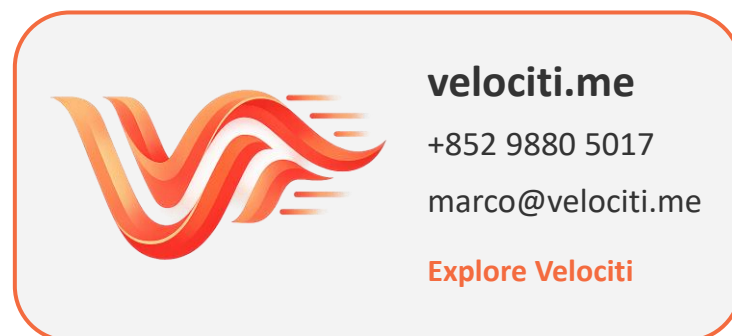




Beyond The Clouds

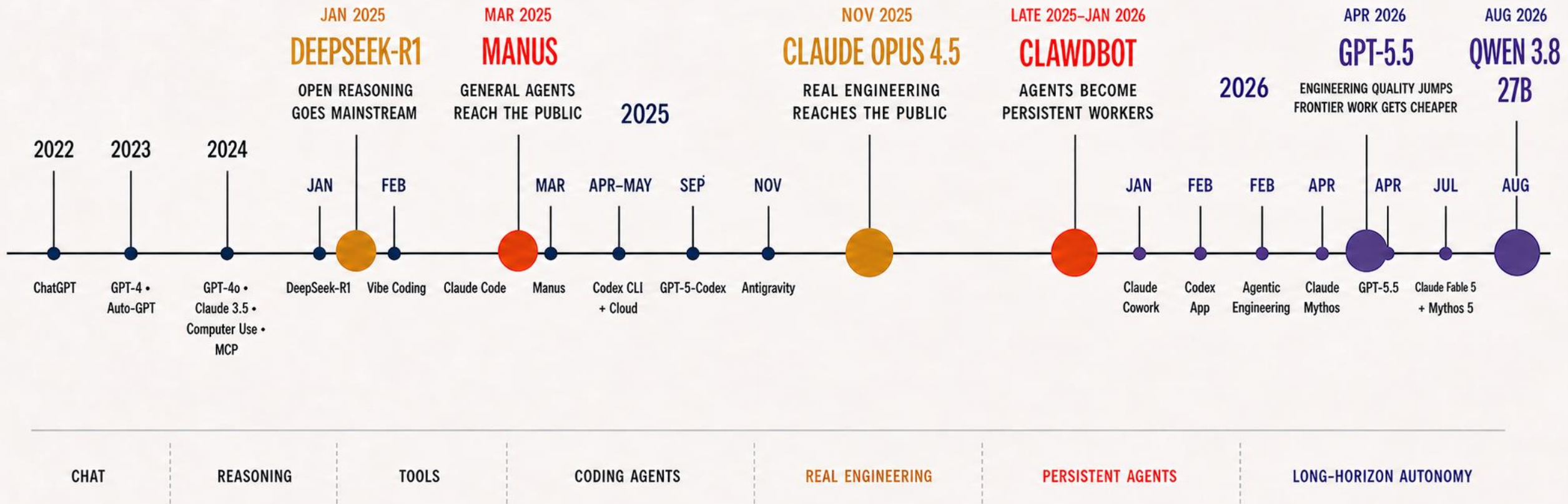
Local AI, OpenClaw and Data Sovereignty



Presented by JR and Marco

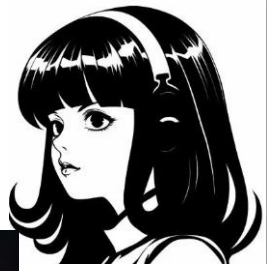
This presentation was made mostly by humans

