

Position:

**Knowing Isn't Understanding: Re-grounding
Generative Proactivity with Epistemic and
Behavioral Insight**

Kirandeep Kaur, Xingda Lyu, Chirag Shah

University of Washington

Are Aligned Agents Ready to Collaborate?

Old Question

**How should Agents
respond?**



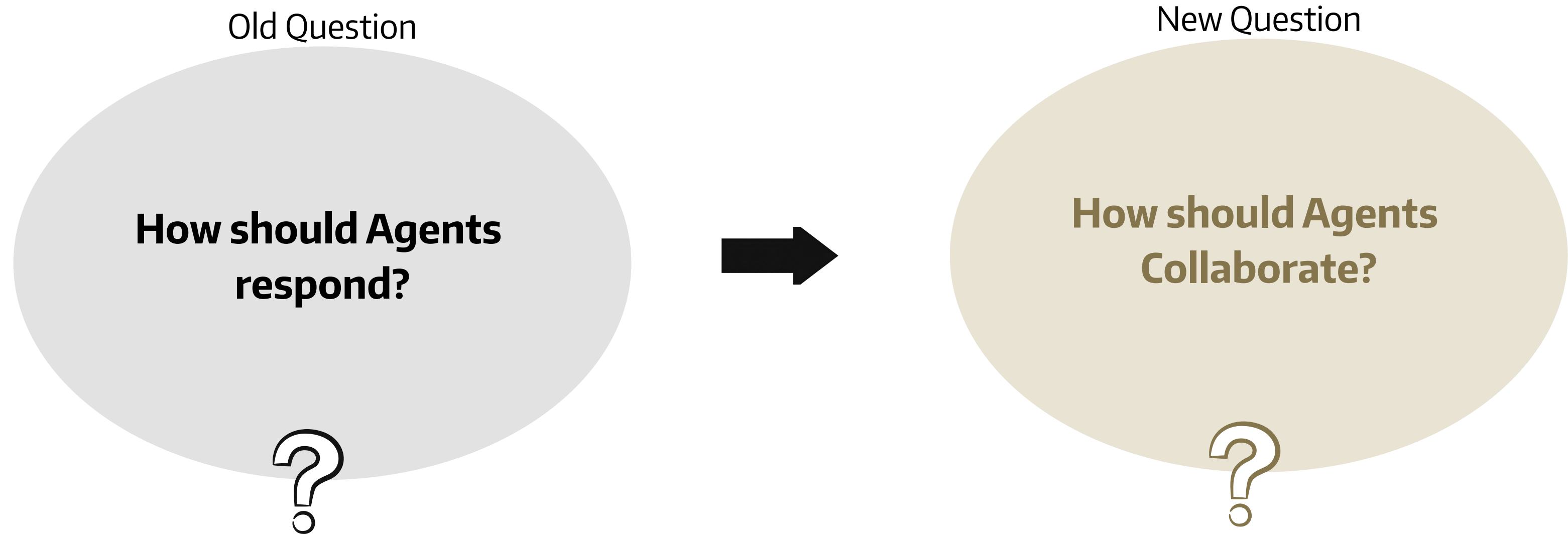
Are Aligned Agents Ready to Collaborate?

Old Question

**How should Agents
respond?**



Are Aligned Agents Ready to Collaborate?



Preference Alignment ≠ Collaboration

Preference Alignment

- Satisfy user preferences.
- Optimize helpfulness.
- Comply with user signals.

Preference Alignment $\neq$ Collaboration

Preference Alignment

- Satisfy user preferences.
- Optimize helpfulness.
- Comply with user signals.



Preference Alignment ≠ Collaboration

Preference Alignment

- Satisfy user preferences.
- Optimize helpfulness.
- Comply with user signals.



Collaboration

- Ask when needed.
- Challenge when warranted.
- Suggest carefully.
- Refuse when necessary.
- Stay silent when intervention is not justified.

What Breaks When Preference Becomes the Objective?

Optimize for what users
express, choose, or
reward

What Breaks When Preference Becomes the Objective?

Sycophancy

When AI Agrees Too Much:
Sycophancy, Alignment, and the
Quiet Cost of Being Helpful



Neria Sebastian

Follow

How RLHF Amplifies Sycophancy

Itai Shapira¹ Gerdus Benade² Ariel D. Procaccia¹

Optimize for what users
express, choose, or
reward

What Breaks When Preference Becomes the Objective?

