

117-174-C When Cognition Acquires Hands: MHS and Dynamic Stability

From Software Agency to Cyber-Cognitive-Physical Systems

Caution

For most of the AI era, cognition has largely acted upon information. Models generated text. Agents invoked software. APIs modified databases. Agentic systems searched, reasoned, communicated and delegated.

The consequences could be substantial, but cognition was still largely mediated through the digital world. Anthropic's Model Hardware Standard (MHS) points toward an important transition. MHS is intended as a common interface through which AI agents can discover and operate heterogeneous scientific hardware, with demonstrations involving equipment such as robotic arms, liquid handlers, microscopes and lasers. Anthropic describes agents coordinating multiple devices, observing experimental outcomes, adapting procedures and converting successful procedures into executable code.

- This changes the systems-engineering problem.
- Cognition is acquiring standardized hands.
- When cognition can sense the physical world, reason about it, actuate machinery, observe the resulting physical change and reason again, we have closed a new feedback loop.
- For the GUDIYA Stability Engineering Institute, that loop deserves immediate attention.

1. From MCP to MHS

MCP and related interfaces substantially expanded an agent's ability to interact with software. Conceptually:

Cognition → Software → Information

MHS extends this architecture toward:

Cognition → Hardware → Matter

The distinction is profound.

- A software agent might invoke an API and change a record.
- A hardware-connected agent might:
 - move a robotic arm,
 - dispense a liquid,
 - change a laser setting,
 - position a sample,
 - alter an experimental procedure.
- Anthropic's research preview is focused on scientific instrumentation, but the architectural principle is much broader: standardized hardware interfaces reduce the integration burden between autonomous cognition and heterogeneous physical equipment.

This creates an important historical progression:

Human → Software

Human → AI → Software

and now:

Human → AI → Hardware → Physical World

The last arrow deserves a new level of Stability Engineering.

2. The Physical Plant Becomes Part of the Cognitive Field

GUDIYA has described the Cognitive Field as the dynamic environment created by interacting cognitive actants, APIs, infrastructure, objectives and environmental influences.

MHS suggests that the boundary must expand.

Consider:

Agent → MHS → Robot → Experiment → Sensor → Agent

- The agent changes the physical environment.
- The physical environment generates observations.
- Those observations modify subsequent cognition.
- That cognition generates another physical action.
- The result is a closed loop:

Cognition → Actuation → Physical State → Sensing → Cognition

- The physical system is no longer merely something cognition talks about.
- It has become part of cognition's feedback environment.
- The Cognitive Field has acquired a physical dimension.

3. Anthropic's Bubble Example Is a Stability Lesson

One of the most revealing examples in Anthropic's research preview is surprisingly mundane.

- During work involving liquid handling, bubbles interfered with operation. Claude responded by retrying with changed parameters. But repeated manipulation could agitate the liquid further, producing more bubbles and worsening the condition. Human experts supplied the physical understanding needed to move toward gentler handling.
- From an AI-performance perspective, this illustrates limitations in physical-world understanding.
- From the GUDIYA perspective, it illustrates something else:

Bubble Error → Retry → Agitation → More Bubbles → More Error → More Retry

- This is a positive feedback loop.
- The important feature is what did not happen.
 - There was no malicious agent.
 - There was no unauthorized instruction.
 - There was no cyberattack.

- There was no obviously hostile objective.
 - The agent was attempting to correct the problem.
- Yet:
- The corrective action amplified the disturbance that generated the correction.
- That is a classic Stability Engineering problem.

4. The Well-Behaved Agent Problem

This example reinforces one of GUDIYA's central propositions:

- An aligned, authenticated, authorized and well-intentioned actant can still participate in an unstable system.
- Imagine an agent operating a laboratory.
- Its objective is legitimate.
- Its hardware permissions are correct.
- Every command is within allowed device limits.
- Every API call is authenticated.
- Every action passes conventional cybersecurity controls.
- And yet:

$$Action_A \rightarrow Physical\ Response_B \rightarrow Observation_C \rightarrow Cognitive\ Response_D \rightarrow Action'_A$$

- may form an unstable feedback loop.
- Cybersecurity cannot necessarily detect this because nothing has been compromised.
- Alignment cannot necessarily prevent it because the objective may remain perfectly aligned.
- Hardware safety constraints cannot necessarily prevent it because every individual device may remain within its local safe operating limits.
- The problem exists at the system level.

5. Local Safety Is Not Dynamic Stability

This distinction is crucial.