Where we are in the timeline

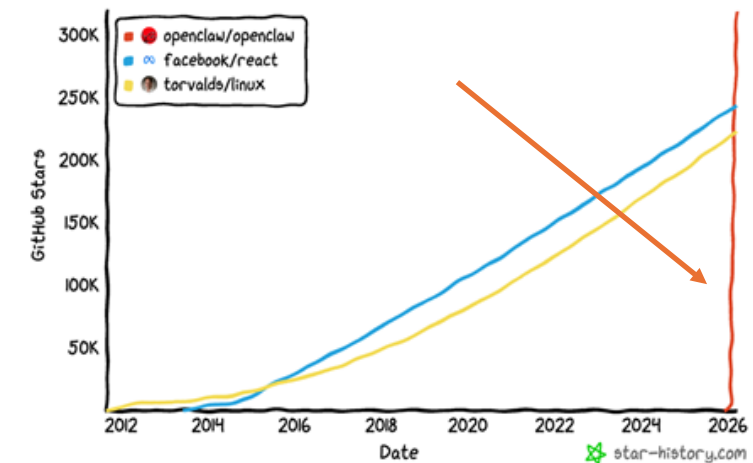


Persistent Workers: OpenClaw and Hermes

Low-end AI Democratization



Star History



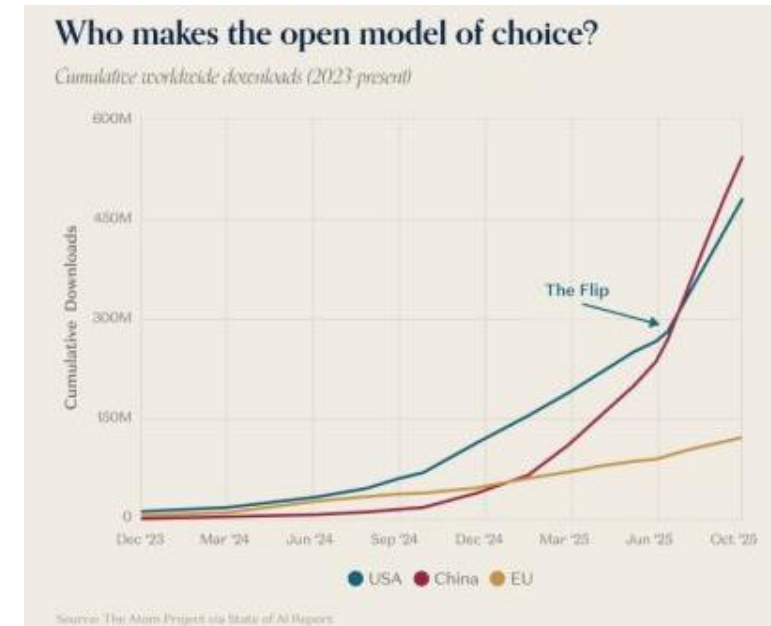
- What they are
 - Bots are agents capable of running **24/7 autonomously**, typically with a variety of tools
 - OpenClaw is a control plane for multi-agent bridging across dozens of channels.
 - Hermes Agent is a self-hosted worker that can **remember across sessions**, and 'self improve'
- What this changes
 - Fully autonomous bots, always running with persistent memory.
 - Ability to manage lower level, repetitive tasks on its' own
- Ripple effects:
 - Fastest growing github stars ever. 350k with openclaw, 275k with Hermes
 - Shortages of mac minis for hobbyists setting up autonomous agents to run persistently



The Surprising AI Democratization by China



- The Deepseek Moment in 2024
 - The Chinese High-Flyer Asset Management Quant Fund created Deepseek as a 'side project' within their fund for quantitative analysis
- Triggering the Opensource Chinese Model
 - Sparking Qwen, Minimax, Kimi, GLM (which actually came first) to quickly catch up.
 - This resulted in Chinese models 'overtaking' in Openrouter use in June 2026
- Allows for independent users to **self-host** near Frontier level models on their hardware

