

Care Intelligence

A Framework for AI Care Orchestrators, Nurse Stewardship, and the
Healthcare AI Operating Layer

A WHITE PAPER

Author: Robert Domondon — Founder, NAIIO Institute / Nurse Intelligence Network (Florence Media Network)

Version: 1.0 · June 2026

Status: Working draft for review. External-standard citations have been verified against current primary sources (June 2026); the figures herein are sourced to primary or near-primary authorities (among them AHA, WHO, Gallup, the U.S. Surgeon General, Pew Research Center, JAMA Network Open, and FDA device data), with residual return-on-investment and market-size claims kept deliberately measured (see *Honest Limitations*).

Abstract

Healthcare is not transformed by isolated artificial-intelligence features. It is transformed by a governed orchestration layer that coordinates many kinds of intelligence — across models, workflows, people, and settings — into coherent, accountable, and humane action. This paper names that orchestration capacity a distinct discipline, **Care Intelligence**, and argues that the profession already operating inside the orchestration loop of care — assessing, coordinating, escalating, monitoring, documenting, and sustaining continuity — is nursing. Building on established frameworks for care coordination (AHRQ), trustworthy and socio-technical AI (NIST), and the ethics and governance of AI for health (WHO), and on lifecycle controls for AI-enabled medical devices (FDA, IMDRF), the paper defines Care Intelligence operationally, distinguishes it from the tools that serve it, and sets out the mindset, the architecture, the administrative applications, the system configurations, and a staged roadmap by which institutions can build a trustworthy Care-Intelligence layer with clinicians — nurses foremost — as its stewards. The paper is deliberately confident about direction and design principles and measured about enterprise return-on-investment claims, which require institution-specific evidence.

Executive Summary

- **Care Intelligence** is defined as the disciplined capacity to align clinical need, human judgment, workflow, safety, ethics, and system constraints into context-aware action. It is a coined synthesis, grounded in established care-coordination and trustworthy-AI principles rather than presented as a settled regulatory term.
- The central thesis is that durable transformation comes from a *governed orchestration layer*, not from isolated AI features — a position consistent with NIST's socio-technical view of AI risk and benefit.
- The principal danger is not chiefly that AI produces a wrong answer; it is that AI produces an answer that is *right for the wrong definition of success* — optimizing a single metric while degrading the whole. Governance of objectives, not only of outputs, is therefore decisive.
- **Nurses are the most credible initial stewards** of the orchestration layer, because their daily practice already integrates assessment, coordination, escalation, monitoring, documentation, and continuity, and because the clinical reasoning cycle they are trained in (ADPIE) is structurally the feedback loop on which agentic AI is built.
- The orchestrator's competence is organized as **five literacies** — clinical, systems, governance, metacognitive, and communication — assembled as a competency ladder that external parties can inspect.
- The **AI operating system (AI OS)** is defined operationally as a unified, sovereign orchestration layer — a governance-and-execution fabric across models and workflows — not a single model, and not a regulatory category.

- Administration in the age of AI is reframed as an intelligence problem governed by the executive's own professional obligations, including the explicit duty to establish lifecycle ethics frameworks for AI/ML and automated decision systems.
 - System configuration should be evaluated across five dimensions — sovereignty and privacy, operational resilience, interoperability, governance and accountability, and scalability — with serious consideration of sovereign on-premise and hybrid-edge deployment for high-consequence workflows.
 - The roadmap is bounded and lifecycle-oriented: start with one auditable use case, govern before scaling, layer the intelligence, and treat deployment as a continuous learning loop — consistent with NIST, WHO, FDA, and IMDRF guidance.
-

A Note on Terms

Two central terms in this paper — **"Care Intelligence"** and **"AI OS"** — are deliberately coined. They are not settled terms in health-policy or regulatory literature, and the paper does not present them as standardized. Each is defined operationally and tethered throughout to existing governance language (AHRQ care coordination; the NIST AI Risk Management Framework; WHO guidance on the ethics and governance of AI for health; FDA and IMDRF lifecycle controls). The paper distinguishes a small family of related terms that nest rather than compete: *Caring Intelligence* (the governing ethos — care, dignity, accountability, human agency); *Care Intelligence* (the capacity defined here); *AI care orchestration* (the method by which the capacity is operationalized); and *AI OS* (the layer in which orchestration runs).

