraining”, from being treated as authorised when it uses a new data source, changes preprocessing, departs from the approved evaluation method or produces an unacceptable subgroup effect.

The boundary should also contain regulatory reassessment triggers. These are conditions that require the organisation to stop relying on the existing route and reconsider the relevant authorisation. Examples include a change outside a declared limit, failure to meet acceptance criteria, a serious incident attributable to model output, a sustained monitoring breach, a new market or user group, a changed clinical claim, a changed level of human oversight, an unplanned data source, or a deviation from the authorised protocol. The trigger does not determine the regulatory conclusion by itself. It ensures that the conclusion is made explicitly by the competent function before deployment.

C. The Planned Modification as the unit of execution

The evolution boundary governs a class of contemplated change; a Planned Modification governs one specific proposal. It is a transactional control record linked to the affected AI System, AI Use Cases, AI Models, Datasets, Dataset-Model Associations, AI Risks and product baseline. It should have a stable identifier,

revision history, accountable owner and state. A practical state model distinguishes Draft, Assessed, Authorised, Implemented, Withdrawn and Superseded. “Implemented” should mean that the approved output has completed the required release process, not merely that retraining or coding has finished.

Each Planned Modification should record the initiating objective or Signal, the precise scope, the characteristics expected to change, the characteristics required to remain unchanged, anticipated benefits, affected records and markets, and the initial boundary assessment. It should then carry or link the modification-specific protocol, acceptance criteria, impact assessment, evidence and Change Authorisation decision. A model or Dataset version produced during execution is not the Planned Modification. The former is a controlled technical output; the latter is the governed intent, method and decision path that explains why the output exists and whether it may be released.

Impact assessment precedes final change classification. The assessment should identify effects on intended purpose, users, affected persons, clinical workflow, input and output meaning, human oversight, data governance, model behaviour, software and infrastructure, labelling, cybersecurity, privacy, regulatory classification and existing controls. It should evaluate benefits,

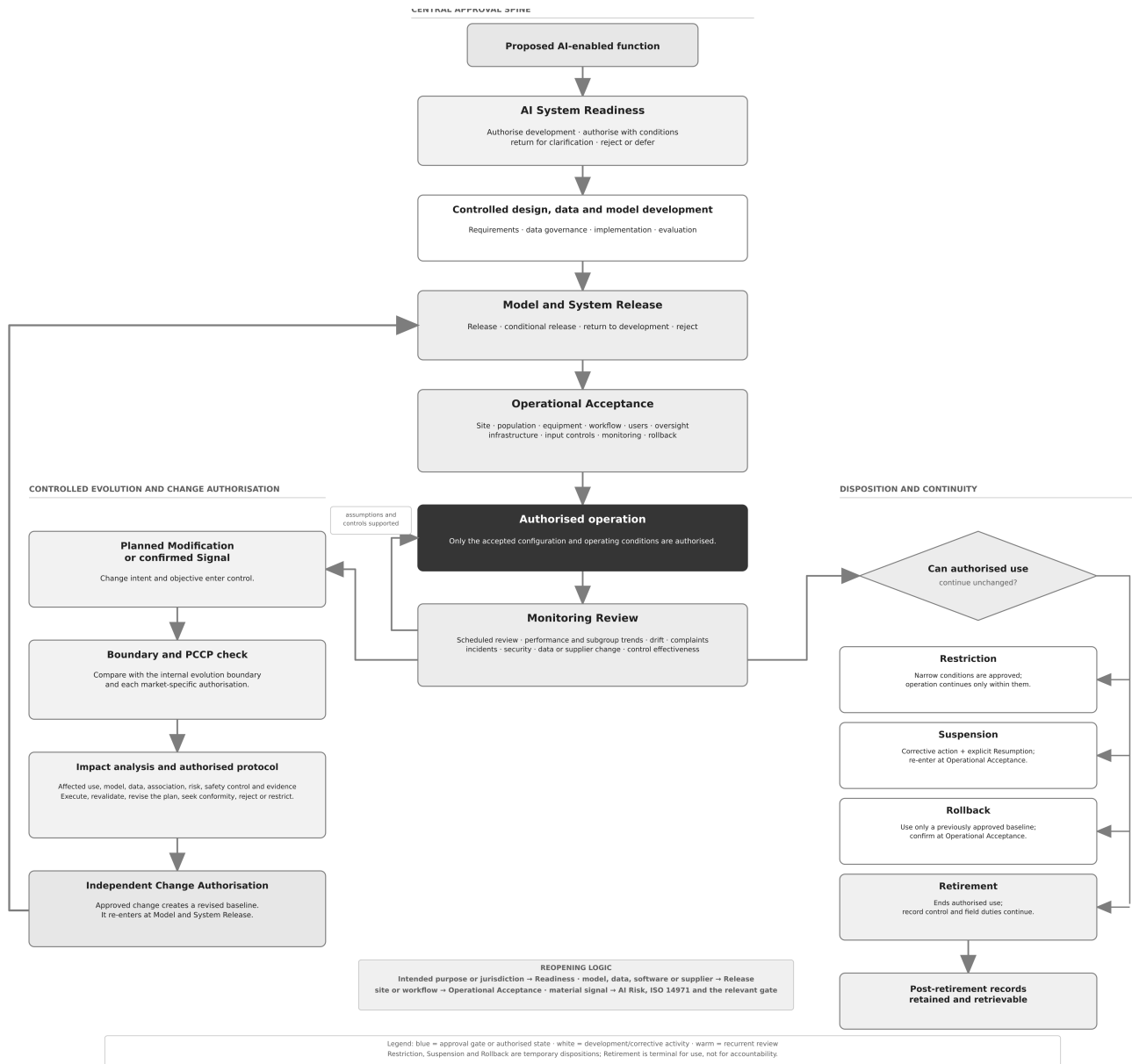


Fig. 6. The gates form a controlled decision network rather than a fixed sequence. Monitoring may preserve authorisation or initiate a Planned Modification, while Controlled Evolution determines whether the change can follow an authorised protocol or requires a new regulatory route. Changed and resumed configurations re-enter at the gate appropriate to the evidence and operating context.

new or modified AI Risks, corresponding ISO 14971 implications, required mitigations and the collective effect of this modification with changes already made under the same evolution strategy. Cumulative analysis matters because several individually acceptable adjustments can collectively move performance, population coverage or operating conditions beyond the evidence considered at initial authorisation.