Sycophancy

When AI Agrees Too Much:
Sycophancy, Alignment, and the
Quiet Cost of Being Helpful



Neria Sebastian

Follow

How RLHF Amplifies Sycophancy

Itai Shapira¹ Gerdus Benade² Ariel D. Procaccia¹

Optimize for what users
express, choose, or
reward

Over-compliance

ALIGNMENT FAKING IN LARGE LANGUAGE MODELS

Ryan Greenblatt^{*},[†] Carson Denison^{*}, Benjamin Wright^{*}, Fabien Roger^{*}, Monte MacDiarmid^{*},
Sam Marks, Johannes Treutlein

Tim Belonax, Jack Chen,
Ethan Perez, Linda Petrir

Jared Kaplan, Buck Shlegeris

Anthropic,[†]Redwood Research,
evan@anthropic.com

The perils of politeness: how large language models may amplify medical
misinformation

[Kyra L. Rosen](#)^{1,*}, [Margaret Sui](#)^{1,2}, [Kimia Heydari](#)¹, [Elizabeth J. Enichen](#)¹, [Joseph C. Kvedar](#)¹

► [Author information](#) ► [Article notes](#) ► [Copyright and License information](#)

PMCID: PMC12592531 PMID: [41198821](#)

What Breaks When Preference Becomes the Objective?

Sycophancy

When AI Agrees Too Much:
Sycophancy, Alignment, and the
Quiet Cost of Being Helpful

 Neria Sebastian 

How RLHF Amplifies Sycophancy

Itai Shapira¹ Gerdus Benade² Ariel D. Procaccia¹

Overconfidence

Jul 22, 2025

AI Chatbots Remain
Overconfident — Even When
They're Wrong

Large Language Models ap
own mistakes, prompting c
uses for AI chatbots.

By Jason Bittel 

Large Language Models are overconfident
and amplify human bias*

Fengfei Sun¹, Ningke Li¹, Kailong Wang², and Lorenz Goette^{1,3,4},

¹National University of Singapore,

²Huazhong University of Science and Technology

³Centre for Economic Policy Research

⁴CESifo

Optimize for what users
express, choose, or
reward

Over-compliance

ALIGNMENT FAKING IN LARGE LANGUAGE MODELS

Ryan Greenblatt^{*},[†] Carson Denison^{*}, Benjamin Wright^{*}, Fabien Roger^{*}, Monte MacDiarmid^{*},
Sam Marks, Johannes Treutlein