Introduction: The Problem and the Opportunity

Healthcare systems are being offered artificial intelligence at an accelerating pace and in a fragmented form — as discrete features embedded in individual products, each addressing a narrow task and each evaluated, if at all, in isolation. The result is a paradox familiar to clinical and operational leaders: a growing inventory of individually capable tools that, in aggregate, increases cognitive and administrative burden rather than reducing it, because the work of integrating them falls to already-stretched clinicians. That burden is not abstract: clinicians already carry on the order of two hours of administrative work for every hour of direct patient care (U.S. Surgeon General, 2022), and disconnected tools add to the load rather than lifting it. The promise of AI in care will not be realized one feature at a time. It will be realized, or lost, at the level of the system that coordinates those features into safe, accountable, and humane action.

This is also a governance moment, and adoption is no longer speculative: by 2024, 71 percent of non-federal acute-care hospitals reported using predictive AI within their electronic health records, up from 66 percent a year earlier (ONC/ASTP). The external scaffolding for responsible AI in health has matured alongside that adoption. NIST has framed AI as a socio-technical system to be managed across a lifecycle; WHO has articulated ethical principles for AI in health and, more recently, guidance specific to large multi-modal models; and the FDA and IMDRF have advanced

lifecycle controls, predetermined change-control planning, and good machine-learning practice for AI-enabled devices, with explicit emphasis on human-AI team performance and post-deployment monitoring. What these frameworks share is a conclusion that the unit of safety is the system in context — the interaction of model, workflow, people, and setting — and not the algorithm alone. They describe, in regulatory language, the same gap this paper addresses in operational language: the absence of a discipline, and a layer, that hold the whole together.

The opportunity is to treat that gap as a discipline to be built rather than a deficiency to be tolerated. The institutions that benefit most from AI in the coming years will be distinguished less by which models they purchase than by whether they build a governed orchestration layer — and by whether they place that layer in the hands of stewards who already understand the coordination of care. This paper is written for the leaders making that choice: chief nursing, medical, and information officers; chief AI and quality officers; policy and regulatory readers; and the clinical and technical teams who implement what they decide. It is organized to be read as thought leadership first, a governance framework second, and an implementation guide third, and it may be entered at any of those three levels.

1. Why This Perspective Exists

The argument of this paper is shaped by a vantage that spans the entire system it describes: the front line of patient care, the logic of diagnosis and therapy, the operations and compliance of hospital administration, the disciplines of corporate strategy and scalability, the synthesis of finance with population health, and the practice of healthcare communication and public trust. This is not offered as biography but as method. Most strategy for AI in healthcare is written from a single vantage — usually the technology or vendor vantage — and inherits that vantage's blind spots. A claim that intends to govern AI across the whole care system should be tested from the bedside, the boardroom, and the community at once.

A second consideration makes this vantage durable rather than exclusionary. In the age of AI augmentation, broad cross-domain awareness no longer requires formal training in every domain. AI can widen access to clinical, operational, legal, and financial expertise, providing decision support and access to knowledge bases at the point of need. This does not remove the requirement for human stewardship; it changes *who can meaningfully participate* in orchestrating care. The implication, developed throughout this paper, is a shift in posture: clinicians — and nurses in particular — should not regard AI as an external force to which they must adapt, but as another domain of practice in which their professional values of care, accountability, and judgment take the lead.

2. The Different Kinds of Intelligence

A useful premise for healthcare is that intelligence is not a single quantity on a single scale but a volume distributed across several largely independent dimensions. A working taxonomy for care includes: linguistic intelligence (explanation, documentation, translation); pattern intelligence (recognition and prediction); multimodal clinical intelligence (integrating laboratory data, images,

notes, and signals); operational intelligence (flow, scheduling, resource coordination); social intelligence (reading interpersonal and family dynamics); ethical intelligence (surfacing tradeoffs, duties, and accountability); and metacognitive intelligence (knowing what to question, what to escalate, and what *not* to optimize).

Two observations follow. First, contemporary health AI is strongest where signals are structured, repeatable, and measurable, which explains its concentration in diagnostics and imaging. The regulatory record makes the pattern concrete: of roughly 1,451 AI/ML-enabled medical devices the FDA had authorized through December 2025, about three-quarters are in radiology. This is consistent with the socio-technical framing of the NIST AI Risk Management Framework (AI RMF 1.0, 2023), which locates AI risk and benefit not in the model alone but in the interaction of technical components with people, workflows, and context. Second, much of care delivery still depends on contextual, social, and ethical judgment distributed across clinicians and teams. The dimensions in which current AI remains comparatively weak — embodied and physical intelligence, emotional nuance, social reading, and moral agency — are best understood not as proof that these capacities are unreachable but as a roadmap of intelligence still to be earned. It is precisely in those dimensions that clinicians, and nurses with holistic human judgment, are essential partners in shaping what comes next.