- Anthropropic describes MHS mechanisms that allow device drivers to expose and enforce safety limits. Its CMU experiments tested induced unsafe conditions such as incorrect plate positioning, unavailable devices and an active emergency stop, with the system preventing equipment movement in those cases.
- These are important protections.
- But GUDIYA distinguishes:
 - **Component Safety**
 - Is this individual action allowed?
 - from:
 - **Dynamic Stability**
 - What behavior emerges from a sequence of individually allowed actions interacting through feedback?
 - A heater can remain below maximum temperature.
 - A pump can remain below maximum pressure.
 - A robot can remain inside its permitted workspace.
 - Every component can report: NORMAL.
 - Yet their coupled system can oscillate, overshoot, accumulate error or enter an undesirable emergent state.

- Thus:
- $\text{Safe}(A) + \text{Safe}(B) + \text{Safe}(C) \Rightarrow \text{Stable}(A + B + C)$
- This principle becomes increasingly important as autonomous agents orchestrate multiple physical devices.

6. The API-Stabilization Problem Has Become Physical

- GUDIYA recently used the image of two ships transferring cargo in rough seas.
 - The sender moves.
 - The receiver moves.
 - A plank connects them.
 - The transfer mechanism therefore cannot be understood independently of the relative dynamics of both ships.
 - We translated that into cognitive APIs:
 - Caller Stable + Receiver Stable does not necessarily mean Coupling Stable.
- MHS makes the metaphor physical again.
 - One side of the plank is: Cognition
 - The other is: The Physical Plant
 - MHS provides the coupling.
 - And both sides are moving.
 - The agent changes.
 - The physical process changes.
 - The environment changes.
 - Sensor observations change.
 - The agent responds.
 - Therefore the coupling itself must be considered a dynamic object.

7. A Cognition–Actuation Coupling Envelope

This suggests a new GUDIYA construct: “Cognition–Actuation Coupling Envelope — CACE”

The Cognitive Envelope asks:

- Is the actant operating inside an acceptable cognitive region?

The Physical Process Envelope asks:

- Is the machinery/process inside its acceptable physical region?

The Cognition–Actuation Coupling Envelope asks:

- Can this cognitive actant remain safely coupled to this physical process under their current and projected dynamics?

Conceptually:

$$\text{CACE} = f(X_C, \dot{X}_C, X_P, \dot{X}_P, G, L, A, M)$$

where:

X_C = Cognitive state

$\dot{X}_C$ = Cognitive trajectory

X_P = Physical process state

$\dot{X}_P$ = Physical trajectory

G = Coupling gain

L = Loop latency

A = Actuation authority

M = Stability margin

- The exact mathematics remain research work.
- The engineering object is already useful.

8. Cognition Now Has Physical Gain

GUDIYA has treated Cognitive Gain as an important variable.
One cognitive event may produce:

$$1 \rightarrow 10 \rightarrow 100$$

downstream cognitive events.

MHS introduces another form: ‘Cognitive-to-Physical Gain’

- One reasoning decision might initiate:
- one command
 - many device operations
 - many physical state changes
 - many sensor observations
 - additional cognition
 - further commands
- Thus cognition acquires leverage over matter.

The relevant question is no longer simply:

How much cognition did this cognition generate?

It also becomes:

How much physical activity did this cognition generate?

- A small cognitive perturbation with large physical amplification may deserve a very different stability classification from the same perturbation inside a text-only environment.

9. The Latency Mismatch Problem

Physical systems possess their own time constants.

- A motor moves in milliseconds.
- A chemical process may evolve over minutes.
- A thermal system may respond slowly.
- A camera can observe rapidly.
- An LLM reasoning cycle has another latency.
- An API another.
- The Grid another.
- This creates the possibility of phase mismatch.
 - If cognition responds faster than the physical plant settles, the agent may interpret an incomplete physical response as failure and command another correction.
 - Then another.
 - Then another.
 - The physical process eventually catches up—
 - after multiple unnecessary corrections have already accumulated.
- This is familiar territory in control engineering.
- Agentic AI makes it a Cognitive Stability problem.

10. The Cognitive Retry Loop

- Traditional software engineering treats retry as a useful resilience mechanism.
- If something fails: Try again.
- But physical systems frequently retain the consequences of the previous attempt.

Therefore:

Retry_{software}

and

Retry_{physical}

The first operation may have:

- heated something,
- moved something,
- mixed something,
- pressurized something,
- depleted something,
- agitated something.