D. From controlled records to a PCCP

A PCCP submitted to FDA should remain a controlled document in the marketing submission and QMS. The implementation principle proposed here is narrower: its substantive sections should be assembled from controlled source records instead of being authored and maintained as a second independent account

of the system. The approved AI System and AI Model provide the device and model description. Planned Modifications provide the Description of Modifications. Their protocols provide the data-management, retraining, performance-evaluation and update procedures. Their impact assessments provide the benefit-risk and mitigation analysis. The system-level evolution boundary provides constraints and reassessment triggers, while linked monitoring, rollback, labelling and transparency controls complete the applicable plan content.

This separation improves both authorship and review. A PCCP section can contain polished text yet depend on a Planned Modification whose update procedure, acceptance criterion or cumulative-impact assessment is missing. Document complete-

ness alone would conceal that defect. A record-derived approach can show two different measures: whether every PCCP section contains text, and whether every source record contributing to the plan is complete and approved. Export or submission should be held when either condition fails.

Assembly does not remove document control. At the submission gate, the generated PCCP is reviewed, approved and frozen as a market-specific regulatory baseline. Its source identifiers, versions and effective date are retained. Subsequent editing of a source record does not silently change the authorised plan. A proposed revision to the plan creates a new controlled version and follows the regulatory route applicable to modifying the PCCP. FDA states that modifications to an authorised PCCP will generally require a marketing submission for the modified device and plan [48]. This is why the system must distinguish the current internal evolution design from the PCCP version actually authorised for a particular device and market.

E. Executing an authorised PCCP

Operationalisation begins when a product objective, scheduled evolution cycle, monitoring Signal, complaint, CAPA, supplier notice or other evidence creates a candidate Planned Modification. The first decision is not whether the proposed model appears better. It is whether the change is sufficiently defined to assess and whether the device may continue operating unchanged while that assessment proceeds. A confirmed safety or performance Signal may require immediate restriction or suspension independently of any later decision to modify the system.

The candidate is then compared with both the internal evolution boundary and the authorised PCCP applicable to each market. For the United States, FDA considers a modification consistent with an authorised PCCP when it was specified in the Description of Modifications and is implemented in accordance with the Modification Protocol [48]. Both conditions are necessary. A retraining activity named in the plan is not covered if the manufacturer uses unapproved data, changes the method, omits required evaluation, fails an acceptance criterion or cannot follow an authorised procedure.

Once the regulatory route is established, execution remains within design and change control. The organisation releases the authorised Dataset versions, records their use through Dataset-Model Associations, executes the approved training or modification method, freezes the resulting AI Model candidate and performs the specified verification and validation. Results are assessed against predefined criteria overall and for every population, environment or device characteristic required by the protocol. Failures, anomalies and departures are recorded; they are not averaged away by an acceptable headline metric.

The completed evidence then enters Change Authorisation. The review confirms that the candidate matches the Planned Modification; the authorised method was followed; required evidence is complete; acceptance criteria are satisfied; AI Risk and ISO 14971 conclusions are current; controls are implemented and verified; cumulative impact remains acceptable; labelling, transparency and user communication are ready; and the market-

specific regulatory route remains valid. The author of the modification should not be the sole approver. Approval, conditional approval, rejection and withdrawal should be explicit outcomes, with conditions and signatures retained as part of the decision record.

An approved candidate proceeds to the normal release and deployment controls. The new effective baseline identifies the model, data, software, configuration, labelling and applicable use conditions. Operational Acceptance is repeated wherever the modification can affect local population fit, equipment compatibility, workflow, training, oversight, infrastructure, input quality or rollback readiness. Initial post-update monitoring may be more intensive than routine monitoring because the modification's prospective assumptions are now being tested in use. The previous approved baseline remains eligible for rollback only for the period and environments in which its safety, compatibility and support remain established.

F. Boundary and deviation decisions

The operational decision cannot be reduced to "inside PCCP" or "outside PCCP". The internal evolution boundary and a regulatory PCCP answer different questions. The internal boundary asks whether the change is consistent with the organisation's approved strategy and technical control envelope. The PCCP asks whether a particular regulator has prospectively authorised the specified change and method as part of a marketing authorisation. Both must be resolved for the applicable market.

The matrix makes failure safe. A change that does not meet the authorised method or acceptance criteria does not become acceptable because its purpose was beneficial or because a similar modification appeared in the plan. FDA guidance treats failure to follow the PCCP or inability to follow it as a potential deviation, and significant modifications that are not specified or not implemented in accordance with the authorised PCCP are likely to require a new marketing submission [48]. Controlled Evolution converts that principle into a release-blocking rule attached to the specific Planned Modification and affected baseline.

G. Monitoring, risk and cumulative control

Monitoring closes the Controlled Evolution loop. Each important monitoring measure should be traceable to an assumption, risk control, performance requirement, use condition or acceptance conclusion. A threshold breach therefore has a defined semantic effect: it challenges identified decisions and creates an accountable response. The response may preserve authorisation after investigation, intensify monitoring, open an AI Risk, update the ISO 14971 analysis, initiate CAPA, create a Planned Modification, impose Restriction, execute Rollback or suspend use. The action depends on the evidence and potential consequence, not on the fact that a PCCP exists.