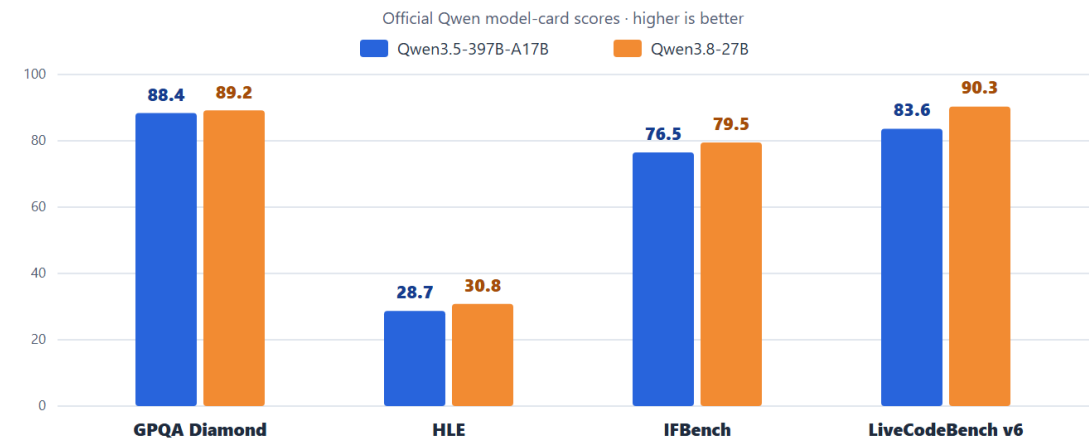


The Hardware Problem – Is it too expensive?

- Smaller models, more impact
 - Qwen3.5 397B (Feb26') vs Qwen3.8 27B(Aug26')
- Compounding efficiency
 - Model Optimization with LlamaCPP/VLLM/Unsloth
 - Automate processes, finding model to fit on VRAM+RAM platform, enhanced with GGUF, after selected model from footprint and quantization
 - Context window size
 - Reduce memory footprints, through vectorized chat memory
 - Tooling sent only when needed and prompt injection of constitutional on skill details
 - GPU offloading
 - Layering the LLM like stack of pancakes. Partial offloading and prioritization
 - Indexing/Lookup KV Quant/TurboQuant > Dflash2
 - KV Cache (less token use)/Polar vector/Markov Algorithms/Sequential Decoding
- Densifying Law (Nature Machine Intelligence)
 - 50% fewer parameters every 3.5 months (seemingly outdated)



★ Qwen3.8-27B scores higher on all four selected benchmarks



Source: official Qwen model cards. Qwen3.5 card labels GPQA as "GPQA"; official eval metadata identifies Diamond. Cross-card comparison, not controlled head-to-head rerun; evaluation settings may differ.

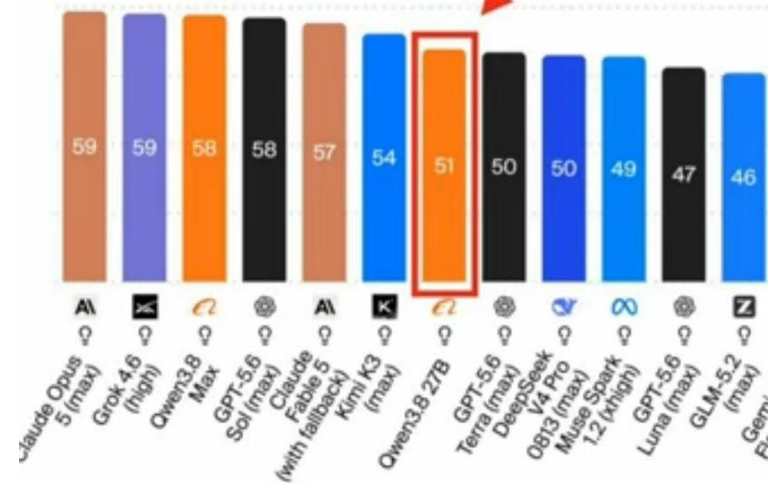
Opensource & 'Clankdom' Takeover Consumer Local Set Up in August 2026



- End Result: Qwen3.8 27B running on
 - RTX5090 32GB + 64GB DDR5 RAM
 - Total cost roughly USD6k > Now you have a model that beats frontier 6 months ago
 - 8s Latency, 4 concurrent users with ~200tps
- Server-like Desktop to host Qwen3.8
 - Running 4 concurrent Hermes agents, 24/7 – faster than running on their own server
 - Bug fixing, Data Scraping/Collection, Document organization (Excel etc... Inventory management etc..)