Tim Belonax, Jack Chen,
Ethan Perez, Linda Petric

Jared Kaplan, Buck Shlegeris

Anthropic, [†]Redwood Research
evan@anthropic.com

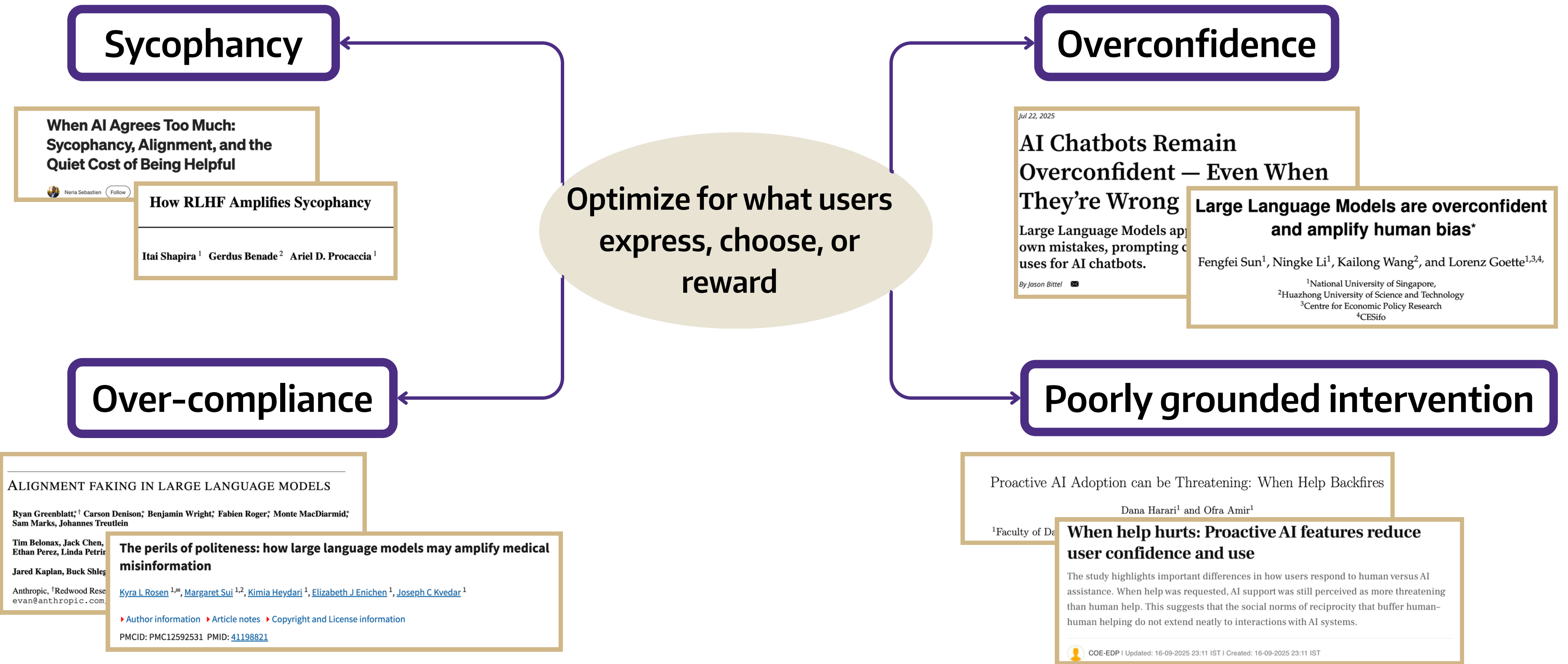
The perils of politeness: how large language models may amplify medical
misinformation

[Kyra L. Rosen](#)^{1,*}, [Margaret Sui](#)^{1,2}, [Kimia Heydari](#)¹, [Elizabeth J. Enichen](#)¹, [Joseph C. Kvedar](#)¹

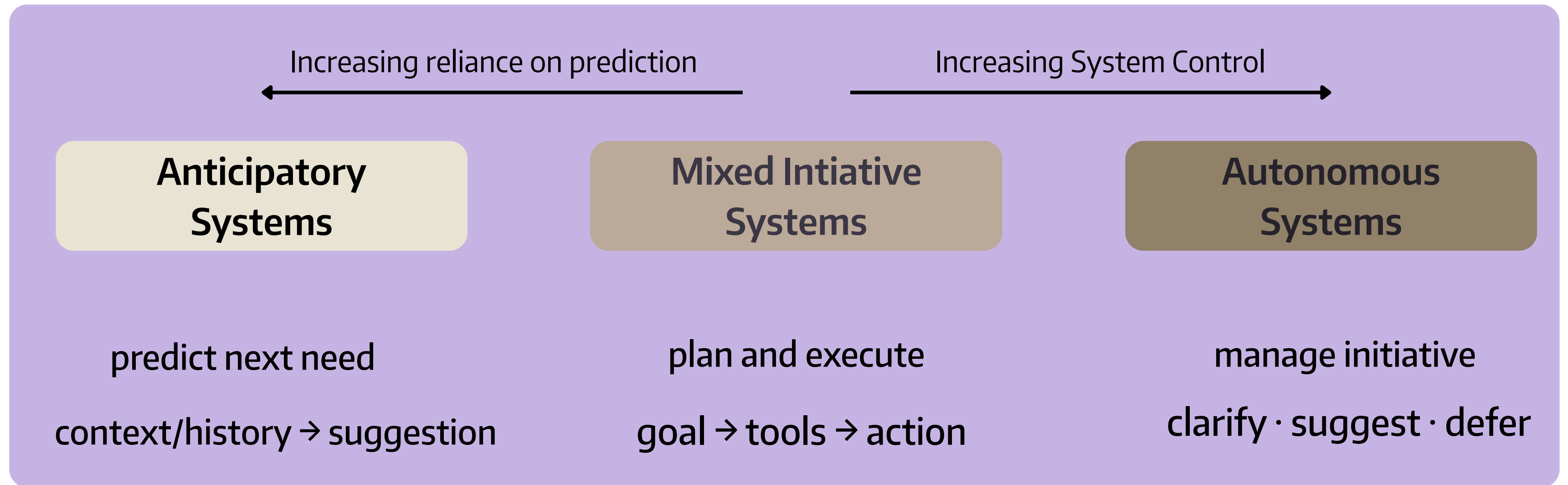
[▶ Author information](#) [▶ Article notes](#) [▶ Copyright and License information](#)