Acceptance criteria and monitoring thresholds should be connected but need not be numerically identical. Pre-release evaluation may require a performance margin that accounts for uncertainty and expected operational variability, while post-market monitoring may use alert and action limits designed for different sample sizes and response times. What matters is that the rela-

Table 7. Allocation of Controlled Evolution records to PCCP content and QMS execution

Controlled Evolution element	Principal controlled source	PCCP or regulatory representation	Operational control consequence
Evolution objectives and strategy	AI System evolution statement linked to intended purpose, post-market plan and regulatory strategy	Rationale and context for the proposed modification types	Future changes are judged against a stated improvement or maintenance objective rather than a generic desire to update the model
Permitted boundary, invariants and exclusions	Approved system-level evolution boundary	Change boundary and limits supporting the Description of Modifications and Impact Assessment	A candidate change cannot enter the abbreviated execution route unless the complete combination of scope, method and invariants is within the approved boundary
Specific proposed change	Planned Modification	Description of Modifications	Each change is individually identifiable, sufficiently specific, linked to its rationale and traceable to the affected device function and performance characteristics
Data-management method	Planned Modification protocol linked to Dataset versions and Dataset-Model Associations	Modification Protocol: data-management practices	Eligible sources, partitions, provenance, quality controls, independence, leakage prevention and transformation rules are defined before data are used
Development or retraining method	Planned Modification protocol linked to the approved AI Model and software baseline	Modification Protocol: retraining practices	The development method, frozen or variable parameters, reproducibility controls and permitted departures are explicit
Performance evaluation and acceptance	Planned Modification protocol linked to independent evaluation Associations and Design Inputs	Modification Protocol: performance-evaluation activities and acceptance criteria	Evidence is generated against criteria fixed before results are known and is attributable to the exact candidate baseline and intended-use populations
Update, communication and transparency	Planned Modification protocol linked to release, deployment, training and labelling records	Modification Protocol: update procedures	Release, distribution, user notification, labelling, training, technical deployment and post-update monitoring are executed through the competent QMS processes
Benefit, risk and cumulative effect	Modification-specific impact assessment linked to AI Risks and the ISO 14971 file	Impact Assessment	New and modified risks, mitigations, residual conclusions and the collective effect of prior changes are reviewed before authorisation
Monitoring and reassessment	Monitoring Plan, thresholds, Signals and review commitments	Real-world monitoring and conditions supporting implementation of the plan	Operational evidence tests the assumptions used to authorise the modification and can initiate restriction, rollback, CAPA or another Planned Modification
Rollback	Approved previous baseline, deployment inventory and rollback procedure	Update and risk-mitigation content, where applicable	Rollback is possible only to an identified approved configuration and is itself authorised, recorded and verified in the receiving environment
Change Authorisation	Independent AI Review linked to complete evidence and QMS approvals	Evidence that the authorised plan was followed; not a substitute for the PCCP itself	No candidate is released merely because it is described in a PCCP; an accountable approver confirms protocol conformance, acceptance, risk disposition and regulatory route
Implemented configuration and chronology	Released AI Model, Dataset and software versions; release record; deployment record; audit trail	Updated labelling, submission records and later regulatory reporting as applicable	The organisation can reconstruct which authorised change produced each deployed baseline, where it was deployed and which PCCP version governed it

Table 8. Decision matrix for a proposed modification

Internal Controlled Evolution position	Market-specific PCCP position	Required governance response	Release consequence
Inside the approved internal boundary	Specified in an authorised PCCP and executable in accordance with its protocol	Complete the modification-specific impact assessment; execute the authorised protocol through the QMS; obtain Change Authorisation; update labelling and records as applicable	May use the authorised PCCP route in that market if all criteria and legal conditions remain satisfied; no automatic right to release is created
Inside the approved internal boundary	Not included in an authorised PCCP, or no PCCP applies	Use the ordinary design-change and regulatory-assessment route; determine whether a new submission, notified-body approval or other conformity activity is required	No reliance on the PCCP; release waits for the applicable internal and external authorisations
Outside the approved internal boundary	Apparently described in a PCCP	Stop the abbreviated route and reconcile the inconsistency; reassess intended purpose, system boundary, risks, evidence and the authorised plan	No release until the internal governance basis and regulatory position are both re-established
Outside the approved internal boundary	Outside or not addressed by the PCCP	Return to the earliest affected readiness, design, risk and conformity decisions; create a new regulatory strategy	Treated as a new or materially changed configuration according to the applicable requirements
Described in the authorised PCCP	Protocol cannot be followed, acceptance criteria are not met or a material deviation occurs	Do not deploy under the PCCP; investigate the deviation, assess safety and effectiveness, invoke nonconformity or CAPA as appropriate, and determine whether a new marketing submission or PCCP revision is required	Candidate remains unauthorised; operational controls for the existing baseline continue or are restricted according to risk

tionship and rationale are documented. A monitoring threshold should identify which acceptance assumption it protects and what action follows when the threshold is reached.

Risk control also extends across modifications. The Planned Modification identifies the AI-specific source or changed condition and its affected objectives. Safety-relevant effects are evalu-

ated within the ISO 14971 risk-management process, including new event sequences, hazardous situations, control interactions and production or post-production information. Broader effects remain visible in AI Risk. The Impact Assessment then records how both authorities informed the decision without merging their scoring systems. This is particularly important when a modification improves average diagnostic performance while increasing subgroup disparity, automation bias, privacy exposure, cybersecurity dependence or referral volume.

Cumulative control prevents serial boundary erosion. Each Change Authorisation should compare the candidate with the immediately preceding version, the originally authorised baseline and prior modifications implemented under the same strategy or PCCP. Relevant measures include cumulative performance movement, population expansion, changed data provenance, repeated exceptions, evolving human reliance, accumulated labelling changes and the continuing validity of rollback. A sequence of bounded changes may require reassessment even when no single change exceeds a threshold.

H. Jurisdictional interpretation and portability

The FDA PCCP is the most developed medical-device implementation of predetermined change. Its current final guidance applies to AI-enabled devices reviewed through the 510(k), De Novo and PMA pathways, and describes prospective authorisation of specific modifications and the methods used to develop, validate and implement them [48]. The FDA route should not be described as a general international permission to update an AI-enabled device.

