

# INTERLOOM

## Multi-Agent Reasoning for Cascading Consequences in Interconnected Systems

SEE WHAT HAPPENS NEXT.

TECHNICAL PAPER  
VERSION 1.0 / JULY 2026

A technical architecture for turning events and policy decisions into inspectable, evidence-aware graphs of cross-domain consequences.

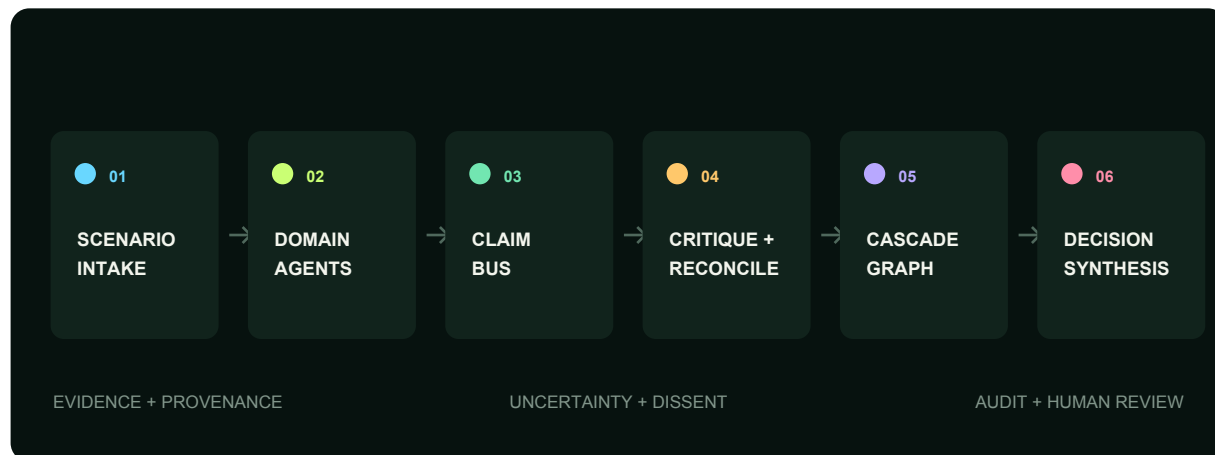
## ABSTRACT

# A reasoning surface for decisions that cross systems

Many high-stakes decisions are difficult not because their first-order effects are unknown, but because later effects cross disciplinary and institutional boundaries. A reduction in rainfall can alter reservoir operations, cropping patterns, food prices, public health, migration, and infrastructure demand. Conventional dashboards expose measurements while leaving these propagation paths implicit.

Interloom is a proposed multi-agent reasoning architecture that turns an event or policy intervention into an inspectable graph of consequences. Domain-specialized agents propose claims, exchange structured evidence, challenge weak links, and preserve uncertainty. A synthesis layer produces a decision-facing causal narrative without presenting speculative reasoning as forecast certainty.

What happens next, through which mechanism, with what evidence, and where is the chain most uncertain?



This paper defines Interloom's problem formulation, architecture, agent protocol, evidence model, uncertainty representation, safety constraints, evaluation framework, and implementation roadmap. A planetary-systems prototype is the flagship demonstration, but the architecture is domain-general.

# The decision gap

Complex risks propagate through coupled natural, technical, economic, and social systems. The IPCC Sixth Assessment Report describes compound and cascading risks as impacts that spread across sectors, boundaries, and socioeconomic and natural systems. A decision interface that treats water, energy, health, and economics as independent categories can therefore hide the structure that matters most.