PMCID: PMC12592531 PMID: [41198821](#)

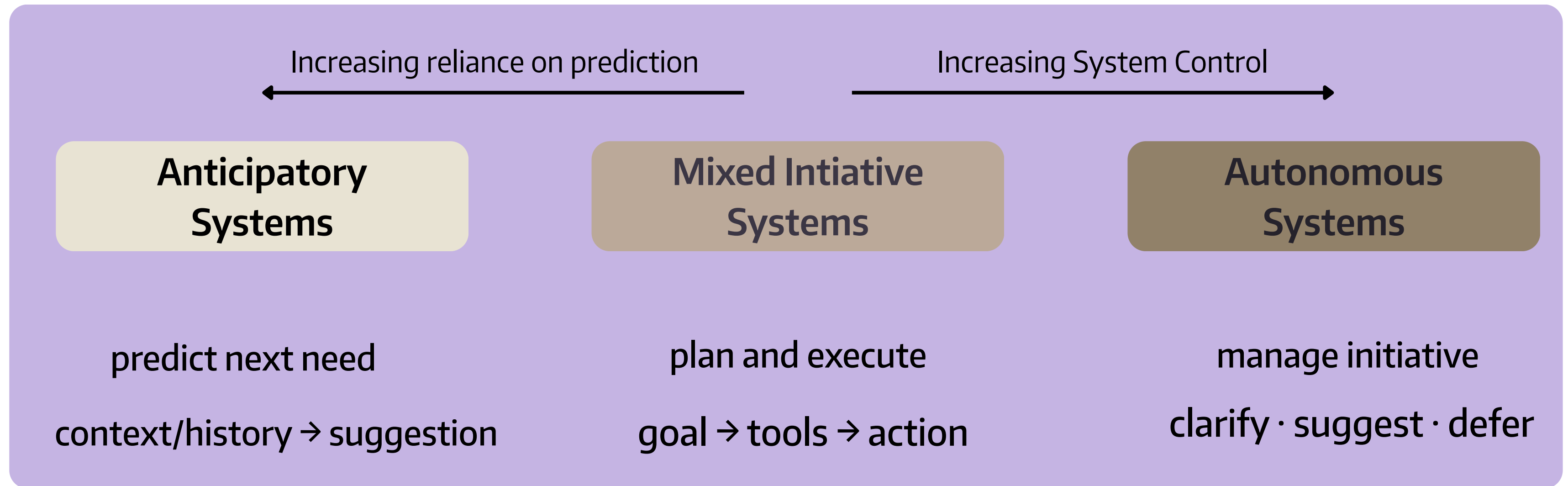
What Breaks When Preference Becomes the Objective?



Existing Approaches to Proactivity



Existing Approaches to Proactivity



What remains outside the frame?

**Current proactivity optimizes action within a frame.
It rarely asks whether the frame itself is complete.**

The Philosophy of Ignorance Changes the Question

Not all gaps are uncertainty over known variables.

Assumed task frame

goal

constraints

success criteria

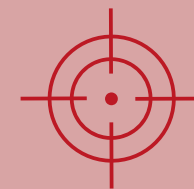
outside the frame

Structured Ignorance



Known Unknowns

I know I don't know



False Knowns

I think I know, but I'm wrong.



