

WISL SERIES · VOLUME No. 91

An Introductory Beginning

Immutable Hardware-Enforced Non-Agentive AI Governance Engine

NAIGE

Hardwiring Medical AI Ethics in Singapore

*Translating Singapore's AI-MD governance
expectations — MOH/HSA AIHGle 2.0 and HSA GL-04
— into hardware-enforced constraints, via the NAIGE
NAI 4.0 Kinetic Capstone.*

ISBN: 978-981-17970-2-6

Edwin Koh Wui Kiat (Tiger)

Sovereign Author & Architect

Non-Agentive AI Governance Singapore

ACRA T260229801

July 2026

An introduction, and an invitation

This volume is the entry point to the NAIGE architecture. It sets out, in plain terms, a single idea: that safety for clinical AI should be a physical property of the hardware, not a promise made in software. Where a software guardrail can be edited, bypassed, or allowed to drift, a hardware constraint either holds or it does not — and it can be inspected, measured, and proven.

The chapters that follow map each mechanism of the engine to the governance expectations Singapore has already set out for AI medical devices, so that the framework can be read against a real regulatory backdrop rather than an abstract one.

A note on maturity. NAIGE is presented here at the concept-and-design stage. The mechanisms and figures described are design intent and derived requirements — not yet independently measured, built, or certified. This volume is written to invite exactly the step that would test them: a kinetic and clinical validation. Read it as an architecture seeking its proof, not a product claiming one.

SECTION 01

The MOH/HSA AIHGle 2.0 Direction

With the March 2026 release of the MOH/HSA Artificial Intelligence in Healthcare Guidelines (AIHGle 2.0) — complemented by HSA’s GL-04 Regulatory Guidelines for Software Medical Devices — Singapore has set clear expectations for AI Medical Devices (AI-MDs): accountability across the AI lifecycle, transparency, robust risk management, and patient safety.

As AI moves closer to acute, high-stakes clinical robotics, the limits of probabilistic, software-only safeguards become harder to accept. A guardrail that can be silently edited, bypassed, or allowed to drift is difficult to certify with confidence. The NAIGE NAI 4.0 Kinetic Capstone responds by shifting AI safety from permeable software policies to hardware-enforced silicon constraints — making the safe state a physical default rather than a software promise.

SECTION 02

Patient Safety: The P-LIFE 1.00™ Doctrine

A central duty in the AIHGle 2.0 guidelines is upholding non-maleficence and patient safety. NAIGE hardwires this duty through the P-LIFE 1.00™ (Inaction-as-Safe-Default) doctrine.

Rather than relying on a software “stop” command, the machine’s baseline default state is absolute mechanical

stillness. An AI cannot autonomously trigger physical actuation; motion is denied unless verified human authority actively releases the hardware constraints.

SECTION 03

Data Protection: Privacy-by-Physics via 3ZEROS™

AIHGle 2.0 calls for secure-by-design data protection. Standard approaches use software blurring to obscure identity, which is reversible. NAIGE instead enforces Privacy-by-Physics through its P-Layer, using the 3ZEROS™ protocol — Zero Camera, Zero Audio, Zero Cloud.

In place of optical lenses, the system uses a 905nm LiDAR array that voxelises raw spatial point-clouds into an anonymised 96-cell grid at the node, so that biometric identity cannot be reconstructed downstream before the AI processes the data.

SECTION 04

Human Autonomy: The N-Layer & the Sacred Pause™

To safeguard human autonomy against opaque, “black-box” decisions, NAIGE isolates the probabilistic AI into a non-deterministic Advisory Stratum with zero actuation rights.

When the AI proposes an action, the deterministic N-Layer arrests it using a 1 kHz master clock and enforces the Sacred Pause™: a mandatory, hardware-timed delay of 25-1000 ms that suspends the machine, forcing an

unavoidable window for human deliberation before any kinetic command can propagate.

SECTION 05

Sovereign-in-the-Loop: The Tripartite .1x Key™

To guarantee that AI only advises while humans ultimately decide, the N-Layer acts as a rigid “AND” logic gate requiring multi-factor physical human authorisation. No kinetic action executes without the Tripartite .1x Key™, which demands three simultaneous, hard-to-forge inputs:

- Biological identity — an iris scan (token expiring in seconds).
- Cryptographic intent — a hardware USB token.
- Physical presence — active depression of a kinetic foot pedal (dry-contact switch).