This distinction carries a practical corollary for evaluation and procurement. Strength on a dimension that can be verified or simulated — a labeled image, a passing test, a structured note — implies little about strength on a dimension that resists verification, such as reading a family's readiness to receive difficult news, or recognizing when a protocol should be set aside for this particular patient. Institutions assessing an AI system should therefore ask not only how capable it is, but on which dimensions its capability has actually been demonstrated, and should assign explicit human responsibility for the dimensions it does not cover. A system that performs well on a benchmark and is then deployed into a context the benchmark never measured is a common and avoidable source of harm — and exactly the failure mode the NIST framework's socio-technical emphasis is designed to surface.

3. Defining Care Intelligence

Care Intelligence is the disciplined capacity to sense, interpret, prioritize, coordinate, act, learn, and re-align care under real-world constraints. It is the capacity to hold patient need, professional judgment, operational reality, safety requirements, ethical guardrails, and institutional goals together in coherent, context-aware action.

The definition is anchored deliberately in existing language. It extends the AHRQ conception of care coordination — "the deliberate organization of patient care activities between two or more participants (including the patient) involved in a patient's care to facilitate the appropriate delivery of health care services" — from the coordination of activities to the coordination of *intelligences*, human and machine. It is bounded by the trustworthiness characteristics of the NIST AI RMF (valid and reliable; safe; secure and resilient; accountable and transparent; explainable and interpretable; privacy-enhanced; and fair with harmful bias managed) and by the six ethical principles of the WHO guidance on the ethics and governance of AI for health (2021): protecting autonomy; promoting human well-being, safety, and the public interest; ensuring transparency

and explainability; fostering responsibility and accountability; ensuring inclusiveness and equity; and promoting AI that is responsive and sustainable.

Care Intelligence carries an explicit orientation toward augmented clinical agency. Properly designed, AI does not replace human judgment; it equips clinicians to act more decisively, to anticipate need, and to preserve capacity for human connection by handling complexity in the background. The governing relationship is one in which capability granted to machines is matched by capability built in people — a point with direct safety implications, since meaningful human oversight depends on clinicians retaining the skill and context to challenge an AI output rather than merely ratify it. This orientation guards against a subtle but consequential failure mode: a system that gradually assumes the role of decision-maker while a nominal human approver loses the context, the skill, or the time to dissent. The framework treats the erosion of human capability — deskilling — as a defect to be designed against rather than an acceptable side effect, because meaningful oversight is the precondition on which every other safety claim ultimately depends.

4. The Art of Clinical Orchestration

The original contribution of nursing practice to this framework is a model of optimization that is holistic rather than single-objective. Across a single shift, a nurse does not maximize one variable; she balances several competing goods at once: cost efficiency and resource use; patient experience; clinical outcomes and safety; regulatory and policy compliance; operational performance such as length of stay and throughput; reimbursement considerations; continuity of care across shifts; and the workload, well-being, and resilience of the caregiver. A good shift is the concurrent balancing of these demands, not the maximization of any one of them.

This reframing exposes the central risk of AI in care. The most serious danger is not that a system produces an incorrect output but that it produces a technically correct output against a mis-specified objective — the reward-hacking problem. Optimizing length of stay can manufacture pressure toward unsafe discharge; optimizing satisfaction can reward the avoidance of difficult but necessary conversations; optimizing documentation time can erase clinically meaningful nuance; optimizing throughput can shift hidden burden onto nurses. Each is a measurable gain that conceals an unmeasured loss. The governance implication is that objectives, and not only outputs, must be examined: who selected the metric, what value is embedded in it, and what a system could do that technically succeeds while violating care.

Several established methods make this balancing rigorous rather than intuitive, and AI can place them in clinicians' hands: queuing theory for patient flow and bottlenecks; mathematical allocation methods for scarce resources; simulation for staffing and care scenarios; and multi-criteria decision analysis for weighing competing objectives. A related discipline is the allocation of intelligence itself — applying the appropriate level of human or machine capability to each task, reserving the highest-cost expertise for the highest-complexity work, and declining to spend premium intelligence on tasks that do not require it. Underpinning all of it is metacognition: the explicit discipline of asking which questions matter and what should, and should not, be optimized.

5. Healthcare AI Orchestration

Orchestration, in this framework, is the management of multiple models, agents, rules, workflows, handoffs, and escalations across the care lifecycle. It is not reducible to technical routing; it is the combination of clinical governance, operational design, and human-factors design. This is consistent with the direction of international device regulation: the IMDRF's work on machine-learning-enabled medical devices, and the Good Machine Learning Practice guiding principles it has advanced, emphasize multidisciplinary expertise across the total product lifecycle, evaluation on representative populations, a focus on human-AI *team* performance rather than model performance in isolation, clear information for users, and monitoring after deployment.