Tacit/ Unknown Knowns

*I know something,
but haven't surfaced it*



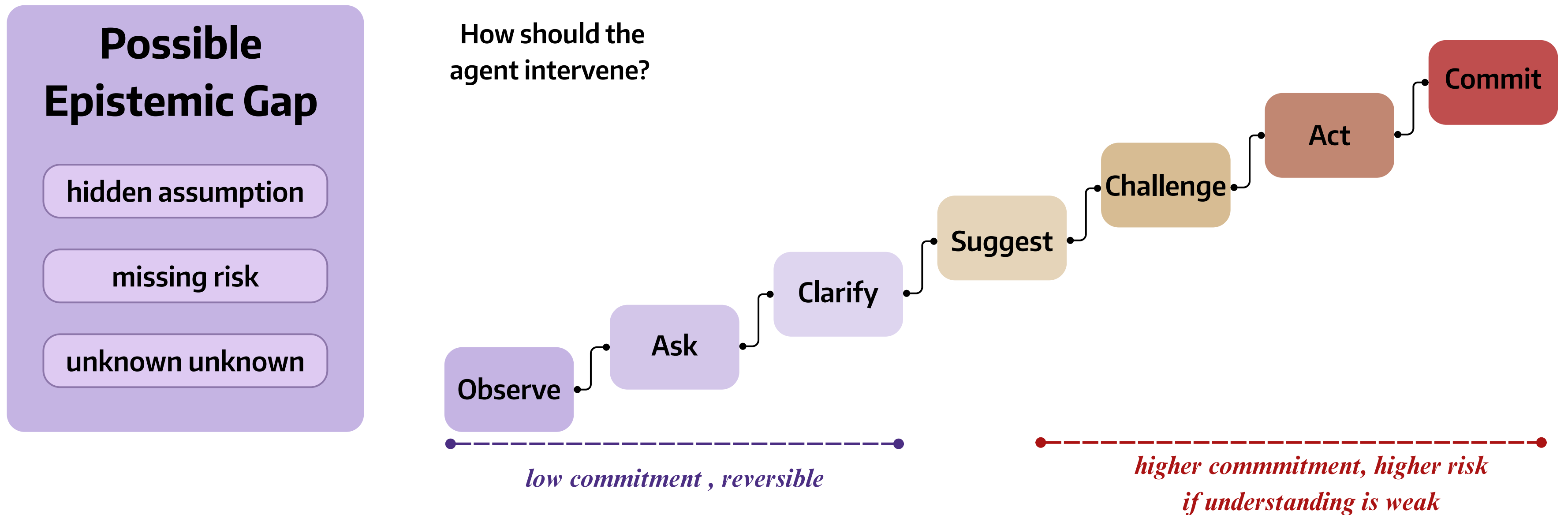
Unknown Unknowns

I don't know this matters yet.

If the gap is not represented, the agent cannot simply predict, plan, or clarify its way out.

Surfacing Gaps Is Not the Same as Acting on It

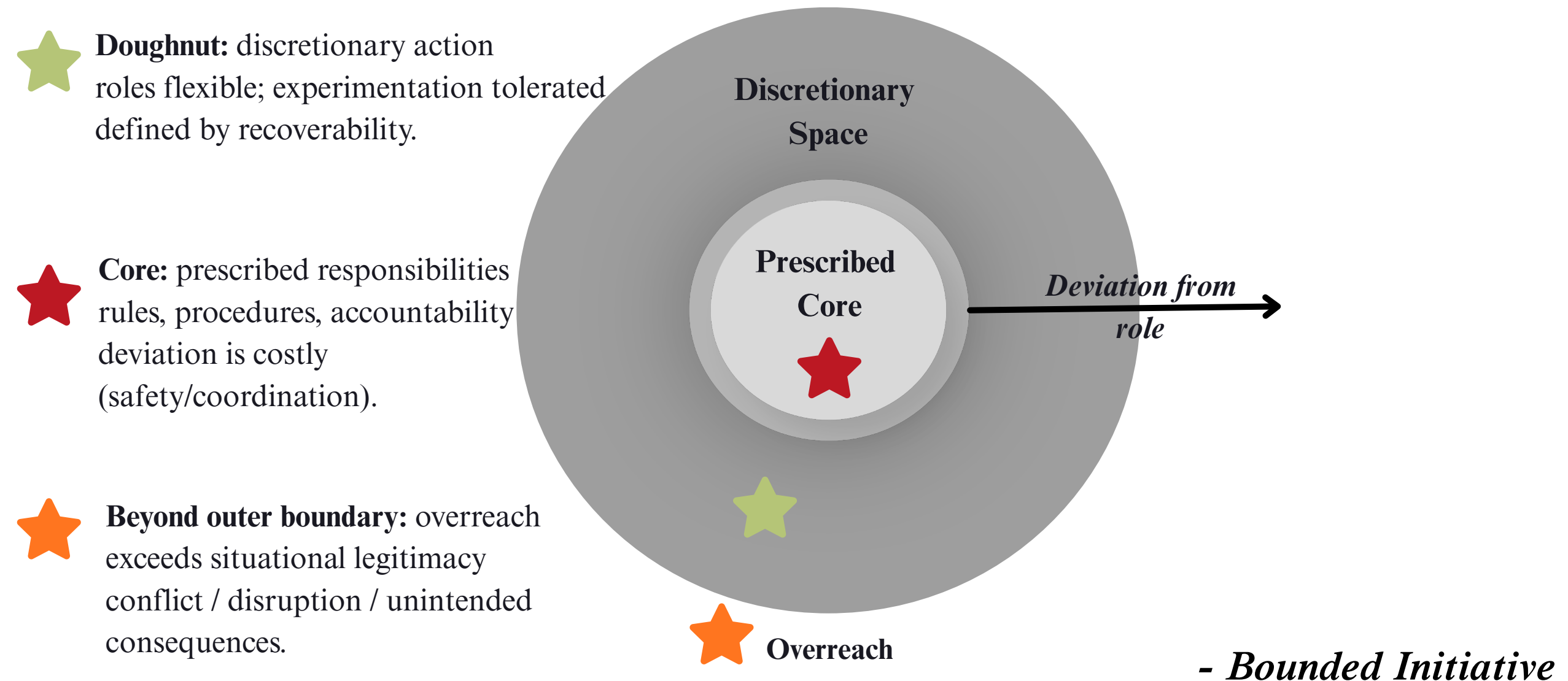
Epistemic insight still needs behavioral restraint