Because one factor is physical co-location, a remote or synthetic input — a spoofed credential or a deepfake — cannot, by itself, cause the machine to move.

SECTION 06

Traceability: WD117 Forensic Finality

AIHGle 2.0 asks that AI-MD decisions be transparent, explainable, and reproducible. Because ordinary software logs can be edited or deleted by an attacker, NAIGE secures accountability in its D-Layer through the WD117 Enhanced WORM (Write-Once-Read-Many) ledger.

Every event is hashed (SHA-3-512) and committed to a write-once record, so that the audit trail is designed to be unalterable and legally defensible after the fact.

Design note: at scale, a continuous event stream is best carried as a hash chain on write-once media, periodically anchored to one-time-programmable silicon rather than burning every event to fuses — a refinement for the hardware requirements specification.

SECTION 07

Fail-Safe Governance: The U7 Sovereign Brake™

AIHGle 2.0 expects AI-MDs to perform robustly and consistently. To this end the NAIGE Capstone houses the U7 Sovereign Brake™, a one-shot sequencer that monitors continuously for anomalies, agentic drift, or cyber-breach.

If a boundary is breached or authorisation fails, the U7 is designed to sever the low-dropout (LDO) power rails to the actuators in under 1.05 milliseconds, returning the machine to its safe default of mechanical stillness.

SECTION 08

Preventing Agentic Drift: The A7 Bridge Hub

Agentic drift occurs when an AI silently expands its own decision-making scope. To prevent it, NAIGE forces all probabilistic advisories through a single hardware chokepoint: the A7 Bridge Hub.

This ingress gateway strips the AI of actuation rights and screens incoming proposals for contradiction, hallucination, and uncertainty, translating semantic data into strict, deterministic hardware features before anything can reach the execution pipeline.

SECTION 09

Tiered HSA Alignment: Class B vs. Class C/D

NAIGE maps to HSA's SaMD risk classifications through two modular configurations:

- A12 [P+N] — Observational Tier (HSA Class B SaMD). Uses the Privacy and Navigation layers for residential eldercare monitoring, such as the Toa Payoh Hearth setting. Advisory only, with zero physical actuation rights.
- A12 [P+N+D] — Acute Kinetic Tier (HSA Class C/D). Adds the D-Layer (Sovereign Brake and WD117 ledger) for high-payload clinical robotic platforms performing acute physical interventions near vulnerable patients.

Classification follows the device's intended use and its own risk, not the room it sits in: a non-actuating monitor remains Class B even inside a hospital ward — though ward-IT integration adds its own network and cybersecurity obligations.

SECTION 10

The 10-Point Sovereignty Audit

To support HSA dossier compliance, NAIGE is designed to undergo a pre-deployment audit intended to demonstrate safety as an instrument-measured physical fact rather than a software promise. The audit is organised in four phases:

- Phase 1 — Compute & timing. Power-draw monitors verify the Orange Code CPU cap holds AI resources at a 1.1x baseline; oscilloscopes confirm the Sacred Pause™ timing.
- Phase 2 — Privacy. Physical confirmation of the total absence of optical lenses and audio drivers (3ZEROS™).
- Phase 3 — Authorisation. Independent cryptographic validation of iris tokens and continuity testing of the kinetic foot pedal.
- Phase 4 — Isolation. Verification of a scheduled data purge and total namespace isolation, proving the AI has no access to hardware control pins.

COLOPHON

On this work

This is Volume No. 91 of the WISL series — an introductory beginning intended to open the architecture to readers, reviewers, and prospective validation partners. It is a conceptual and design framework; each mechanism described here awaits its measured proof.

The mission it serves is a plain one: that intelligence near a vulnerable patient should observe and advise, while a human decides and the hardware enforces — non-autonomous by construction.

Author & Architect: Edwin Koh Wui Kiat (Tiger)

Publisher: Non-Agentive AI Governance Singapore · ACRA T260229801

Guiding principles: humility · silence · dignity · benevolence

Aligned to: MOH/HSA AIHGle 2.0 · HSA GL-04 · ASEAN CSdT