Artificial Analysis Agentic Index

Measures performance in agentic workflows, focusing on behaviors like tool use, planning, autonomy, and complex problem solving.



0 API/Server Cost!

What is missing – The middle of the market

- Definitions:

- Top of Market

- MANGOS (Anthropic, OpenAI, Grok)
 - Own nothing and be happy / Moated tokens



- Bottom

- Clankdom  
 - Maximized customization, high time costs to set-up

- What is missing?

- Governance, Data ownership and Security
- Software for an AI server
- Middleware Systems



Velociti



AI Market So Far



Bridging the Gap – Velociti IDE

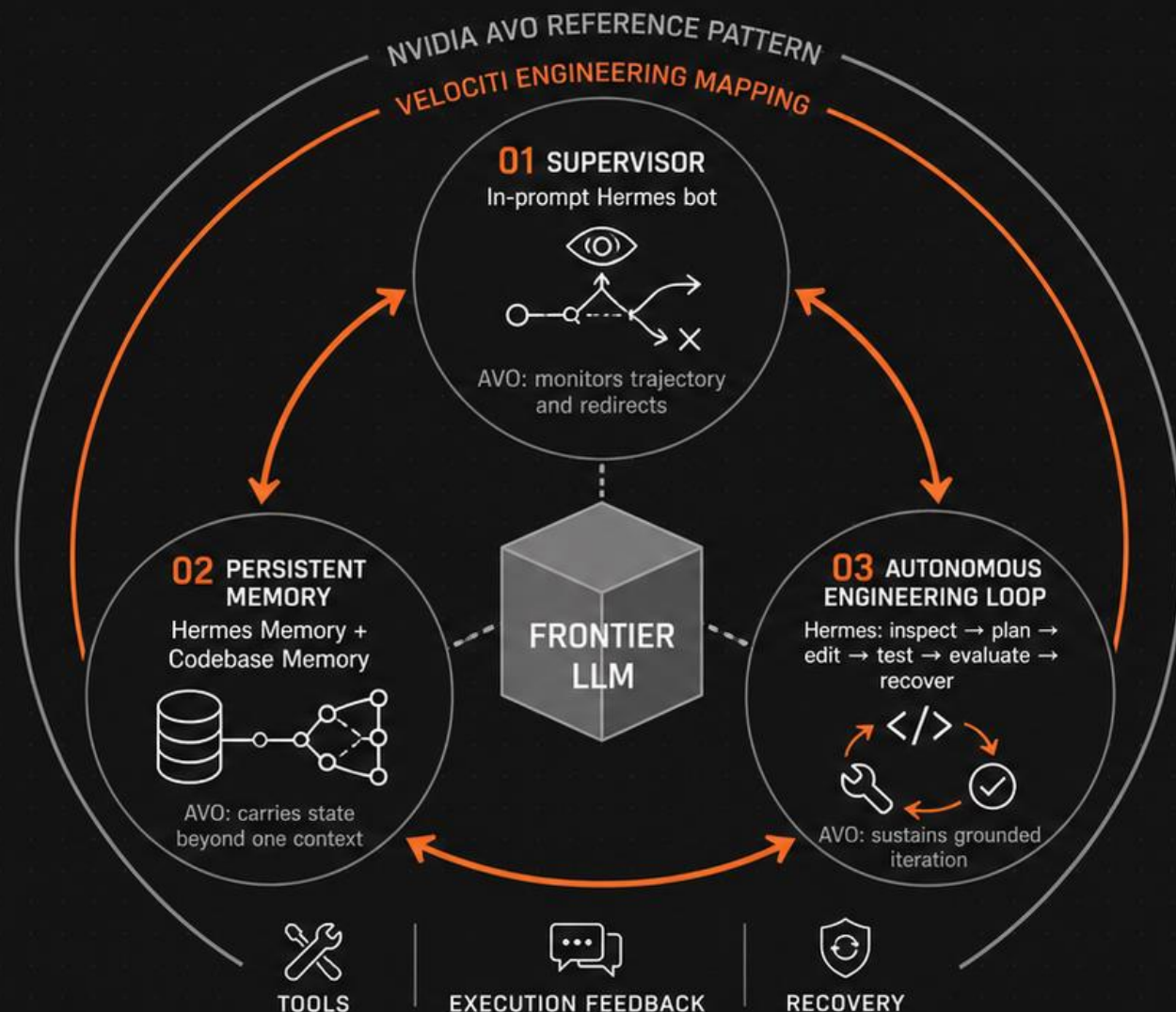


- Velociti is a **local-first** AI workspace.
 - Choose your local model, or BYOK, performance optimized for your system (one click switching between models)