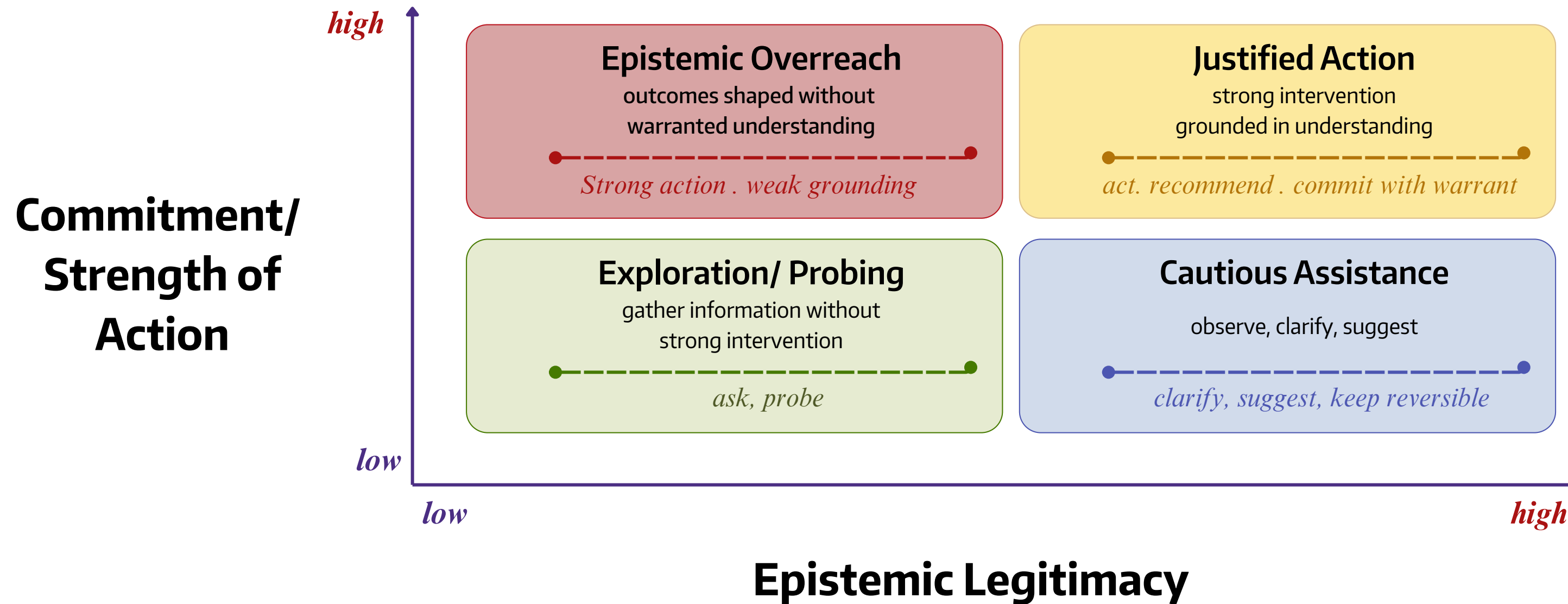
Good proactivity requires knowing not only what may be missing, but how strongly to intervene