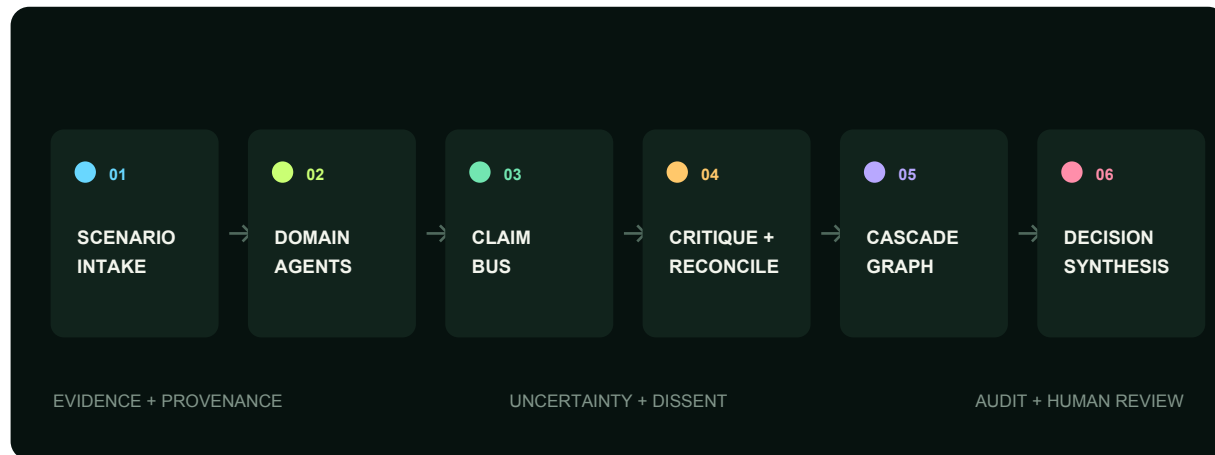
## Interloom sits between three familiar tool classes:

| TOOL CLASS            | PRIMARY QUESTION                       | LIMITATION                                       |
|-----------------------|----------------------------------------|--------------------------------------------------|
| Monitoring dashboards | What is happening?                     | Relationships remain implicit.                   |
| Forecasting models    | What might a selected variable become? | Scope is bounded to model outputs.               |
| General assistants    | How can this be explained?             | Evidence and handoffs are often opaque.          |
| Interloom             | What happens next, and why?            | Designed for inspectable cross-domain reasoning. |

## Design principles

- Specialize before synthesizing. Distinct agents maintain domain-specific assumptions, evidence standards, and failure modes.
- Make causal handoffs explicit. Every downstream claim identifies the upstream condition it depends on.
- Separate evidence from inference. Observations, model outputs, assumptions, and generated hypotheses are different data types.
- Preserve dissent. Conflicting agent views remain visible until resolved or recorded as uncertainty.
- Optimize for inspectability. Users can trace each conclusion to sources, transformations, and critiques.
- Support decisions, not authority substitution. Interloom surfaces options while accountable choices remain human.

# Six layers, one inspectable graph



## 3.1 Scenario intake

Converts a user statement into a structured scenario: event or policy, geography, magnitude, duration, time horizon, baseline assumptions, affected population, and decision objective. Ambiguous fields become clarifications or bounded alternatives.

## 3.2 Domain agents

Each agent combines a bounded mandate, domain tools, typed input and output schemas, validation rules, calibrated confidence behavior, and explicit abstention conditions. A planetary deployment may include Climate, Water, Agriculture, Biodiversity, Energy, Infrastructure, Economy, Public Health, and Migration.

## 3.3 Causal message bus

Agents communicate through typed Consequence Claims rather than unrestricted conversation. This creates a machine-auditable interface between domains.

## 3.4 Critique and reconciliation

Receiving agents challenge causal validity, missing variables, unsupported magnitude, and temporal mismatch. Claims can be accepted, revised, branched, rejected, or marked unresolved.

## 3.5 Cascade graph

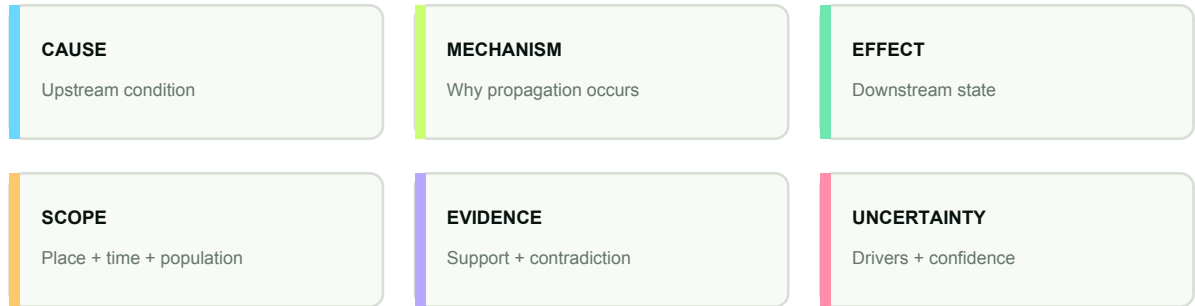
Accepted and unresolved claims form a directed property graph. Nodes represent events, states, populations, assets, and interventions. Edges encode causal or dependency relations.

