



AIDEN: Sovereign AI Control Tower

The Problem & The Architecture

The Problem

Enterprise teams do not fail because they lack software. They fail because critical work is spread across too many disconnected systems.

Plans live in spreadsheets. Official records live in systems of record. Decisions happen in meetings, email, and chat. Evidence is assembled after the fact. By the time leadership asks what happened, teams are often reconciling yesterday's reality by hand.

That gap creates four recurring problems:

- **Process drift:** the operating plan and the official record slowly move apart.
- **Manual reconciliation:** skilled people spend time checking rows, tags, dates, values, and approvals instead of making decisions.
- **Weak evidence trails:** it's hard to reconstruct the reason behind a decision later.
- **AI trust risk:** a model may sound confident even when the underlying data has not been checked.

Aiden is designed to solve this at the process layer and the compute layer.

At the top, Aiden makes enterprise AI more controlled by wrapping work in deterministic tools, schemas, approvals, and evidence.

At the bottom, Aiden runs on a private, CUDA-independent inference path, so the intelligence layer can remain modular, protected, and independently governable.

Aiden Architecture In One View

- **Upper Layer:** governed, human-in-the-loop process controls that connect operating plans to official records.
- **Lower Layer:** private, CUDA-independent local inference on Intel acceleration.
- **Adoption Path:** low-burden 1-3-month BPR-in-a-Box before heavy IT integration.



The Aiden Control Tower

Aiden is a Sovereign AI Control Tower for enterprise operations.

It does not replace the client's systems of record. It sits above them as a governed operating layer.

The basic model is simple:

- Operating systems hold intent.
- Systems of record hold truth.
- Aiden governs the loop between them.

In our current **Metrc** proof case (i.e., heavily regulated industry **Track & Trace**):

- **Smartsheet** (no-code work management platform) is the Digital-Twin operating cockpit where operations stage work.
- **Metrc** is the compliance record and chain-of-custody ledger.
- **Aiden** pulls, compares, routes for review, governs execution, confirms results, and preserves evidence.

The same pattern applies beyond regulated industry verticals/supply chains:

- Real estate underwriting and asset management.
- Insurance review and pricing workflows.
- FMCG Supply Chain and Inventory operations.
- Regulated finance, legal, healthcare, and compliance processes.

Why This Is Different from a Chatbot

Aiden separates probabilistic reasoning from deterministic execution.

The language model can explain, summarize, compare, and advise. But it does not blindly write to enterprise systems.



Every controlled action moves through a governed path:

Pull current records

-> Compare against operating intent

-> Preview / dry run

-> Preflight checks

-> Human approval

-> Execution

-> Confirmation

-> Evidence write-back

This creates a safer division of labor:

- **Deterministic Code owns the facts:** schemas, calculations, row parsing, validation, and controlled system calls.
- **Aiden owns the synthesis:** explaining what the evidence means and where attention is needed.
- **Humans own the release decision:** no production write proceeds without explicit approval.

This is the core trust mechanism.

Private, CUDA-Independent Inference

Aiden's application layer talks to a private local model endpoint through a familiar model-serving interface. Underneath that endpoint, the inference path avoids the NVIDIA CUDA acceleration stack and runs through an Intel XPU / Level Zero / oneAPI-class execution path on bare metal.

In simpler terms:

- The business app gets a “normal” way to call a model.
- The compute layer runs on our “private” Intel acceleration path.

CUDA is not being emulated. It is being avoided. The system uses a different acceleration path.

This matters because it gives the architecture a stronger independence story:



- Private local inference instead of default dependency on external model APIs.
- CUDA-independent acceleration instead of dependence on only one GPU software ecosystem.
- Modular application workflows above the inference layer.
- Clearer governance over where data, memory, and execution live.

Predictable inference economics: local inference shifts the cost model away from variable external API/token usage and toward fixed infrastructure capacity that can be planned, tuned, and amortized.

The critical point is **independence**: Aiden shows that enterprise AI workflows can be built on a governed, private, CUDA-independent acceleration path.

What Aiden Delivers & How to Start

What Aiden Delivers

Aiden gives leadership a practical control layer for high-value, high-risk workflows.

