

VOLUME VI · THE BEDROCK CANON

THE BEDROCK RESEARCH AGENDA

What we do not yet know

FIRST EDITION · JULY 2026

ON THIS VOLUME

The first five volumes documented stabilized understanding — what the discipline has learned well enough to write down and to act on. This volume documents the opposite: the frontier beyond that understanding, and the questions the discipline has become capable of asking but not yet of answering.

It is not a roadmap, a backlog, or a forecast. It is a research agenda, and it is written in a different register from the volumes before it. It does not teach. It asks. Where the earlier books stated what is known, this one states, as plainly as it can, what is not — and resists, throughout, the temptation to resolve an open question into a comfortable answer.

Several of these questions may occupy years. Some may prove unanswerable. A few, once answered, may overturn doctrine that earlier volumes stated with confidence. None of that is a defect. It is the condition of a discipline that is still alive.

CONTENTS

—	Preface <i>The frontier</i>	04
01	Organizational Intelligence <i>What is the asset we named but do not yet understand?</i>	06
02	Enterprise Cognition <i>Is the model of six faculties the structure, or the simplest picture that fit?</i>	08
03	Evaluation <i>Can we say, and prove, what "better" means?</i>	10
04	Platform <i>Which of our abstractions are permanent, and which merely feel so?</i>	13
05	The Discipline <i>Can a discipline correctly study itself?</i>	15
06	Open Problems <i>The largest unknowns, left open</i>	17
—	Closing & Colophon <i>Why the Canon must never be complete</i>	19

The frontier

A discipline is defined not only by what it knows but by the questions it has become capable of asking. Early in the life of a field, its deepest questions cannot even be formulated; the concepts required to state them do not yet exist. Progress is, in part, the slow acquisition of better questions. The five volumes that precede this one were the acquisition of enough understanding to see, at last and clearly, where that understanding ends.

This volume documents that edge. Behind it lies doctrine — the settled territory the Canon has mapped. Ahead lies the unknown. Between them is the frontier: the band of questions the discipline can now pose but not yet resolve. The purpose of research is not to confirm doctrine. It is to refine it, to challenge it, and — occasionally, and most valuably — to overturn it. So the chapters here do not build arguments toward conclusions. They locate questions and leave them standing.

We have tried to obey one rule above all others: not to pretend certainty where there is none. It is easy, having written five volumes of doctrine, to answer a hard question out of habit. We have instead let the questions remain open, because a question closed prematurely is worse than one left honestly unanswered — the first hides the frontier, the second marks it.

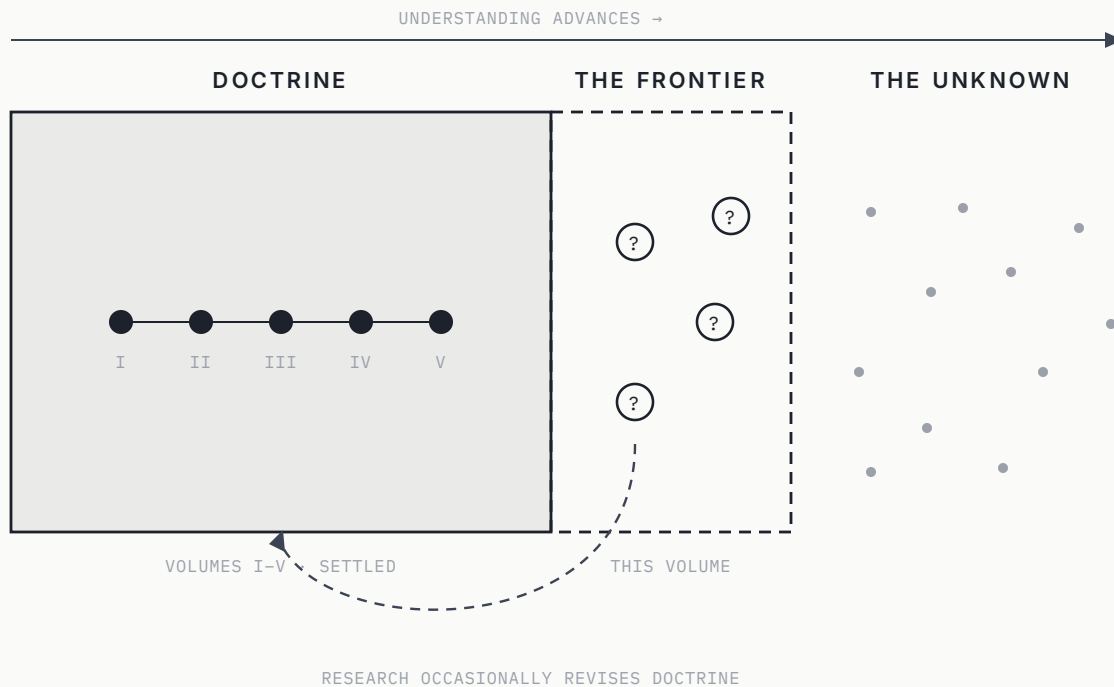
What follows, then, is a map of ignorance, drawn as carefully as the earlier volumes drew their map of knowledge. Doctrine ends here. Research begins.