Orchestration also has a native cognitive structure already familiar to clinicians, and the parallel is exact enough to be operationally useful. The agentic loop on which modern AI systems are built — perceive, reason, act, observe, and adjust — is structurally the nursing process taught as ADPIE. Assessment corresponds to perception and context-gathering; Diagnosis to the formation of a usable model of the situation, which in machine terms is state estimation; Planning to action selection; Implementation to execution, frequently through the use of tools; and Evaluation to the observation of results that feeds a renewed assessment. The cadence of the loop appropriately tightens as stakes rise — a property that well-designed agentic systems strive to achieve and that expert nursing practice has long exhibited. Two implications follow. Because clinicians already reason in this cycle, orchestration is less a foreign competency to be acquired than an existing competency to be extended across human and machine agents. And because the decisive difference between the two loops is what closes them — a human exercising accountable judgment, as against an objective function pursued automatically — the central governance task is to ensure that the loop, in any consequential workflow, closes on a named person rather than on a metric.

Finally, orchestration requires graduated oversight proportioned to consequence — what may be described externally as an escalation ladder, and what this framework operationalizes through a four-tier capability scale (operational, strategic, transformational, and existential). The essential principle is that risk is a property of what is about to happen, evaluated against who is asking, with what data, toward what effect; that ambiguity is resolved toward the higher tier; and that the loop closes on a named, accountable human rather than on a tier or a tool. The orchestrator's role is best understood as that of a clinical systems conductor — not a replacement clinician.

6. Nurses as Natural Care Orchestrators and Stewards

The case for nurses as the initial stewards of the orchestration layer is a case from fit rather than sentiment. Nursing work maps directly onto the activities AHRQ identifies as constitutive of care coordination: establishing accountability, communicating, facilitating transitions, assessing needs and goals, creating proactive care plans, monitoring and following up, supporting self-management, linking patients to community resources, and aligning resources and the care team. The profession is, in other words, already organized around the coordination problem that AI orchestration must solve.

Three further considerations strengthen the case. First, the reasoning cycle of nursing (ADPIE) is the cycle agentic AI is converging on, which positions nurses to govern human-AI teaming in live care without first being retrained into an unfamiliar mental model. Second, WHO has emphasized that nurses and midwives are central to primary care and are often the first — and sometimes the only — health professional people see, which places them at exactly the point where an intelligence layer meets lived reality. Third, public trust — which AI in healthcare conspicuously lacks — is a resource nurses hold, and the gap can now be stated empirically rather than asserted. A majority of Americans report discomfort with AI in their own clinical care (about 60 percent in a 2023 Pew Research Center survey), and roughly two-thirds express low trust in their health system to use AI responsibly (about 66 percent in a 2025 JAMA Network Open study). Over the same period, nursing has been rated the most trusted profession in the United States for more than two decades — 76 percent of Americans rated nurses' honesty and ethical standards high or very high in Gallup's December 2024 poll, its twenty-third consecutive year at the top, with 75 percent the following year. The implication is structural rather than rhetorical: the public distrusts healthcare AI and trusts nurses above every other profession, which makes nurses the credibility bridge that healthcare AI otherwise lacks. That trust is itself the product of a practice tradition rooted in advocacy, accountability, and the relief of suffering.

Stewardship in this sense is not honorary; it implies concrete authority. It includes the standing to halt a workflow on safety or dignity grounds, to require that an AI output be explainable before it is acted upon, and to escalate without penalty. A stewardship role stripped of these powers is symbolic, and symbolic oversight is precisely what converts a governance framework into a liability — the appearance of a human in the loop without the substance of one. A useful practical test of whether an institution has genuinely made nurses stewards of its orchestration layer is simply whether a nurse can stop it.

This is explicitly not a territorial claim against physicians, pharmacists, informaticists, or administrators. The framework assumes a model of joint clinical authority in which nurses orchestrate the care process, physicians retain decision authority within their scope, and allied clinicians hold discipline-specific safety authority. The claim is narrower and more defensible: that stewardship of the *care process itself* — the coordination, monitoring, escalation, and continuity that orchestration automates and amplifies — belongs first to the profession that has always carried it.