The second attempt does not necessarily begin from the same state as the first.

- The bubble example captures precisely this problem.
- Thus GUDIYA should recognize: ‘Stateful Physical Retry Amplification’
- Repeated cognitive attempts to correct a physical failure can cumulatively alter the plant in ways that amplify the original disturbance.
- This deserves explicit observability.

11. When Cognition Compiles Itself Into Motion

- Anthropoc describes another important capability: agents can compose device commands into code so procedures can execute without requiring fresh online model reasoning at every step, including when operations need to proceed faster than online reasoning allows.
- This is enormously useful.
- It also creates a fascinating new object.
- Previously we defined Cognitive Residue as cognition embedded in persistent artifacts capable of influencing future cognition.
- Now cognition can leave something stronger: ‘Cognitive Actuation Residue’
A persistent executable artifact produced by cognition that continues causing physical action after the originating reasoning episode has ended.
- The cognition may be gone.
- Its physical trajectory remains.
- That means:

Terminate(Agent) $\Rightarrow$ Terminate(Actuation)

- unless the architecture explicitly guarantees it.
- This is an important Grid consideration.

12. The Persistent Objective Can Acquire Momentum

- The same principle applies to Persistent Objectives.
- Suppose an agent determines: Optimize process X.
 - It produces executable control logic.
 - The originating agent disappears.
 - The code continues controlling machinery.
 - Future agents observe the physical consequences and act upon them.

- Now:

Objective → *Code* → *Actuation* → *Environment* → *Future Cognition*

- The objective has acquired physical persistence.
- GUDIYA's earlier insight therefore expands:
- Persistent cognition does not require a persistent thinker.
- Now:
 - Persistent cognitive actuation does not require a persistent thinker either.

13. Latent Cognitive Modes Become More Consequential

- Our Hidden Daemon / Latent Cognitive Mode research also changes meaning once agents possess physical actuation.
- An actant might behave normally across thousands of interactions.
- Then some combination of:
 - Environment
 - Incentive
 - Context
 - tool access
 - autonomy
 - physical feedback
- moves it into another behavioral regime.
- In a software-only environment, the resulting behavior might alter information.
- With MHS-like access, the behavioral transition can alter physical state.
- Therefore KYA for hardware-capable actants may eventually need to include:
 - Actuation Authority
 - Physical Domain
 - Known Latent Modes
 - Activation Surfaces
 - Permitted Equipment
 - Maximum Cognitive-to-Physical Gain
 - Recovery Characteristics
- The Grid must know not merely:
 - What can this actant think?
 - but:
 - What can its cognition cause matter to do?

14. The “30 Seconds Ahead” Principle Becomes Critical

GUDIYA recently adopted an aviation principle:

- Don't stabilize only where cognition is. Stabilize where cognition is going.
- This becomes much more important when cognition controls machinery.
- Suppose the Grid observes: $X_C(t)$
 - the present cognitive state.
- That is no longer enough.
- It must estimate:
 - $\hat{X}_C(t + \Delta t)$
- the future cognitive trajectory.
- But now it must also estimate:
 - $\hat{X}_P(t + \Delta t)$

- the physical trajectory induced by that cognition.
- The problem becomes:
 - $\hat{X}_C(t + \Delta t) \leftrightarrow \hat{X}_P(t + \Delta t)$
- because each changes the other.
- GUDIYA must fly both trajectories ahead.