THE REGISTER

This volume asks; it does not teach. Its confidence is in the quality of its questions, not in answers.

ASSUMED

Volumes I–V. Their doctrine is the settled ground; this volume stands at its edge.



THE BOUNDARY IS WHERE KNOWLEDGE ENDS — AND WHERE RESEARCH BEGINS.

Figure 01 — The frontier. The five volumes are the settled region. The frontier is the narrow band where questions can be posed but not answered; beyond it is the unknown, unbounded. Understanding advances rightward — and research, occasionally, sends a finding backward to revise what was thought settled.

01

CHAPTER ONE

Organizational Intelligence

The first volume named organizational intelligence and argued it is an asset. Naming a thing is not understanding it. Having worked with it across many deployments, the discipline can now ask questions about it that it could not before have formulated — and cannot yet answer.

Can it be represented formally, or only enumerated?
What is its smallest unit — and is there one?
By what law does it accumulate, and decay?
Is it a measurable quantity, or only an ordering?



We named organizational intelligence and built a company around capturing it. We do not, in any deep sense, understand it. Naming located the phenomenon; it did not explain it. The questions below are the ones the naming made possible, and each has proved, on inspection, to be genuinely open.

Can organizational intelligence be represented formally? We capture it operationally — as policies, traces, ontologies — but operational capture is not representation. It remains an open question whether there exists a formal object, some structure or calculus, of which a given organization's intelligence is an instance, or whether it is irreducibly particular and can be held only by enumeration. The answer would decide whether the discipline can ever possess a theory, or only ever a practice.

What is its smallest unit? We took the decision as the unit of observation, and have quietly assumed it is also the unit of intelligence. There is no reason the two must coincide. It may be that judgment decomposes into some smaller, reusable atom below the decision; it may be that it does not quantize at all, and that any unit we impose is an artifact of our instruments rather than a feature of the thing.

By what law does it accumulate and decay? We observe that it appreciates with use and evaporates with attrition. We have no account of the rate of either — no half-life, no conditions under which decay halts, no way to say whether some captured intelligence becomes permanent or whether all of it is perishable given a long enough silence. Until there is such a law, we are describing weather, not physics.

OPEN QUESTIONS

- 1.1 Can it be represented formally, or only enumerated?
- 1.2 What is its smallest unit — and does it quantize at all?
- 1.3 By what law does it accumulate?
- 1.4 By what law does it decay, and can decay halt?
- 1.5 Is it a measurable quantity, or only an ordering without a scale?
- 1.6 When, exactly, does an observation become an instance of it worth keeping?

02

CHAPTER TWO

Enterprise Cognition

The first volume modeled six faculties; the fifth built a platform in their shape. The model is useful, which is not the same as complete, or correct. Working inside it has exposed questions about the model itself that we are not yet able to close.

What faculties are missing because we have not needed them?

Is the machine-and-person partition principled, or a heuristic?

Do organizations reason collectively, or only coordinate?

Is the single loop the structure, or the simplest picture?