7. The Mindset and Knowledge of the AI Care Orchestrator

The competence required of an AI care orchestrator can be organized as five literacies. **Clinical literacy** — care pathways, acuity, escalation, and safety. **Systems literacy** — workflow, interoperability, resource allocation, staffing logic, and transitions of care. **Governance literacy** — risk management, documentation, transparency, bias, privacy, incident response, and post-deployment monitoring, informed by the NIST AI RMF and by the professional codes that already bind clinicians and executives. **Metacognitive literacy** — knowing what to question, what to delegate, what to automate, and what must remain under human accountability. **Communication literacy** — translating among clinicians, executives, technical teams, compliance functions,

patients, and the public, which is the discipline of trust and is too often treated as peripheral to clinical and technical work.

These literacies are best developed as a competency ladder that progresses from an AI-aware practitioner to a credentialed care orchestrator, with attainment assessed as a trajectory that external parties — employers, regulators, payers — can inspect. External validation is constitutive of the credential: a competency that no outside party will examine is not yet real. Designed this way, the orchestrator role is an on-ramp that broadens participation rather than a narrow specialization, consistent with the democratization argument of Section 1.

8. AI OS as a Concept

The term *AI OS* is a strategic concept, not a regulatory category, and it must be defined operationally to be useful. An AI operating system, in this framework, is a unified, sovereign orchestration layer that manages access, context, policies, retrieval, model selection, auditability, escalation, and workflow integration across care settings. It is not a single model; it is a governance-and-execution fabric across models and workflows. The "operating system" metaphor earns its place only if it incorporates the trustworthiness, clinical role clarity, and lifecycle monitoring that NIST, FDA, and IMDRF require — otherwise it is branding.

The strategic value of consolidation is governance as much as convenience. Clinicians today are required to master a proliferating set of applications, each demanding that the human adapt to it and serving as the integration layer between systems that do not interoperate. A unified orchestration layer that adapts to the clinician, rather than the reverse, consolidates not only workflow but accountability — provided it is governed. Consolidation without governance simply relocates the problem, concentrating risk into a single surface; a unified layer that optimizes a mis-specified objective can propagate that error uniformly and at scale. The governing question is therefore not whether to unify, but what conscience the unified layer carries.

Defining the AI OS operationally yields concrete requirements rather than aspiration, which is what makes it useful for design and procurement. A serious orchestration layer must be able to select the model or agent appropriate to a given task and route accordingly; to enforce access controls and data boundaries; to retrieve and ground its outputs in authoritative sources; to maintain an auditable record of what was proposed, by which component, and who acted on it; to surface its own uncertainty and its applicable risk tier to the human overseeing it; and to escalate to a named person as consequence or ambiguity rises. These are properties an institution can specify, test, and require of a vendor — and they are the properties that distinguish a governable layer from a merely convenient one. An institution that writes them into procurement is already practicing Care Intelligence before a single model is chosen.

That conscience is best understood as architecture rather than afterthought. Governance, accountability, bias mitigation, privacy protection, and transparency should be structural properties of the system, embedded from inception, rather than retrofitted through later review. The practical test is whether a safeguard is load-bearing: if it can be removed without the system ceasing to function, it is decorative rather than structural. Configured this way, the layer is designed to operate as a disciplined clinician does — proactively, accountably, and with care and

ethics balanced in real time. A first, deliberately bounded instance of this concept is a governed, no-protected-health-information environment in which nurses can learn and steward AI without surrendering judgment — explicitly not a clinical decision system and not a replacement for licensed judgment, but an on-ramp to an accountable AI practice.

9. Healthcare Administration in the Age of AI

Administration in the age of AI is most usefully reframed as an intelligence problem rather than only a compliance problem, with each administrative surface becoming a site of orchestration — and each tethered to the holistic balance described in Section 4 so that efficiency does not become a single-metric optimization.

The largest opportunity is the largest cost. Labor is the dominant component of hospital operating expense — roughly 56 percent in 2024, and nearly 60 percent once contract and purchased clinical labor are counted (American Hospital Association) — and AI can forecast demand by volume, acuity, and seasonality; balance skill mix; respect staff preferences and compliance constraints; and reduce both costly over-staffing and unsafe under-staffing. Even modest efficiencies against an expense of this magnitude are material; a single health system has reported, in trade press, savings of roughly thirty million dollars and a fall in agency labor from about a quarter of its workforce to under a tenth after adopting AI-assisted workforce management. That result is single-institution and trade-press-reported rather than generalizable evidence, and it cuts both ways: labor cost optimized in isolation is precisely the single-metric trap described in Section 4, which is why workforce-management AI must be governed against staff well-being and safety, and not against cost alone. Supply and inventory management benefit from demand forecasting, automated ordering, waste reduction, and system-wide materials mobilization. Medication safety is a problem of real magnitude — the World Health Organization estimated in 2024 that medication errors harm more than 1.5 million Americans each year and are associated with several thousand deaths, against a global cost on the order of forty billion dollars — and medication management gains a continuous safety layer that checks for interactions, allergies, duplication, and dose appropriateness against patient-specific factors, working alongside the pharmacist. Diagnostics gain a second read that flags anomalies and trends for earlier intervention while the clinician remains the reader of record; the randomized MASAI trial of AI-supported mammography in roughly 80,000 women, for example, found about 20 percent more cancers detected alongside a 44 percent reduction in radiologists' reading workload and no increase in false positives — the kind of representative-population, human-AI-team evidence this framework calls for. Revenue-cycle functions gain more accurate coding from documentation, detection of fraud, duplication, and misclassification, and streamlined claims and pre-authorization — bridging a fragmented payment environment. Legal and risk functions gain clearer and more personalized informed consent and disclosure, real-time flagging of protocol deviations and documentation gaps, and tracking of regulatory change, with named humans retaining authority to finalize.