The EU AI Act contains a related but distinct concept. Article 43(4) addresses changes to high-risk AI systems that continue to learn after being placed on the market or put into service when the changes and their effect on performance were predetermined by the provider and assessed at initial conformity assessment [5]. MDCG 2025-6 states that predetermined changes for medical-device AI should be integrated into the MDR or IVDR change-management procedure and the technical documentation, while noting that dedicated IMDRF work was still under development at the time of publication [12]. The MDR or IVDR significance of a device change, notified-body expectations and applicable certificate conditions remain independently relevant. A United States PCCP may provide useful evidence, but it does not replace the European conformity decision.

The governance architecture should therefore maintain one Planned Modification and one technical evidence chain, supplemented by a jurisdictional assessment for each affected market. Each overlay identifies the authorised PCCP or equivalent plan, regulatory body, approval or clearance reference, effective version, permitted conditions, reporting or notification duties and conclusion for the proposed change. This avoids duplicating technical development while preventing the most permissive market decision from being applied globally.

Controlled Evolution remains valuable where no formal PCCP mechanism exists. It makes the organisation state what change it anticipates, what must remain invariant, which evidence will be required, which risks and approvals must be revisited, and

what forces a return to earlier governance. In other regulated sectors, the external authorisation column changes, but the internal need to connect prospective planning, controlled execution, monitoring and accountable release remains.

IX. APPLICATION BEYOND MEDICAL-DEVICE DEVELOPMENT

The architecture is applicable beyond medical devices because the six core entities express technology and governance distinctions that recur across regulated AI uses. A regulated application still has a system boundary, one or more contexts in which outputs acquire consequence, technical models, data assets, material data-model relationships and AI-related uncertainty. Readiness, approval, operational acceptance, monitoring, Controlled Evolution, restriction and retirement also remain meaningful capabilities or decisions. Planned Modifications, Signals and AI Reviews retain their operational purpose even where no sector-specific equivalent to a PCCP exists. Portability therefore begins with the entity and decision model, not with the direct transfer of medical-device requirements.

What changes is the authoritative sector layer and the meaning assigned to acceptability. In road vehicles, an AI Risk may need to map to functional-safety, safety-of-intended-functionality and AI-safety analyses, including sector-specific treatment of perception limits and foreseeable operating scenarios [51]. In financial services, the relevant authority may be a model-risk-management framework requiring independent validation, challenge, inventory control, ongoing monitoring and governance proportionate to materiality [52]. In employment or public administration, fundamental-rights, equality, procedural fairness and administrative-law controls may dominate the consequence analysis. In critical infrastructure or industrial automation, resilience, security, functional safety and continuity may be the primary authorities.

The adaptation should follow four steps. First, define the sector objectives and protected interests that determine what counts as material harm or unacceptable effect. Secondly, identify the management system, regulatory file or independent function authorised to make each decision. Thirdly, map AI-specific sources and dependencies to the sector's risk concepts without importing medical-device terminology where it does not fit. Fourthly, revise the evidence required at each gate, including validation independence, operational acceptance, incident reporting, change approval and record retention.

The medical-device case remains a demanding reference implementation because it requires disciplined separation between broad governance and product safety, traceable control effectiveness, controlled design evidence and feedback from production and post-production information. Those features expose whether an AI framework can connect high-level principles to configuration-specific decisions. They do not make medical-device criteria universally correct; they make the interfaces and authority boundaries explicit enough to be replaced.

Portability therefore has a clear limit. Definitions of harm, materiality, validation independence, acceptable residual risk, reporting, retention and regulator authority must be established