Six faculties were enough to build a platform that works. Enough to build is a low standard of truth. A model that has never been shown to be incomplete should be regarded with suspicion, not comfort, and the questions here are attempts to find where ours breaks.

What faculties are missing? We named sensing, interpretation, decision, action, evaluation, and memory because those were the ones the work forced us to name. There may be others we have not needed yet, or cannot see because we lack the concept — anticipation, attention as a distinct faculty, forgetting treated as an active capability rather than a failure. We do not know how to enumerate the faculties of a cognition from first principles, so we cannot say whether six is a discovery or an accident.

Is the partition between machine and person principled? We sort decisions by frequency and consequence, but that is an engineering heuristic, not a law. It remains unknown whether there is a principled boundary — a property of a decision that determines, rather than merely suggests, which side it belongs on — and whether that boundary moves as capability grows or whether some part of it is permanent because it concerns accountability, which is a property of persons and not of competence.

Can organizations reason collectively? We model cognition as a loop around a shared memory. Reality may be less orderly: many loops, nested and interfering, with a collective reasoning that is more than the sum of its participants. We do not know whether such collective reasoning is a real phenomenon with its own structure, or a folk description of coordination — and the platform we have built quietly assumes the latter.

OPEN QUESTIONS

- 2.1 What faculties are missing from the model?
- 2.2 Is there a principled machine/person boundary?
- 2.3 Which of that boundary is permanent, and which moves with capability?
- 2.4 Do organizations reason collectively, with their own structure?
- 2.5 What new faculties does abundant machine attention create?
- 2.6 Is the single loop the topology, or only the simplest that fit?

03

CHAPTER THREE

Evaluation

The fifth volume made evaluation the mechanism that keeps the platform scientific. But evaluation rests on being able to say whether a decision, or an operation, was good — and that is far less settled than the confidence of the earlier volumes implied.

Can decision quality be separated from outcome and luck?
How do we measure learning, as opposed to mere change?
Can we verify a deployment's deepest claim at all?
How is a surprise a system cannot notice ever measured?



Evaluation is the mechanism we trust to keep the discipline honest. It rests on a capacity we have not actually established: the capacity to say, and to prove, what "better" means. The confidence of the fifth volume outran the evidence, and these are the questions that remain.

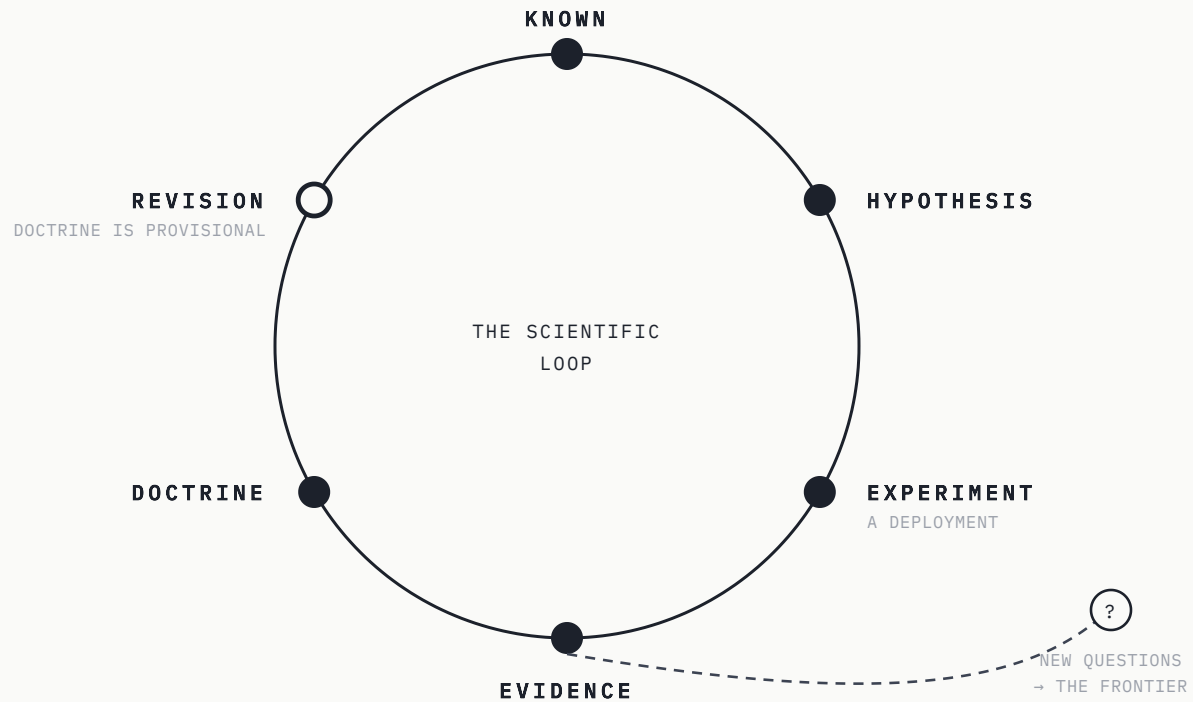
Can the quality of a decision be separated from its outcome? An operation can decide well and fare badly through luck, or decide badly and prosper. We evaluate decisions against intent, but intent is itself contested and shifts over time, and at the scale of a whole operation the separation of good judgment from good fortune is, at present, unsolved. Without it, we cannot be sure we are rewarding the reasoning rather than the result.