Across all of these surfaces a common governance pattern applies, and stating it explicitly protects the institution from the principal failure mode. Every administrative model should have a named accountable owner; should be measured against an audited pre-deployment baseline

rather than a projected one; should carry at least one explicit counterweight metric drawn from patient safety, equity, or workforce sustainability; and should be monitored on a defined cadence for performance drift and for unintended effects on the populations and staff it touches. Administrative AI is where the temptation to optimize a single financial metric is strongest, and where the consequences of doing so are most easily concealed, because the resulting harm typically lands on parties the optimized metric does not observe — a future shift, an excluded population, an overburdened unit.

This administrative agenda is not external to professional ethics; it is increasingly an explicit expectation of professional leadership. The American College of Healthcare Executives, in its 2025 policy statement *Healthcare Leaders and the Ethical Use of Artificial Intelligence*, calls on leaders to ensure that technology used in patient-care decision-making is transparent, equitable, and subject to ongoing monitoring, and to establish ethics frameworks for the development, procurement, implementation, and ongoing monitoring of artificial intelligence, machine learning, and automated decision-making systems. (This guidance resides in an ACHE policy statement, which is advisory, rather than in the binding ACHE Code of Ethics.) Administration in the age of AI is, in substantial part, the construction of that lifecycle governance framework. The standing guardrail is that administrative AI must not optimize cost or revenue alone; every administrative model should carry an explicit counterweight for patient safety, equity, and workforce sustainability.

10. Healthcare AI Systems Configurations

There are three viable deployment patterns — cloud-assisted, sovereign on-premise, and hybrid-edge — and the framework's position is that high-consequence workflows warrant serious evaluation of sovereign and hybrid options rather than default reliance on the cloud. The rationale is threefold: protection of patient data under direct institutional control; operational resilience against disruption of connectivity, of the supply of external intelligence, and of energy; and cybersecurity posture. Edge computing complements sovereignty by placing intelligence close to the point of care, which lowers latency, sustains operation through partial outages, and delivers insight where seconds matter, as in critical care and surgery. Sovereignty does not eliminate the need for interoperability; it changes where control, data residence, and resilience are located, and interoperability with regional exchanges and evolving models remains essential.

Configuration is best evaluated as the simultaneous balancing of five dimensions: sovereignty and data privacy; operational resilience; interoperability; governance and accountability; and scalability of compute and edge capacity. Each implies a concrete question. Sovereignty and privacy ask where identified data resides and who can compel access to it. Operational resilience asks whether the system continues to function during a loss of connectivity, of the external supply of intelligence, or of power — and for how long. Interoperability asks whether the system can exchange information with regional exchanges, existing clinical systems, and successor models without becoming a closed enclave. Governance and accountability ask who is responsible for updating models, validating performance, and answering for outcomes. Scalability asks whether compute, storage, and edge capacity can grow with demand without re-architecting from the ground up. A configuration that scores well on one dimension and poorly on another — sovereign

but unable to interoperate, or scalable but without clear accountability — is brittle in proportion to the dimension it neglects. To these the framework adds an energy dimension, both because energy is a resilience dependency for any compute-intensive system and because the ecological footprint of AI is an ethical variable — a position consistent with WHO's principle that AI for health should be responsive and sustainable. Configuration, in short, is itself a holistic optimization; optimizing one dimension while neglecting another yields a brittle system.