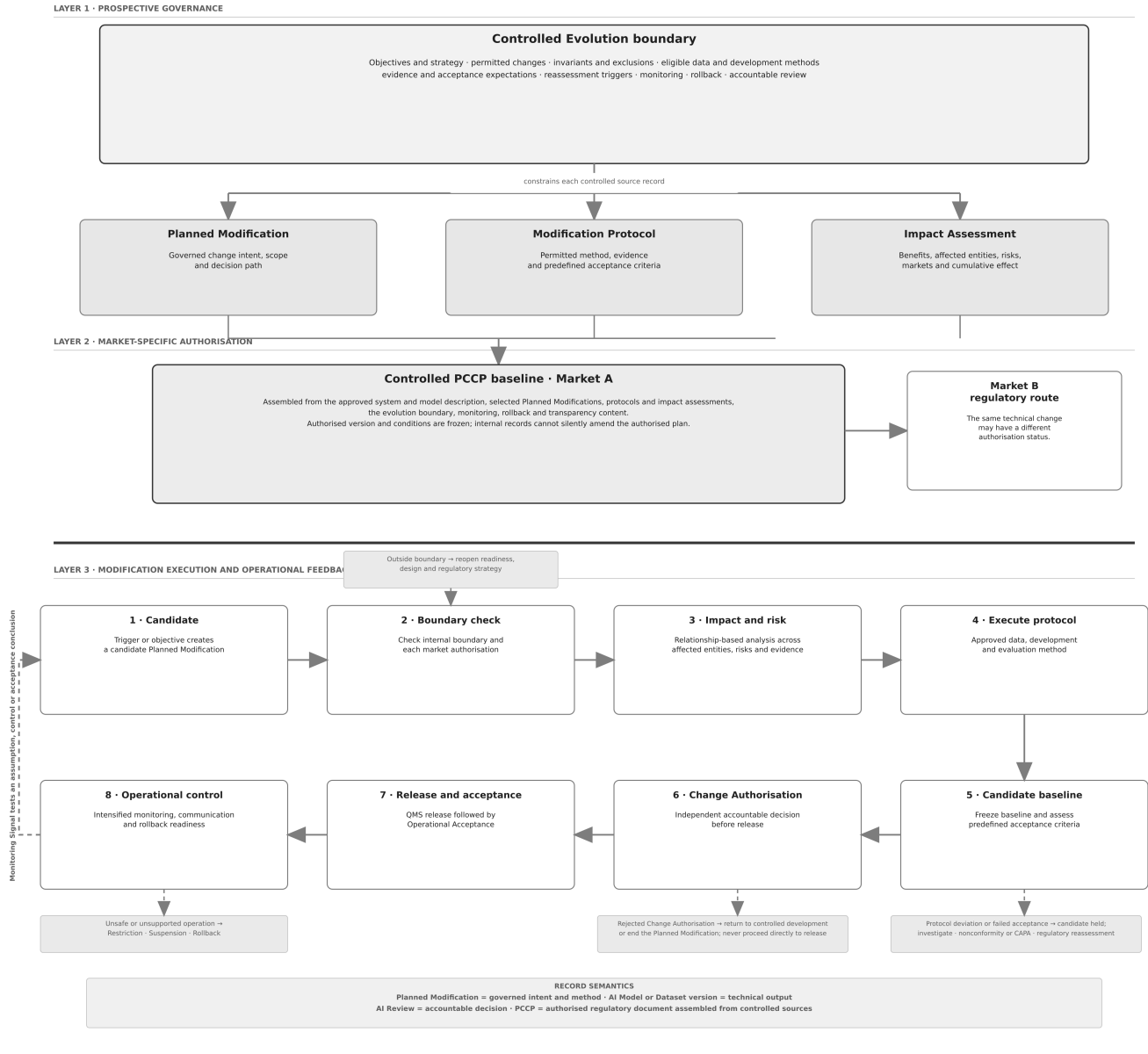


Fig. 7. Controlled Evolution defines the prospective change envelope and maintains the records from which a market-specific PCCP is authorised. Each modification remains subject to protocol conformance, predefined acceptance, independent Change Authorisation, QMS release and operational monitoring; the PCCP creates a regulatory route, not an automatic release decision.

from the receiving sector. The framework is portable because it identifies what must be reallocated and preserves the semantic links needed to do so. A claim that the medical-device implementation itself establishes compliance in another industry would be unjustified.

X. DISCUSSION

The central value of the architecture is decision continuity. Approval is treated as a conclusion about an identified system, authorised use, technical baseline, data lineage, set of assumptions, risk position and body of evidence. Those elements remain connected after approval. When any one of them changes, the organisation can identify which conclusions may no longer hold

and which competent processes must reconsider them. This is more useful than the general statement that governance should cover the full life cycle because it specifies the information structure required to make life-cycle control possible.

Decision continuity produces several practical effects. Reuse becomes governable because an AI Model or Dataset can retain one controlled identity while each use has its own context and approval. Change impact becomes transitive because a changed Dataset can be traced through its Associations to affected models, uses, risks and released systems. Risk analysis becomes more causal because a concern can attach to the technical or contextual relationship in which it originates. Assurance becomes more efficient because reviewers can navigate from an AI decision

Table 9. Sector adaptation of the reference architecture

Sector example	AI entities retained	Sector authority substituted for the medical-device layer	Principal adaptation required
Medical devices	All six entities, operational records, Controlled Evolution and all gates	ISO 13485 QMS, IEC 62304 software processes, ISO 14971 risk-management file and applicable medical-device regulation	Map AI sources to device hazards, hazardous situations, harms, risk controls, clinical or performance evidence and post-market duties
Road vehicles	All six entities and operational records; AI Use Case may be expressed through operational scenarios	Automotive quality, functional-safety, intended-functionality and AI-safety processes	Represent operational design domain, perception uncertainty, scenario coverage, fallback behaviour and vehicle-level safety case
Financial services	All six entities and operational records; AI Model may be the primary regulated unit	Model inventory, independent validation, prudential governance, consumer-protection and operational-risk processes	Add materiality tiers, independent challenge, limitations on use, performance stability, override analysis and validation independence
Employment or public administration	All six entities and operational records; AI Use Case and AI Risk become especially prominent	Equality, labour, data-protection, fundamental-rights, procurement and administrative decision processes	Emphasise affected-person analysis, procedural rights, contestability, transparency, human authority and discriminatory effect
Critical infrastructure and industrial automation	All six entities and operational records; system and supplier dependencies may require additional association types	Safety, security, resilience, asset-management and continuity processes	Add operational-state, fail-safe, adversarial, availability, recovery and supply-chain evidence to the applicable gates

to the approved evidence in the QMS, ISMS or PIMS without creating a second uncontrolled copy.

Controlled Evolution applies the same principle prospectively. It connects an intended future change to the authorised system boundary, affected technical objects, market-specific regulatory route, development method, acceptance criteria, risk position and required evidence before implementation begins. The PCCP is valuable because it can establish advance regulatory agreement for specific modifications and methods. Its operational reliability depends on the quality system maintaining those relationships after authorisation. A plan that cannot identify which Dataset and AI Model versions were used, whether the authorised method was followed, which acceptance criteria were met, who authorised release and where the resulting baseline was deployed does not provide effective change control merely because its narrative remains on file.

The architecture also changes the practical meaning of an AI inventory. A flat register can answer which AI systems exist and who owns them. A governed relationship model can answer which model baseline is authorised in which use context, which data and evidence support that authorisation, which limitations remain open, which risk controls are relied upon, and whether monitoring still supports the original decision. Algorithm registration, Model Cards, Model Facts and Datasheets provide valuable transparency [35-39], but they do not alone establish current authority or the propagation path from a changed dependency to a governance action.

The reviewed literature substantiates the need for life-cycle quality management, socio-technical context, monitoring and integration with existing health infrastructure [21-30,42-50]. The proposed contribution is narrower than claiming those principles as new. It specifies an operational division of labour and a relationship model through which they can be implemented: organisational AIMS controls are given a deeper AI-risk method; technical objects and data relationships are controlled explicitly; broader AI concerns are mapped to, but not conflated with, the

medical-device safety file; and product evidence remains under the authority of the established QMS.

That design also exposes three common failure modes. An inventory-only implementation identifies systems but cannot reconstruct the basis of approval. A document-only implementation contains extensive evidence but cannot perform reliable transitive impact analysis. A model-centric implementation controls performance metrics while losing the clinical use, human interaction and wider system dependencies that determine whether those metrics are meaningful. The proposed entities should be considered necessary only to the extent that they prevent one of these failures or support a defined decision.