The value is not simply "AI answers." The value is **governed operational movement**.

Aiden helps teams:

- See current records from the systems that matter.
- Compare plans against official state.
- Identify mismatches before they become operational or compliance problems.
- Route sensitive decisions through human review.
- Execute approved actions through controlled pathways.
- Preserve evidence for review, audit, and future process design.

In the Metrc proof case, Aiden can trace a package back through source plant, plant batch, harvested plants, harvest, and final package. It can also pull Metrc and Smartsheet evidence into Diagnostics, run an evidence-bound Analyst review, and explain the operating meaning in plain language.

That same pattern can be applied to CRE-style workflows:



- Underwriting intake and review.
- Rent roll and operating statement reconciliation.
- Asset management task governance.
- Insurance review/pricing workflows.
- Evidence packs for investment committee, leadership review, or client reporting.

Why It Matters Now

Most enterprise AI programs face two bottlenecks.

First, the **business-process bottleneck**: AI is impressive in conversation but weak when it must interact safely with real workflows, real approvals, real documents, and real systems of record.

Second, the **infrastructure-dependence bottleneck**: teams want AI capability, but they do not always want sensitive processes coupled to public APIs, hyperscaler defaults, or a single acceleration ecosystem.

Recent U.S. policy direction is moving in the same direction: AI, future computing, information management, cybersecurity, semiconductors, and microelectronics are increasingly treated as strategic infrastructure, not just software categories. For enterprise leaders, the message is clear: AI adoption needs secure, governable, resilient infrastructure, with human accountability and trusted execution paths built in from the start.

Aiden addresses both:

- Deterministic control above.
- Private CUDA-independent inference below.

That is the unified stack.

The top of the stack protects the business process. The bottom of the stack protects infrastructure choice and data posture.



BPR-in-a-Box: The Practical Adoption Path

Aiden can be introduced without asking the client organization to stop operations, replace core systems, or involve IT before the business case is clear.

The starting offer is a focused **BPR-in-a-Box** engagement.

This is a 1–3-month Business Process Review that produces:

- An As-Is business process map
- A Future State workflow design
- A **working prototype** in a **secure operating box**.
- Clear evidence of where time, risk, and rework can be reduced.
- A business case for change before deeper integration begins.

For client organizations (CRE example), the practical approach would be:

- Carve out a secure workspace on the bare-metal (BMC) side and Smartsheet (no-code work management platform) side.
- Focus on one or two priority workflows such as underwriting, asset management, or insurance review/pricing.
- Require only light internal availability at first: i.e., 2-3 hours per week from a senior analyst or process owner, plus 1 hour per week for project governance and decision review.
- No heavy IT involvement until the business is ready.

This gives the client organization a low-burden way to see the future-state workflow working before making a large platform or integration commitment.

SaaS to Insourcing Path

The model can start as SaaS (Software-as-a-Service) / Service-as-a-Software and mature into a deeper client-owned operating capability.



Phase 1: Plug-and-Play BPR Box	Phase 2: Workflow Pilots	Phase 3: Enterprise Hardening
• Secure workspace. • Process mapping. • Prototype. • Weekly feedback loop. • No major internal build required.	• Live or near-live operating workflows. • Controlled data pulls. • Governed evidence packs. • Human approval gates. • Measurable time/risk reduction.	• Deeper IT/security review. • Production integration planning. • Client governance model. • Long-term operating handoff.

This path respects the reality of corporate bandwidth. The client does not need to solve the whole architecture on day one. They only need to sponsor a small, controlled box where the business case can prove itself.

Executive Takeaway

Aiden is not another chatbot layer.

It is a Governed Enterprise Control Tower that connects operating intent, official records, human approval, and evidence.

Its differentiation is the combination of:

- Deterministic process control at the workflow layer.
- Evidence-bound analysis before recommendations.
- Human-gated execution for sensitive actions.
- Private CUDA-independent inference on Intel acceleration.
- A pragmatic BPR-in-a-Box path that lets clients start small and scale once value is proven.

The promise is simple: Make regulated and high-value work faster and easier to run, simpler to explain, and harder to lose control of.