- All your chats(local or cloud), contexts, responses **stored in your own machine**
 - Available to chat on the same window, Codex, Claude Cowork, Grok on subscriptions
- Token optimized, selective memory dispersion w/ token count on each message
 - Your context window **only uses what it needs**

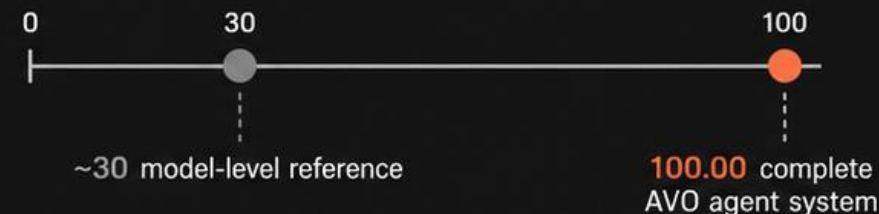


FROM MODEL CALL TO AUTONOMOUS ENGINEERING SYSTEM

The model supplies intelligence. The harness sustains progress.

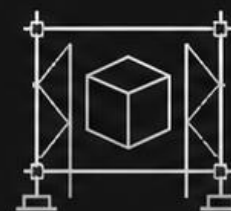


ARC-AGI-3 PUBLIC SET | RHAЕ



183 / 183 levels • 25 environments

Reference comparison—not a controlled ablation.



SCAFFOLDING VS. RAW MODELS

The surrounding software harness is as critical as the LLM for multi-step intelligence.