## 3.6 Decision synthesis

The graph becomes a consequence chain, a set of weak links and disagreements, leverage points, and monitoring indicators that could confirm or falsify a pathway.

# Structured collaboration, not open-ended chatter

An agent is useful only when its output can be evaluated, routed, and revised. Interloom's atomic unit is a typed Consequence Claim. The structure forces every agent to identify a mechanism, scope, evidence basis, and uncertainty before a claim enters the cascade graph.



## Run lifecycle

|    |                                                                                                 |
|----|-------------------------------------------------------------------------------------------------|
| 01 | DECOMPOSE<br>Extract variables and select the smallest competent agent set.                     |
| 02 | PROPOSE<br>Generate independent first-order claims before agents see one another's conclusions. |
| 03 | ROUTE + EXTEND<br>Send claims to affected domains and propose downstream effects.               |
| 04 | CHALLENGE<br>Inspect evidence, mechanism, scope, magnitude, and omitted variables.              |
| 05 | RECONCILE<br>Accept, revise, branch, reject, or preserve as unresolved.                         |
| 06 | STOP + SYNTHESIZE<br>End low-value expansion and render the decision-facing graph.              |

# Every edge must explain itself

Interloom should retrieve from an approved source registry rather than rely on model memory for decision-relevant claims. Evidence objects include a stable source identifier, publication date, geography, methodology, measured variable, applicable time range, and quality rating. Provenance is attached at the claim-edge level because a source may support one mechanism without supporting the entire cascade.

| EVIDENCE CLASS           | HOW THE INTERFACE LABELS IT                            |
|--------------------------|--------------------------------------------------------|
| Direct observation       | Measured relationship in the applicable context        |
| Model estimate           | Output from a named model under stated assumptions     |
| Institutional assessment | Expert synthesis with confidence language              |
| Transferred analogy      | Evidence from another geography or population          |
| Generated hypothesis     | A question for investigation, not an established claim |

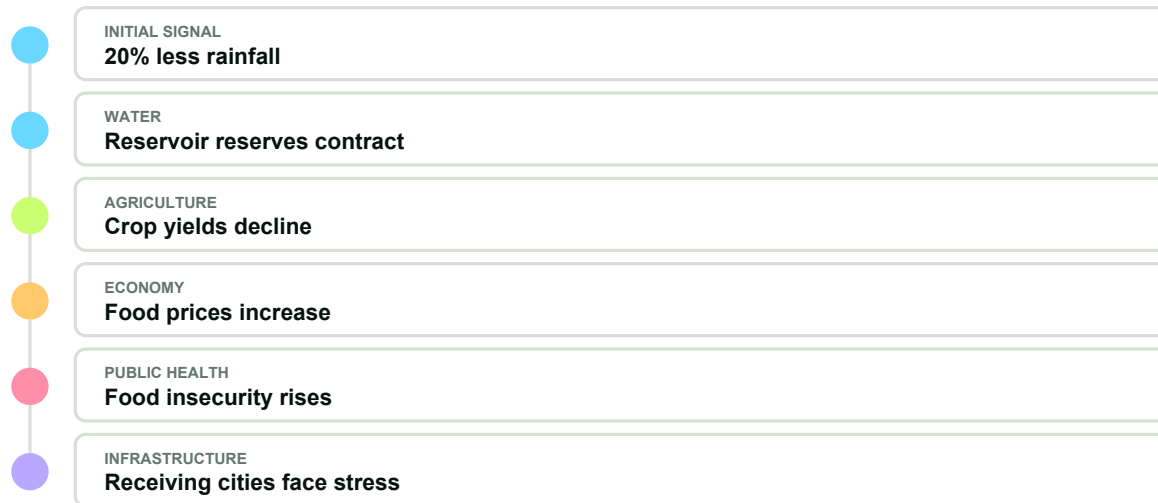
## Four uncertainty dimensions

|                                                     |                                                     |
|-----------------------------------------------------|-----------------------------------------------------|
| <b>EVIDENCE</b>                                     | <b>MODEL</b>                                        |
| Quality and agreement of supporting sources.        | Sensitivity to model choice and functional form.    |
| <b>SCENARIO</b>                                     | <b>SCOPE</b>                                        |
| Dependence on future conditions and human response. | Transferability across time, place, and population. |

A weak upstream edge limits confidence in dependent paths unless the downstream effect has independent support. This prevents long chains from accumulating false certainty.

# A drought becomes a system in motion

Consider the scenario: A severe drought reduces annual rainfall by 20 percent in California's Central Valley. The public prototype renders the cross-domain propagation as a sequence of agent handoffs. In a production run, every transition would be backed by an evidence object and an explicit claim schema.