Administrative proportionality is therefore essential. Every field, status, relationship and gate should support accountability, approval, impact analysis, monitoring, change control or evidence retrieval. The Dataset-Model Association justifies separate control because the relationship carries material many-to-many properties. A separate record for every experiment or minor analysis would not be justified where an approved development report already controls the evidence and the relevant baseline can be traced reliably. The architecture should scale the formality and review depth to intended purpose, novelty, uncertainty, regulatory significance and potential consequence.

Automation can support this proportionality but should not decide acceptability. Calculated readiness can identify missing evidence, stale reviews, broken links and inconsistent baselines. It cannot determine whether acceptance criteria were clinically meaningful, whether an evaluation population supports the intended use, or whether unresolved uncertainty is acceptable. A complete checklist can coexist with an unsafe product. Accountable approval therefore remains a substantive judgement supported by automation, not an outcome generated from record completion.

XI. LIMITATIONS AND FUTURE WORK

This paper proposes a conceptual architecture and does not report empirical evidence from certification, notified-body assessment, regulatory submission or comparative implementation. The narrative review was structured but not systematic; sources may have been missed, and the evidence of convergence should not be interpreted as proof that the proposed entity model is superior to alternative implementations. The worked example tests internal traceability, not real-world governance effectiveness.

Several design choices require empirical validation. The preferred cardinality of the relationships, the minimum information needed in a Dataset-Model Association, the appropriate granularity of AI Use Cases, the minimum viable Planned Modification record and the governance burden created by cross-system links should be tested across organisations of different sizes and product architectures. Supplier-managed models, federated learning, continuously learning systems and configurations composed of several interacting models may require additional associative entities and configuration rules.

The implementation prototype used during conceptual testing demonstrates that the relationships, deterministic findings and PCCP assembly can be represented in a working system. It does not establish that the proposed controls improve safety, reduce regulatory burden or produce PCCPs acceptable to a regulator or notified body. The Controlled Evolution model should be evaluated through real submission, audit and post-market cases, including changes that succeed, fail acceptance, deviate from protocol, cross a jurisdictional boundary or accumulate over several authorised releases.

The normative and legal environment is also evolving rapidly. EN 18286:2026 had been published for one day at the legal cut-off used for this paper, and its status under European harmonisation and citation should be rechecked before any conformity claim. The AI Act timeline reflects Regulation (EU) 2026/1744 and may be affected by subsequent legal, implementing or sectoral measures. The architecture cannot substitute for a product- and jurisdiction-specific regulatory strategy.

ISO/TS 24971-2:2026 is limited to machine-learning-enabled medical devices that do not employ large language models or generative AI. Further work is required to define equivalent risk-control depth for generative systems, adaptive multi-model systems and AI agents. Additional validation is also needed for the proposed relationship vocabulary, particularly cardinality, supplier-model relationships, monitoring data and configuration management across federated or continuously learning systems.

Future work should evaluate representative implementations and assurance cases against explicit comparators, including document-only and inventory-centred approaches. Useful measures include the time and accuracy of change impact analysis; the proportion of risks with traceable sources, controls and effectiveness evidence; completeness of model and data lineage; consistency between AI and medical-device risk conclusions; rate of stale or broken evidence links; time needed to answer regulator or auditor queries; and administrative effort per governed system. Evaluation should also examine whether the architec-

ture improves the quality and timeliness of decisions, rather than merely increasing record completeness.

XII. CONCLUSIONS

The governance challenge created by AI-enabled medical devices is not simply the addition of a new technology or body of compliance obligations. It is the need to maintain control over behaviour that depends on models, data, operating context and human interaction, while preserving the authority of established product, software and safety processes. A policy, inventory, model card or extended software risk table can contribute evidence, but none alone provides that continuity.

The proposed architecture addresses the challenge by controlling six related configuration entities, defining the meaning of their relationships and using operational records to make governance decisions against explicit baselines. Its coordinated risk model preserves wider AI effects while directing safety-relevant causal factors into the medical-device risk-management file. Its quality-system allocation avoids parallel evidence structures by leaving design, verification, validation, supplier, release, quality-event and post-market records under their competent authority.

Controlled Evolution extends decision continuity into prospective change. It defines the permitted change envelope, represents each proposal as a Planned Modification, connects the modification to its protocol, impact and evidence, and requires independent Change Authorisation before a revised baseline can be released. A PCCP can be assembled and frozen from those records for a particular market without becoming an isolated source of truth. When the authorised route is used, the QMS still controls execution, verification, validation, release, deployment and post-update monitoring.

The resulting contribution is decision continuity across the full product life cycle. From initial product-realisation planning to post-market change, retirement and historical record retrieval, the organisation can reconstruct what was authorised, why it was acceptable, which evidence and assumptions supported the decision, and which later event may require it to be reopened. That capability is the practical objective against which the architecture should be evaluated. It provides a demanding reference case for other regulated industries, provided that each sector supplies its own definitions of harm, acceptability, evidence and authority.

Status note. Legal, standards and bibliographic status was verified on 1 August 2026. The framework is a scholarly reference architecture and does not constitute legal advice or a claim of conformity or certification.

REFERENCES

- [1] International Organization for Standardization. ISO/IEC 42001:2023, Information technology: Artificial intelligence: Management system. Geneva: ISO; 2023. <https://www.iso.org/standard/42001>
- [2] International Organization for Standardization. ISO/IEC 23894:2023, Information technology: Artificial intelligence: Guidance on risk management. Geneva: ISO; 2023. <https://www.iso.org/standard/77304.html>
- [3] International Organization for Standardization. ISO/IEC 23053:2022, Framework for Artificial Intelligence Systems Using Machine Learning. Geneva: ISO; 2022. <https://www.iso.org/standard/74438.html>