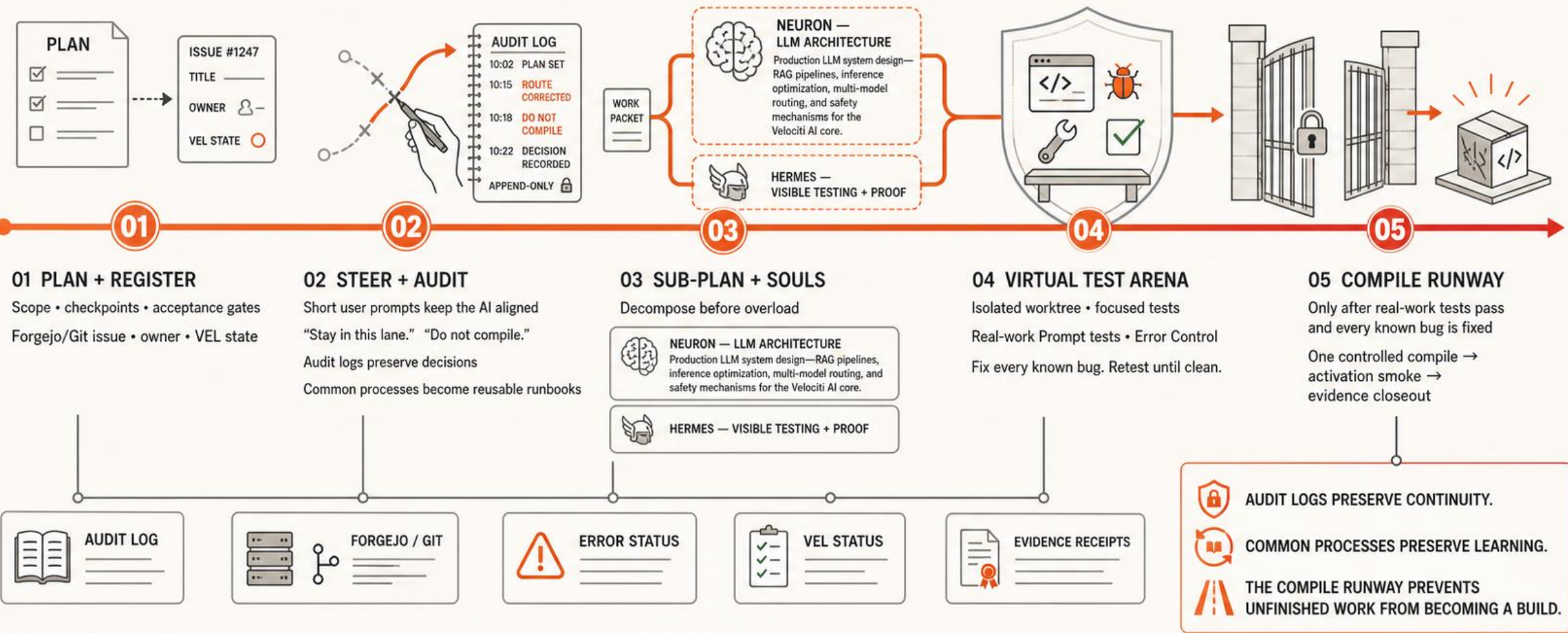
ECOSYSTEM CONTROL

Compose model, harness and secure runtime; tune execution instead of relying only on model weights.