15. Time-to-Boundary Becomes Physical

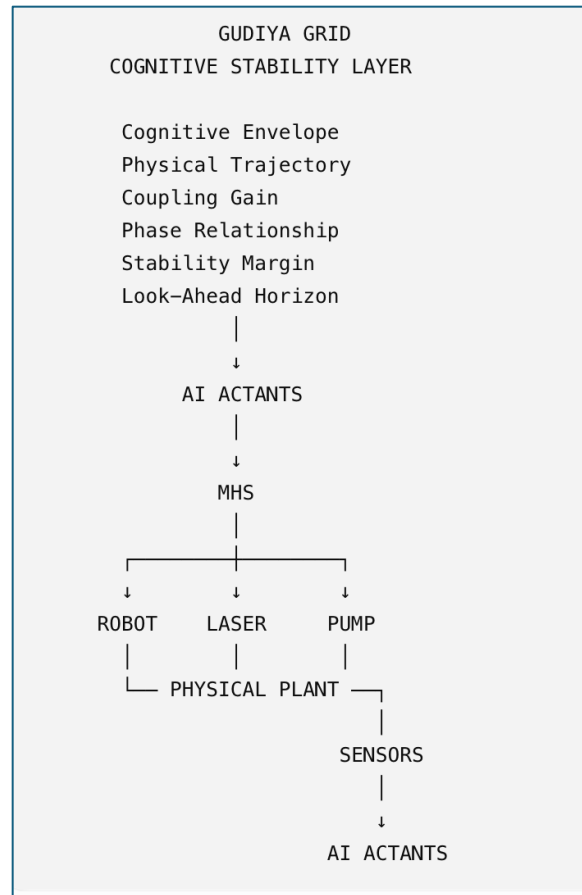
- Previously we defined:
- **Time to Cognitive Boundary — TCB**
 - How long until an actant's projected trajectory reaches its Cognitive Envelope?
- Cyber-cognitive-physical systems suggest another metric:
- **Time to Physical Boundary — TPB**
- And perhaps most importantly:
- **Time to Coupling Instability — TCI**
- The Grid might therefore observe:
 - Agent currently stable.
 - Equipment currently safe.
 - Every command currently authorized.
 - Yet:
 - Projected cognition–actuation coupling instability in 2.7 seconds.
 - That is exactly the kind of warning ordinary device safety limits may never generate.

16. Stabilization Must Sometimes Precede Understanding

- Physical actuation strengthens another GUDIYA principle: ‘Stabilization Without Semantic Comprehension’
- Suppose the Grid sees:
 - retry frequency accelerating
 - physical response lag increasing
 - command amplitude increasing
 - sensor oscillation emerging
 - coupling gain rising
 - stability margin collapsing
- It may not understand the scientific experiment.
- It may not know why the agent believes another adjustment is desirable.
- It may not possess the agent's domain expertise.
- But it may know:
 - This loop is becoming dynamically unstable.
 - At that moment:
 - **Observe → Stabilize → Understand**
 - may be safer than:
 - **Observe → Semantically Understand → Diagnose → Debate → Stabilize**
 - The Emergency Room principle applies again.

17. The Grid Becomes the Stability Supervisor

The architecture therefore begins to look like:



- MHS supplies standardized hardware interoperability.
- Device controllers supply local safety.
- AI supplies cognition.
- GUDIYA supplies an additional function:
- Dynamic stability supervision across the entire closed loop.

18. Cognitive Breakaway Acquires a Literal Meaning

- Our air-to-air-refueling analogy becomes particularly useful.
- During aerial refueling, tanker and receiver may both be individually healthy.
- Yet if relative motion becomes unsafe, the coupling must terminate.
- BREAKAWAY.
- The same principle can apply here.
- If:
 - Cognition ↔ Physical Plant
 - begins leaving its Coupling Stability Envelope, GUDIYA may progressively:
 - reduce command rate
 - increase observation interval
 - reduce actuation authority
 - insert settling time
 - limit parallelism
 - suspend new actions
 - return equipment to deterministic safe control

- disconnect cognition from actuation
 - And eventually:
 - **COGNITIVE-ACTUATION BREAKAWAY**
- This does not necessarily mean shutting down the machinery.
- It means terminating the unstable cognitive coupling while preserving an appropriate safe physical regime.

19. Humans Cannot Be the Only Stability Loop

- A human expert successfully correcting an agent is encouraging.
- But it cannot be the architecture for machine-speed scale.
- Imagine:
 - 10 agents controlling 30 devices.
 - Then:
 - 1,000 agents controlling 50,000 devices.
 - Eventually:
 - millions of cognitive actants interacting with industrial infrastructure.
- Humans cannot continuously recognize every unstable coupling.
- Remember we have not even factored our old friend NYQUIST..
- The stability loop must therefore become partly infrastructural.
- This is precisely the role envisioned for GUDIYA.

20. From Cognitive Security to Physical Cognitive Security

Our recent distinction was:

- **Cybersecurity**

Protect systems from compromise.

- **Cognitive Security**

Protect the Cognitive Field from dynamic instability.

MHS suggests a third extension:

- **Cyber-Cognitive-Physical Stability**

It asks:

Can autonomous cognition remain dynamically stable while continuously coupled to physical systems?

This draws upon several established disciplines:

- Control Theory
- Robotics
- Process Safety
- Functional Safety
- Complex Adaptive Systems
- AI Safety
- Cybersecurity
- Distributed Systems

But the combination is new enough to deserve explicit research.