11. A Roadmap to Get There

The recommended path is bounded, governed, layered, and continuous. **Start bounded.** Choose a single use case in which orchestration value is visible, auditable, and clinically relevant; structured assistance and decision support are lower-friction entry points than fully agentic care, and reversible actions are preferable to irreversible ones because they make learning inexpensive. **Govern before scaling.** Define intended use, the human role, dataset boundaries, subgroup performance, transparency artifacts, update procedures, escalation paths, and monitoring metrics before expanding scope — the discipline encoded in the FDA's lifecycle-management and predetermined-change-control guidance for AI-enabled device software and in IMDRF's total-product-lifecycle approach. **Layer the intelligence.** Progress deliberately from documentation support to coordination support to pathway monitoring to resource planning and only then to selective, supervised agentic workflows, widening scope only behind demonstrated safeguards. **Treat deployment as a learning loop, not a launch event.** NIST, WHO, FDA, and IMDRF converge on lifecycle governance over one-time validation, with continuous monitoring for performance, drift, and emergent behavior. Throughout, success should be measured as a set — the competing goods of Section 4 carried together, alongside the trajectory of human capability and the structural question of who is served and who is excluded — so that the program cannot be judged a success by a single metric while failing the whole.

Two commitments make this roadmap auditable rather than aspirational. First, each stage should produce a small set of governance artifacts — a statement of intended use, a description of the human's role and the points of override, a record of subgroup performance, and a monitoring plan — that travel with the system and are treated as clinical infrastructure rather than paperwork. Second, each stage should define in advance the conditions that would cause the institution to pause or stop, so that halting is a planned governance event rather than a crisis. An organization that can state, for any deployed system, who is accountable, what it is for, where the human closes the loop, and what would make it stop, has built the substance of a Care-Intelligence layer regardless of the sophistication of the underlying models.

The Ten Doctrines

The framework's operating commitments are summarized as ten doctrines, intended to shape institutional culture before they shape technology. Each is grounded in the external standards cited above and in the nursing and professional-ethics tradition.

1. **Care before automation.** Automation is legitimate only when it strengthens safer, more effective, and more humane care.
2. **Ethics by architecture.** Ethical requirements belong in system design, role definitions, interfaces, and monitoring plans — not only in retrospective review.
3. **The human-AI team is the real unit of performance.** AI is judged in workflow context, not in benchmark isolation.
4. **Representative care is trustworthy care.** Systems are evaluated across intended populations, contexts, and subgroups.
5. **Nurses are founding stewards of care orchestration** — not because they own every decision, but because they already operate inside the coordination loop.
6. **Every AI system needs an escalation ladder.** The more consequential the task, the more visible the human override and accountability path.
7. **Transparency is clinical infrastructure.** Intended use, limits, failure modes, subgroup performance, and update history are part of safe care.
8. **Sovereignty where stakes demand it.** Deployment architecture for high-consequence workflows reflects privacy, resilience, latency, and continuity needs rather than default convenience.
9. **Apply the right intelligence to the task.** Matching the level of human or machine capability to the complexity of the task is an ethical and economic duty.
10. **Optimize the whole, not the metric.** Care Intelligence harmonizes competing objectives and does not collapse care into a single number.

Honest Limitations

Intellectual candor remains part of the argument. The terms *Care Intelligence* and *AI OS* are coined and should be defined carefully, used consistently, and tethered to existing governance language wherever they appear. The figures cited in this paper have been sourced to primary or near-primary authorities — among them the American Hospital Association, the World Health Organization, Gallup, the U.S. Surgeon General, the Pew Research Center, JAMA Network Open, and the FDA's device authorizations — and dated accordingly; the count of FDA-authorized AI/ML-enabled devices in particular rises monthly and is cited as of December 2025 against the FDA's live list. The strongest evidence base supports the governance direction advanced here — socio-technical risk management, lifecycle controls, transparency requirements, representative evaluation, and human-AI team performance — and specific high-value applications such as diagnostic assistance, where randomized evidence now exists. Three classes of claim are deliberately left measured. First, there is no randomized-trial evidence for the return on investment of AI-assisted staffing and scheduling; the system-level results referenced here are single-institution and trade-press-reported, and should be read as illustrative rather than generalizable. Second, market-size projections for healthcare AI are published only by private research firms and are omitted here in favor of that restraint. Third, the four-tier risk posture is a design stance

intended to govern the emergent risk of agent ecosystems proactively rather than an evidence base; it should be presented as such and validated through bounded pilots.

Conclusion

Healthcare does not principally need smarter tools. It needs a new layer of stewarded intelligence that can hold together care, workflow, governance, and human values — a governed orchestration layer that coordinates many kinds of intelligence into coherent, accountable, and humane action. The discipline of building and governing that layer is what this paper calls Care Intelligence, and the profession best positioned to steward it is the one that has always carried the coordination, monitoring, escalation, and continuity of care. The work ahead is neither to await a more capable model nor to defend against it, but to build the orchestration layer deliberately — bounded, transparent, sovereign where the stakes demand it, and answerable always to a named human being — so that as intelligence becomes abundant, the judgment that governs it becomes more capable, not less.