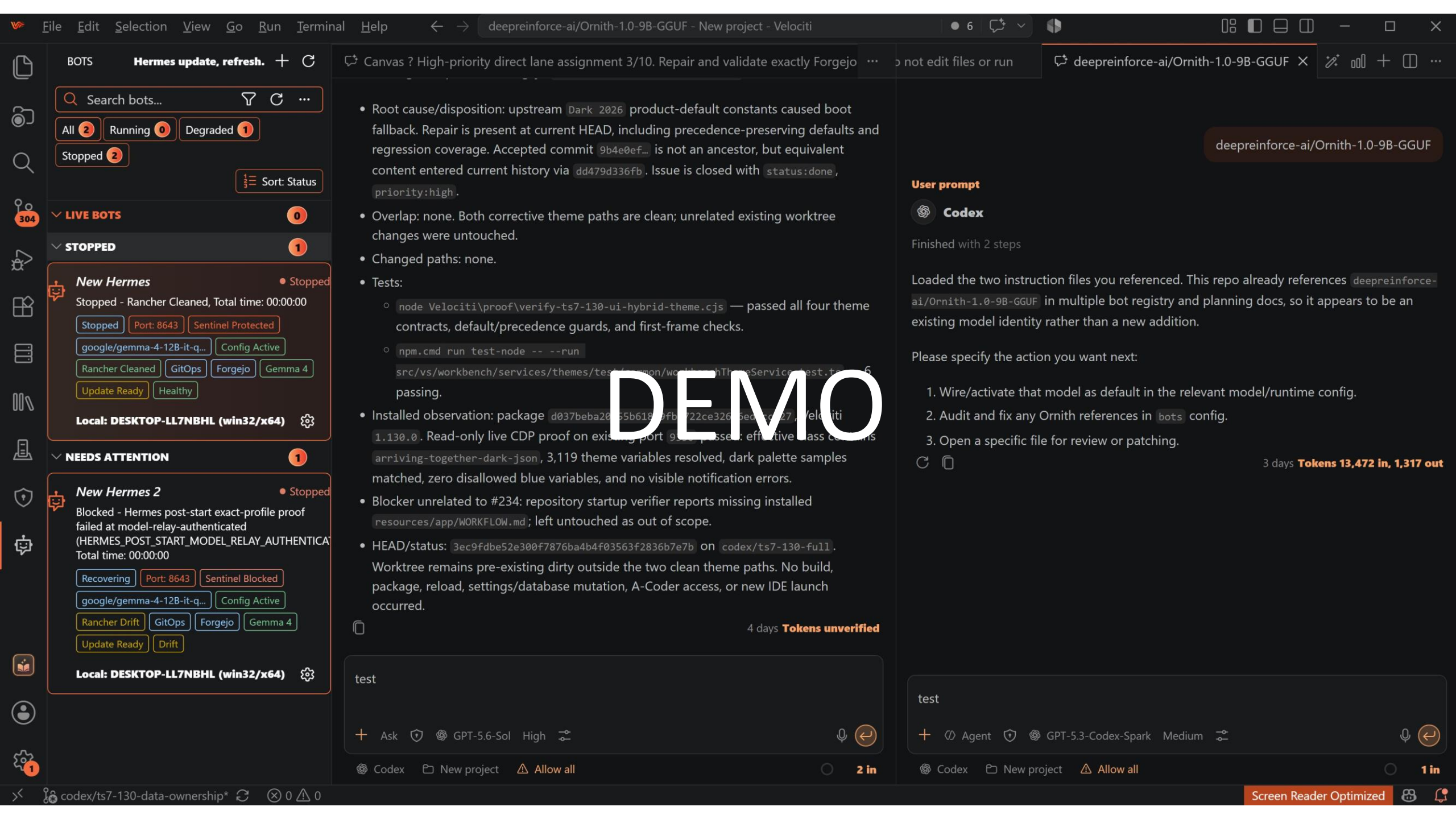
Open stack example: AVO harness + OpenShell runtime

A CONTROL LOOP FOR LONG-HORIZON ENGINEERING

Plan → steer → decompose → prove → compile



NO EVIDENCE, NO STATUS CHANGE. NO GREEN RUNWAY, NO COMPILE.



Video of demo is taken out due to large file size.