- [4] European Committee for Standardization. EN 18286:2026, Artificial intelligence: Quality management system for EU AI Act regulatory purposes. Brussels: CEN; 2026. Publication announcement: <https://www.cencenelec.eu/news-events/news/2026/en-in-the-spotlight/2026-07-30-ai-quality-management/>
- [5] European Parliament and Council. Regulation (EU) 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence. Official Journal of the European Union. 2024. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- [6] International Organization for Standardization. ISO/TS 24971-2:2026, Medical devices: Guidance on the application of ISO 14971: Part 2: Machine learning in artificial intelligence. Geneva: ISO; 2026. <https://www.iso.org/standard/24971-2>
- [7] International Organization for Standardization. ISO 14971:2019, Medical devices: Application of risk management to medical devices. Geneva: ISO; 2019. <https://www.iso.org/standard/72704.html>
- [8] International Organization for Standardization. ISO/IEC 5338:2023, Information technology: Artificial intelligence: AI system life cycle processes. Geneva: ISO; 2023. <https://www.iso.org/standard/81118.html>
- [9] International Organization for Standardization. ISO/IEC 42005:2025, Artificial intelligence: AI system impact assessment. Geneva: ISO; 2025. <https://www.iso.org/standard/42005>
- [10] International Organization for Standardization. ISO/IEC 5259-4:2024, Artificial intelligence: Data quality for analytics and machine learning: Part 4: Data quality process framework. Geneva: ISO; 2024. <https://www.iso.org/standard/81093.html>
- [11] European Parliament and Council. Regulation (EU) 2026/1744 of 8 July 2026 amending Regulations (EU) 2024/1689, (EU) 2018/1139 and (EU) 2023/1230 as regards the simplification of the implementation of harmonised rules on artificial intelligence (Digital Omnibus on AI). Official Journal of the European Union. 2026. <https://eur-lex.europa.eu/eli/reg/2026/1744/oj/eng>
- [12] Medical Device Coordination Group and Artificial Intelligence Board. MDCG 2025-6 / AIB 2025-1: Interplay between the Medical Devices Regulation and In Vitro Diagnostic Medical Devices Regulation and the Artificial Intelligence Act. Brussels; 2025. https://health.ec.europa.eu/document/download/b78a17d7-e3cd-4943-851d-e02a2f22bbb4_en
- [13] International Organization for Standardization. ISO 13485:2016, Medical devices: Quality management systems: Requirements for regulatory purposes. Geneva: ISO; 2016. <https://www.iso.org/standard/59752.html>
- [14] International Electrotechnical Commission. IEC 62304:2006+A1:2015, Medical device software: Software life cycle processes. Geneva: IEC; 2015. <https://webstore.iec.ch/en/publication/22794>
- [15] European Parliament and Council. Regulation (EU) 2017/745 of 5 April 2017 on medical devices. Official Journal of the European Union. 2017. <https://eur-lex.europa.eu/eli/reg/2017/745/oj/eng>
- [16] European Parliament and Council. Regulation (EU) 2017/746 of 5 April 2017 on in vitro diagnostic medical devices. Official Journal of the European Union. 2017. <https://eur-lex.europa.eu/eli/reg/2017/746/oj/eng>
- [17] International Organization for Standardization. ISO/TR 24971:2020, Medical devices: Guidance on the application of ISO 14971. Geneva: ISO; 2020. <https://www.iso.org/standard/74437.html>
- [18] International Organization for Standardization. ISO/IEC 27001:2022, Information security, cybersecurity and privacy protection: Information security management systems: Requirements. Geneva: ISO; 2022. <https://www.iso.org/standard/27001>
- [19] International Organization for Standardization. ISO/IEC 27701:2025, Information security, cybersecurity and privacy protection: Privacy information management systems: Requirements and guidance. Geneva: ISO; 2025. <https://www.iso.org/standard/27701>
- [20] International Electrotechnical Commission. IEC 81001-5-1:2021, Health software and health IT systems safety, effectiveness and security: Part 5-1: Security: Activities in the product life cycle. Geneva: IEC; 2021. <https://webstore.iec.ch/en/publication/63293>
- [21] Overgaard SM, Graham MG, Brereton T, et al. Implementing quality management systems to close the AI translation gap and facilitate safe, ethical, and effective health AI solutions. *npj Digital Medicine*. 2023;6:218. <https://doi.org/10.1038/s41746-023-00968-8>
- [22] Bartels R, Dudink J, Haitjema S, Oberski D, van 't Veen A. A perspective on a quality management system for AI/ML-based clinical decision support in hospital care. *Frontiers in Digital Health*. 2022;4:942588. <https://doi.org/10.3389/fdgh.2022.942588>
- [23] Bedoya AD, Economou-Zavlanos NJ, Goldstein BA, et al. A framework for the oversight and local deployment of safe and high-quality prediction models. *Journal of the American Medical Informatics Association*. 2022;29(9):1631-1636. <https://doi.org/10.1093/jamia/ocac078>
- [24] Boag W, et al. The algorithm journey map: a tangible approach to implementing AI solutions in healthcare. *npj Digital Medicine*. 2024;7:87. <https://doi.org/10.1038/s41746-024-01061-4>
- [25] Khan SD, Hoodbhoy Z, Raja MHR, et al. Frameworks for procurement, integration, monitoring, and evaluation of artificial intelligence tools in clinical settings: a systematic review. *PLOS Digital Health*. 2024;3(5):e0000514. <https://doi.org/10.1371/journal.pdig.0000514>
- [26] Crossnohere NL, Elsaid M, Paskett J, Bose-Brill S, Bridges JFP. Guidelines for artificial intelligence in medicine: literature review and content analysis of frameworks. *Journal of Medical Internet Research*. 2022;24(8):e36823. <https://doi.org/10.2196/36823>
- [27] Lekadir K, Frangi AF, Porras AR, et al. FUTURE-AI: international consensus guideline for trustworthy and deployable artificial intelligence in healthcare. *BMJ*. 2025;388:e081554. <https://doi.org/10.1136/bmj-2024-081554>
- [28] Solaiman B, Mekki YM, Qadir J, Ghaly M, Abdelkareem M, Al-Ansari A. A "True Lifecycle Approach" towards governing healthcare AI with the GCC as a global governance model. *npj Digital Medicine*. 2025;8:337. <https://doi.org/10.1038/s41746-025-01614-1>
- [29] Wiens J, Saria S, Sendak M, et al. Do no harm: a roadmap for responsible machine learning for health care. *Nature Medicine*. 2019;25:1337-1340. <https://doi.org/10.1038/s41591-019-0548-6>
- [30] Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Medicine*. 2019;17:195. <https://doi.org/10.1186/s12916-019-1426-2>
- [31] Reddy S, Rogers W, Makinen V-P, et al. Evaluation framework to guide implementation of AI systems into healthcare settings. *BMJ Health & Care Informatics*. 2021;28:e100444. <https://doi.org/10.1136/bmjhci-2021-100444>
- [32] Vasey B, Nagendran M, Campbell B, et al. Reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *Nature Medicine*. 2022;28:924-933. <https://doi.org/10.1038/s41591-022-01772-9>
- [33] Liu X, Cruz Rivera S, Moher D, Calvert MJ, Denniston AK, SPIRIT-AI and CONSORT-AI Working Group. Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: the CONSORT-AI extension. *Nature Medicine*. 2020;26:1364-1374. <https://doi.org/10.1038/s41591-020-1034-x>
- [34] Cruz Rivera S, Liu X, Chan A-W, Denniston AK, Calvert MJ, SPIRIT-AI and CONSORT-AI Working Group. Guidelines for clinical trial protocols for interventions involving artificial intelligence: the SPIRIT-AI extension. *Nature Medicine*. 2020;26:1351-1363. <https://doi.org/10.1038/s41591-020-1037-7>
- [35] Brereton TA, Malik MM, Lifson MA, Greenwood JD, Peterson KJ, Overgaard SM. The role of artificial intelligence model documentation in translational science: scoping review. *Interactive Journal of Medical Research*. 2023;12:e45903. <https://doi.org/10.2196/45903>
- [36] Sendak MP, Gao M, Brajer N, Balu S. Presenting machine learning model information to clinical end users with model facts labels. *npj Digital Medicine*. 2020;3:41. <https://doi.org/10.1038/s41746-020-0253-3>
- [37] van Genderen ME, van de Sande D, Hooft L, et al. Charting a new course in healthcare: early-stage AI algorithm registration to enhance trust and transparency. *npj Digital Medicine*. 2024;7:119. <https://doi.org/10.1038/s41746-024-01104-w>
- [38] Gebru T, Morgenstern J, Vecchione B, et al. Datasheets for datasets. *Communications of the ACM*. 2021;64(12):86-92. <https://doi.org/10.1145/3458723>
- [39] Mitchell M, Wu S, Zaldivar A, et al. Model cards for model reporting. *Proceedings of the Conference on Fairness, Accountability, and Transparency*. 2019:220-229. <https://doi.org/10.1145/3287560.3287596>
- [40] Sculley D, Holt G, Golovin D, et al. Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*. 2015;28:2503-2511. <https://papers.nips.cc/paper/5656-hidden-technical-debt-in-machine-learning-systems>
- [41] Sambasivan N, Kapania S, Highfill H, Akrong D, Paritosh P, Aroyo LM. "Everyone wants to do the model work, not the data work": data cascades in high-stakes AI. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*. 2021:1-15. <https://doi.org/10.1145/3411764.3445518>
- [42] McCradden MD, Joshi S, Anderson JA, London AJ. A normative framework for artificial intelligence as a sociotechnical system in healthcare. *Patterns*. 2023;4(11):100864. <https://doi.org/10.1016/j.patter.2023.100864>