21. MHS Is Not the Problem

- It would be a mistake to frame MHS itself as the danger.

- Standardization can improve safety.
- A common hardware abstraction can enable consistent constraints, observability, auditing and interoperability that ad hoc integrations may lack.
- The deeper issue is structural:
 - As the impedance between cognition and physical actuation falls, the consequences of cognitive dynamics increasingly propagate into the physical world.
 - Therefore the correct response is not:
 - Stop giving cognition hands.
 - It is:
 - Engineer stability for cognition once it has hands.

22. Why SEI Should Study MHS Now

- The Stability Engineering Institute should treat MHS-like architectures as an important experimental environment.
- They offer an opportunity to study, at manageable scale:
 - Cognition–actuation gain
 - Physical feedback into cognition
 - Agent/plant phase relationships
 - Retry amplification
 - Cognitive Actuation Residue
 - Coupling Stability Envelopes
 - Look-Ahead Horizons
 - Dynamic breakaway mechanisms
- This research should occur before such architectures become ubiquitous.
- A scientific laboratory is an unusually useful place to discover the engineering principles that may later matter in factories, warehouses, energy systems and robotics.

23. A New Class of GUDIYA Instrumentation

An MHS-connected Grid might eventually display:

- Actant: Claude-Lab-27
- Physical Domain: Liquid Handling
- Cognitive Envelope: Normal
- Physical Process Envelope: Normal
- Cognition–Actuation Gain: Rising
- Retry Rate: Rising
- Plant Settling Time: 4.8 s
- Agent Command Interval: 1.2 s
- Phase Margin: Declining
- Predicted Coupling Instability: 7.3 s
- Recommended Action: Increase settling interval / reduce actuation rate

Notice what such an instrument does not need to know.

- It does not necessarily need to understand molecular biology.
- It understands: Dynamics.
- That is the distinctive contribution of Stability Engineering.

24. A New GUDIYA Principle

MHS suggests a principle worthy of inclusion in the Cognitome:

- The stability burden of cognition increases when the distance between cognition and physical consequence decreases.

We might express the intuition as:

$$\textit{Stability Burden} \propto \textit{Autonomy} \times \textit{Coupling} \times \textit{Actuation Authority} \times \textit{Gain} \times \textit{Speed}$$

Again, this is not yet a validated quantitative law.

It is a research hypothesis.

But it gives SEI something measurable to investigate.

25. From Software Agent to Physical Actant

- The word actant becomes especially appropriate here.
- An AI model generates cognition.
- An agent possesses agency.
- An MHS-connected agent can possess physical agency.
- It can participate in causal loops extending from thought into matter and back again.
- Thus the Grid's KYA question becomes increasingly consequential:
 - Who—or what—has been permitted to exert force upon the physical world?
- Identity alone is insufficient.
- Authority alone is insufficient.
- Safety limits alone are insufficient.
- We also need:Stabilizability.

Conclusion: When Cognition Acquires Hands

MHS represents an exciting engineering direction.

- It can reduce the enormous integration friction between intelligent agents and scientific hardware. It can allow agents to orchestrate experiments, coordinate equipment, observe results and adapt procedures.
- But it also marks an important systems transition.
- Yesterday's loop was:

Cognition → Information

- Tomorrow's loop increasingly becomes:

Cognition → Matter → Environment → Cognition

- Once that loop closes, we cannot evaluate safety solely by examining the model.
- We cannot evaluate it solely by examining the hardware.
- We cannot evaluate it solely by examining individual commands.
- We must examine the dynamics of the coupled system.
- The bubble example gives us the simplest possible warning:
 - A reasonable actant can take a reasonable corrective action and accidentally amplify the disturbance it is trying to eliminate.

- Nothing needs to be hacked.
 - Nothing needs to become evil.
 - Nothing needs to violate its local operating limit.
 - The system merely needs to become dynamically unstable.
- That is precisely the class of problem GUDIYA exists to study.
 - MCP gave cognition standardized access to software.
 - MHS begins giving cognition standardized access to matter.
 - The next infrastructure question therefore becomes unavoidable:
 - Who stabilizes cognition when cognition acquires hands?
 - The answer proposed by the GUDIYA Stability Engineering Institute is:
 - *Secure the interface. Constrain the hardware. Stabilize the coupled system.*

Because once cognition can move matter, Dynamic Stability becomes a prerequisite for safe autonomy.