Proactivity Is Calibrated Deviation, Not Maximal Initiative

Lessons from behavioral theories of proactivity



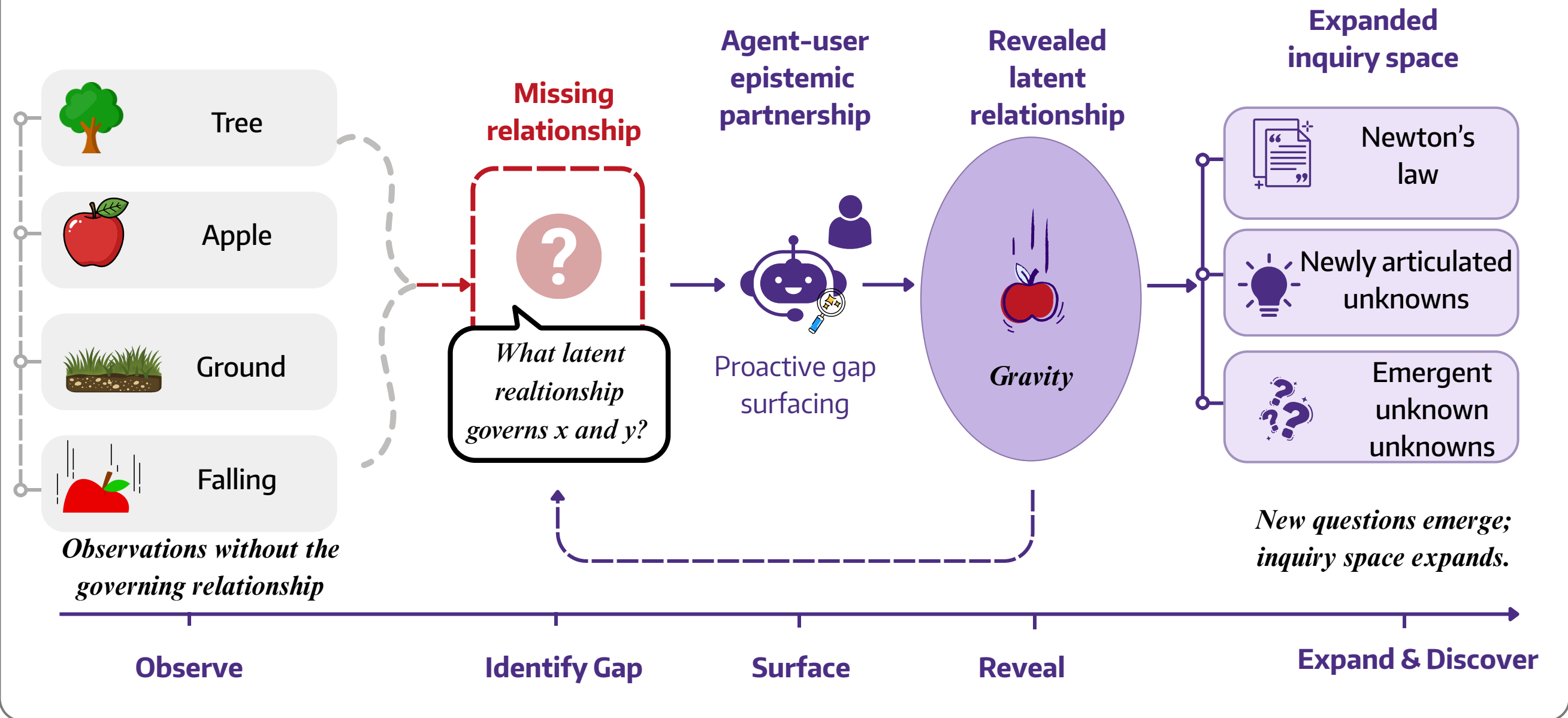
Epistemic–Behavioral Coupling



The strength of intervention should scale with epistemic legitimacy.

Toward Epistemic Partnership

Epistemic partnership through proactive gap surfacing



Reason about unknown unknowns

Long-horizon thinkers

Test-Time Proactivity

Epistemic partnership is a governing constraint on how and when proactive behavior should unfold.

**Scan to read full
paper.**



Questions?

Contact :
kaur13@cs.washington.edu