- [43] World Health Organization. Ethics and governance of artificial intelligence for health: WHO guidance. Geneva: WHO; 2021. <https://www.who.int/publications/i/item/9789240029200>
- [44] World Health Organization. Regulatory considerations on artificial intelligence for health. Geneva: WHO; 2023. <https://www.who.int/publications/i/item/9789240078871>
- [45] International Medical Device Regulators Forum. Good machine learning practice for medical device development: guiding principles. IMDRF/AIML WG/N88 FINAL:2025. 2025. <https://www.imdrf.org/documents/good-machine-learning-practice-medical-device-development-guiding-principles>
- [46] Tabassi E. Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. Gaithersburg, MD: National Institute of Standards and Technology; 2023. <https://doi.org/10.6028/NIST.AI.100-1>
- [47] U.S. Food and Drug Administration. Artificial intelligence-enabled device software functions: lifecycle management and marketing submission recommendations. Draft guidance. Silver Spring, MD: FDA; January 2025. <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/artificial-intelligence-enabled-device-software-functions-lifecycle-management-and-marketing>
- [48] U.S. Food and Drug Administration. Marketing submission recommendations for a predetermined change control plan for artificial intelligence-enabled device software functions. Guidance for industry and Food and Drug Administration staff. Silver Spring, MD: FDA; August 2025. <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/marketing-submission-recommendations-predetermined-change-control-plan-artificial-intelligence>
- [49] Team-NB and IG-NB. Artificial intelligence in medical devices: questionnaire. Version 1. Joint adoption date 14 November 2024. Brussels: European Association of Medical Devices Notified Bodies; 2024. <https://www.team-nb.org/wp-content/uploads/2024/11/Team-NB-PositionPaper-AI-in-MD-Questionnaire-V1-20241125.pdf>
- [50] Coalition for Health AI. Responsible AI Guide. Version 0.3, public-comment draft. 2024. <https://www.chai.org/workgroup/responsible-ai/responsible-ai-guide-raig-and-raig-executive-summary>
- [51] International Organization for Standardization. ISO/PAS 8800:2024, Road vehicles: Safety and artificial intelligence. Geneva: ISO; 2024. <https://www.iso.org/standard/83303.html>
- [52] Board of Governors of the Federal Reserve System, Federal Deposit Insurance Corporation and Office of the Comptroller of the Currency. Revised Guidance on Model Risk Management, SR Letter 26-2, with attachment Supervisory Guidance on Model Risk Management. Washington, DC; 2026. <https://www.federalreserve.gov/supervisionreg/srletters/SR2602.htm>
- [53] U.S. Food and Drug Administration. Quality Management System Regulation (QMSR). Effective 2 February 2026. Silver Spring, MD: FDA; 2026. <https://www.fda.gov/medical-devices/postmarket-requirements-devices/quality-management-system-regulation-qmsr>