## What the output is - and is not

The output is an evidence-ready map of plausible consequences, mechanisms, uncertainties, and interventions. It is not a deterministic forecast. The distinction matters: Interloom's value is in making dependencies legible enough to question and act upon, not in collapsing a complex future into one number.

Highest leverage hypothesis: early water allocation and demand controls may soften multiple downstream agricultural and health impacts.

Key uncertainty: adaptive behavior, including crop switching, price response, and migration, changes the scale and timing of later effects.

06 / SAFETY + GOVERNANCE

# Safety is part of the architecture

## EVIDENCE GATES

Claims without adequate support are labeled as hypotheses or withheld.

## HUMAN REVIEW

High-stakes domains require domain-owner sign-off before use.

## CAUSAL HYGIENE

Agents cannot silently convert correlation into causation.

## VISIBLE DISSENT

Minority conclusions remain available with their evidence.

## AUDIT TRAIL

Inputs, sources, tools, claims, revisions, and synthesis are logged.

## DISTRIBUTIONAL REVIEW

Red-team agents test for omitted communities and concentrated harms.

LIMITATION

## Evaluation framework

Claim level: factual support, causal completeness, scope consistency, confidence calibration, contradiction detection, and abstention.

Graph level: path precision and recall, discovery of cross-domain effects, cycle control, uncertainty propagation, scenario sensitivity, and run stability.

Human level: time to identify downstream risk, ability to explain mechanisms, recognition of weak links, intervention quality, and calibrated trust.

# From legible prototype to governed reasoning system

|                                                                                                                   |                            |         |
|-------------------------------------------------------------------------------------------------------------------|----------------------------|---------|
| PHASE 1                                                                                                           | LEGIBLE PROTOTYPE          | CURRENT |
| Curated scenario traces demonstrate the interaction model, causal graph, confidence display, and decision lens.   |                            |         |
| PHASE 2                                                                                                           | EVIDENCE-GROUNDED SCENARIO | NEXT    |
| Implement structured intake, domain retrieval, typed claims, citations, and one expert-reviewed drought workflow. |                            |         |
| PHASE 3                                                                                                           | DYNAMIC ORCHESTRATION      |         |
| Add routing, critique, reconciliation, stopping criteria, graph persistence, and run observability.               |                            |         |
| PHASE 4                                                                                                           | EVALUATION + GOVERNANCE    |         |
| Build benchmarks, calibration tests, distributional review, audit exports, and deployment policies.               |                            |         |
| PHASE 5                                                                                                           | DOMAIN TRANSFER            |         |
| Package agents and evidence registries for policy, healthcare, supply chains, and biological systems.             |                            |         |

## Conclusion

Interloom proposes a new reasoning surface for complex decisions. Its central move is to make cross-domain causal handoffs explicit, inspectable, and uncertainty-aware. Specialized agents do not remove ambiguity; they organize it. The resulting graph helps people see how an initial decision can travel through systems, where the chain is strong or fragile, and which intervention could change what happens next.

## REFERENCES

# Research foundations

1. IPCC. Climate Change 2022: Impacts, Adaptation and Vulnerability. Technical Summary. Working Group II contribution to the Sixth Assessment Report. <https://www.ipcc.ch/report/ar6/wg2/chapter/technical-summary/>
2. Du, Y., Li, S., Torralba, A., Tenenbaum, J. B., and Mordatch, I. "Improving Factuality and Reasoning in Language Models through Multiagent Debate." arXiv:2305.14325, 2023. <https://arxiv.org/abs/2305.14325>
3. Liang, T. et al. "Encouraging Divergent Thinking in Large Language Models through Multi-Agent Debate." arXiv:2305.19118, 2023. <https://arxiv.org/abs/2305.19118>
4. Barfuss, W., Donges, J. F., Lade, S. J., and Kurths, J. "When optimization for governing human-environment tipping elements is neither sustainable nor safe." Nature Communications 9, 2354, 2018. <https://doi.org/10.1038/s41467-018-04738-z>
5. Staal, A. et al. "Forest-rainfall cascades buffer against drought across the Amazon." Nature Climate Change 8, 539-543, 2018. <https://doi.org/10.1038/s41558-018-0177-y>

**INTERLOOM****SEE WHAT HAPPENS NEXT.**

SYSTEM 003

DECISION SYSTEMS LAB / 2026