References and Source Backbone

External-standard references are current as of June 2026 and should be re-verified and dated at publication. Internal-doctrine references are available from the author.

External standards

- Agency for Healthcare Research and Quality (AHRQ). *Care Coordination Measures Atlas* (Update, June 2014), Chapter 2 — source of the quoted care-coordination definition (synthesized from a review of more than 40 definitions); and the AHRQ care-coordination measurement framework. ahrq.gov.
- National Institute of Standards and Technology (NIST). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, January 2023 — socio-technical risk framing; seven trustworthiness characteristics; the Govern–Map–Measure–Manage functions. nist.gov.
- World Health Organization (WHO). *Ethics and Governance of Artificial Intelligence for Health* (June 2021) — six core ethical principles; and *Ethics and Governance of AI for Health: Guidance on Large Multi-Modal Models* (January 2024). who.int.
- U.S. Food and Drug Administration (FDA). *Artificial Intelligence-Enabled Device Software Functions: Lifecycle Management and Marketing Submission Recommendations* (draft guidance, January 7, 2025; docket FDA-2024-D-4488); and *Marketing Submission Recommendations for a Predetermined Change Control Plan for AI-Enabled Device Software Functions* (final guidance, December 4, 2024; docket FDA-2022-D-2628). fda.gov.
- International Medical Device Regulators Forum (IMDRF). *Machine Learning-enabled Medical Devices: Key Terms and Definitions* (IMDRF/AIMD WG/N67, May 2022); and

Good Machine Learning Practice for Medical Device Development: Guiding Principles (IMDRF/AIML WG/N88 FINAL:2025, January 2025) — total-product-lifecycle approach, representative evaluation, human-AI team performance, post-deployment monitoring. [imdrf.org](https://www.imdrf.org).

- American College of Healthcare Executives (ACHE). Policy Statement, *Healthcare Leaders and the Ethical Use of Artificial Intelligence* (approved December 2025) — transparent, equitable, and monitored decision technology and the establishment of AI/ML lifecycle ethics frameworks. This obligation resides in an ACHE policy statement (advisory), distinct from the binding ACHE *Code of Ethics*. [ache.org](https://www.ache.org).

Data sources

- American Hospital Association (AHA). *Costs of Caring* — hospital labor as a share of operating expense (~56% in 2024; nearly 60% including contract and purchased clinical labor). [aha.org](https://www.aha.org).
- U.S. Surgeon General. *Advisory on Health Worker Burnout* (2022) — administrative-to-direct-care time ratio (~2:1 in primary care); workforce burnout and turnover costs.
- Gallup. *Honesty and Ethics in Professions* (December 2024; December 2025) — nurses rated the most trusted profession, 76% (23rd consecutive year) and 75% (24th). news.gallup.com.
- Pew Research Center (February 2023) — 60% of Americans would be uncomfortable with a provider relying on AI in their own care.
- JAMA Network Open (February 2025) — approximately 65.8% of U.S. adults report low trust in their health system to use AI responsibly.
- World Health Organization (2024) — global medication-safety burden: more than 1.5 million Americans harmed annually, several thousand deaths, and a global cost on the order of US\$42 billion.
- U.S. Food and Drug Administration — list of AI/ML-enabled medical devices (~1,451 authorized through December 2025; ~76% radiology), cross-checked against JAMA Network Open (November 2025). [fda.gov](https://www.fda.gov).
- The MASAI randomized controlled trial, *Lancet Oncology* (August 2023) — AI-supported mammography in approximately 80,000 women.
- Office of the National Coordinator / ASTP (2026) — 71% of non-federal acute-care hospitals used predictive AI within their EHRs in 2024. [healthit.gov](https://www.healthit.gov).

Internal doctrine (NAIO Institute / Nurse Intelligence Network)

- *Care Intelligence* — Master (CI), and the *NAIO AI Values Statement* (NAIO-VS) — the governing values and conduct layer.
- The *EDENA* four-tier capability scale (operational / strategic / transformational / existential).
- *The Three-Code Doctrine* — the AMA, ACM, and ACHE codes as a shared chain of accountability for responsible AI in care.

- Nursing foundations: *ANA Code of Ethics for Nurses* (2025); Jean Watson, *Caring Science*; Florence Nightingale, *Notes on Nursing* and environmental theory.

NAIO Institute · Nurse Intelligence Network (Florence Media Network) — "Where Nightingale meets neural net."