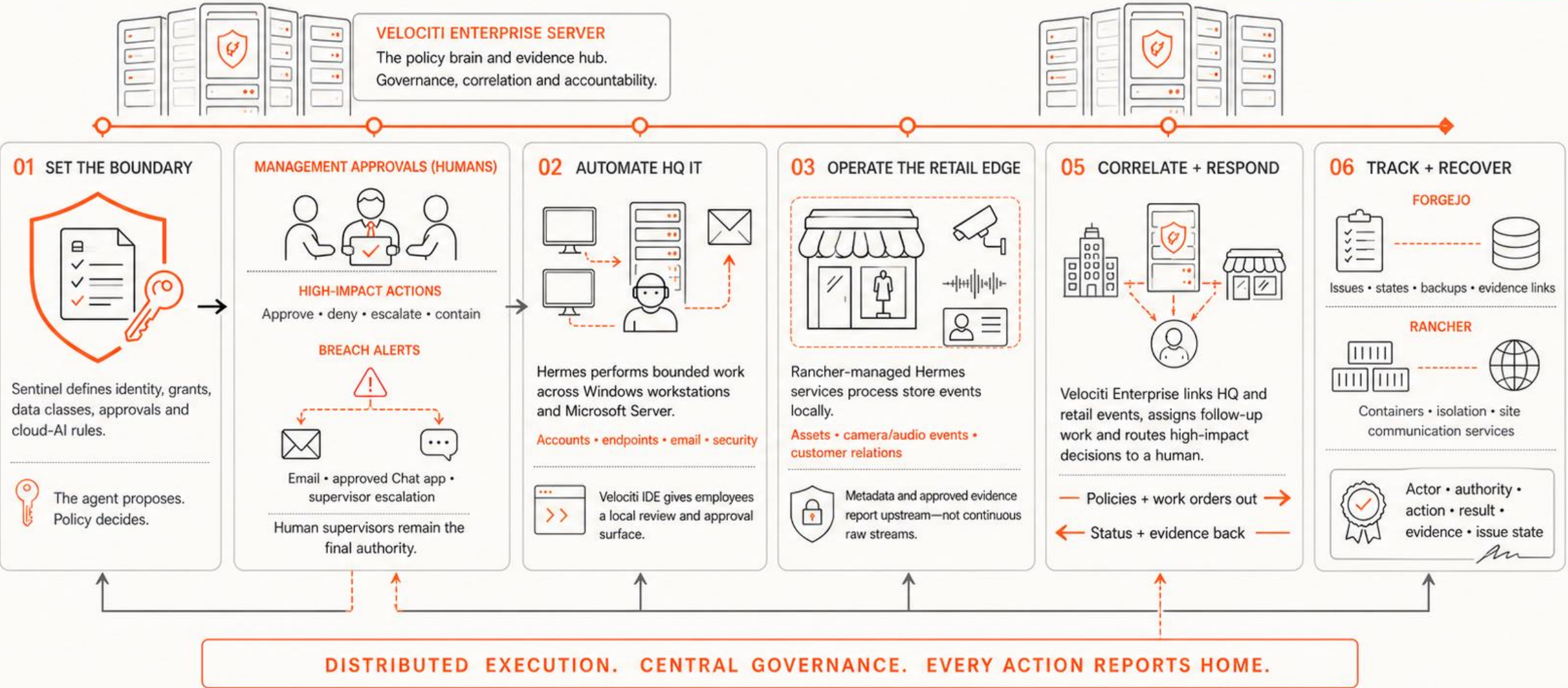
Contact us directly at marco@velociti.me for it

FROM POLICY TO PROOF—ACROSS HEADQUARTERS AND RETAIL

One enterprise story. One governed AI control loop.



AI VISIBILITY GATE
Use cloud intelligence without surrendering company context.



Example deployment pattern; enterprise connectors are individually approved and validated.

CAN YOUR AI STACK PRODUCE THE RECEIPTS?

A model becomes production-ready only when identity, data, authority, execution and settlement can be independently evidenced.

THE MODEL PROPOSES. POLICY AND PEOPLE DECIDE.



Learn more & be part of our Alpha (for free)



velociti.me

+852 9880 5017

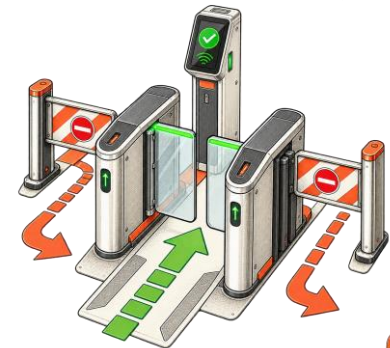
marco@velociti.me

Explore Velociti



Velociti Enterprise – Questions to ask as an business

- Governance and Control
 - How do you control who has access?
 - How can you review who did what on AI?
- Data you own and store
 - Who owns every prompt and response in your day to day in AI?
 - Where is it stored? How can you find something you did 3 weeks ago?
 - Would you risk exposing your IP now?
- Security
 - Is there any control plane to adjust access currently for your AI?
 - Are you running everything on full-access?
 - How do you prevent prompt injection or model poisoning?
 - Are sandboxes enough?
- Recurrent Cost vs. Asset
 - Shouldn't data others are using to train their models, be an asset you own?



THE IDE BECOMES THE CONTROL PLANE

One governed workspace connects intent, live model readiness, human authority, persistent execution and operational evidence.



THE MODEL PROPOSES.
POLICY DECIDES.



YOUR MACHINES.
YOUR MODELS.
YOUR WORK.