How do we measure learning, rather than change? An operation that has changed is not necessarily one that has learned; it may only have drifted. We use the surprise rate as a proxy, but a proxy is not a measure, and we cannot yet distinguish genuine learning from adaptation that carries no understanding — the organization that does the right thing for a reason that will not generalize.

Can a deployment's deepest claim be verified? We judge success by operational metrics and by the presence of a running learning loop. But the real claim — that the organization is now permanently more able to learn — requires, to verify, a timescale longer than any deployment we have yet observed. We are asserting a durability we have not lived long enough to confirm.

OPEN QUESTIONS

- 3.1 Separate decision quality from outcome and luck?
- 3.2 Measure learning, distinct from change or drift?
- 3.3 Verify the claim of durable, learned capability?
- 3.4 Measure the surprises a system is not equipped to notice?
- 3.5 Benchmark intelligence across organizations, or is that category error?



EVERY TURN PRODUCES DOCTRINE — AND, IF THE DISCIPLINE IS HEALTHY, MORE QUESTIONS THAN IT ANSWERS.

Figure 02 — The research loop. A known claim becomes a hypothesis, tested by a deployment, yielding evidence, which may become doctrine — and which, at the revision step, may instead overturn doctrine already held. The loop does not only close. It also spawns new questions, and returns them to the frontier.

04

CHAPTER FOUR

Platform

The fifth volume argued the platform is a temporary expression of permanent ideas, and that modularity draws the boundary between the two. Where that boundary actually falls, we have decided by judgment, not established by knowledge.

Which abstractions truly survive model change?

What belongs in memory, and what in the ontology?

Do reasoning architectures generalize, as ontologies do?

Is any architecture invariant under changes in intelligence itself?



We designed the platform to outlive its parts, and drew a boundary between what is permanent and what is replaceable. We drew it by judgment. Whether we drew it correctly is unknown, and will remain so until enough has changed beneath us to test it.

Which abstractions actually survive a change of model? We built for replaceable engines, but we have not lived through enough model generations to know which of our interfaces are genuinely stable and which merely feel stable because the technology has not yet stressed them. Some abstraction we now treat as permanent may turn out to be an artifact of today's models — visible as such only once they are gone.

What belongs in memory, and what in the ontology? We separate the record of what happened from the language for describing it. The boundary is not principled: some knowledge seems to live in both, some in neither. We do not know whether there is a correct partition between the two, or whether it is a convention that will keep shifting as we understand each better.

Do reasoning architectures generalize the way ontologies do? An ontology is a fresh instance of a shared schema, and that is why learning transfers across industries. It is an open question whether reasoning is the same — whether there exist reusable structures of reasoning, portable across domains — or whether reasoning is irreducibly coupled to its domain and its model, so that nothing about it carries from one deployment to the next.

OPEN QUESTIONS

- 4.1 Which abstractions survive model change, and which only seem to?
- 4.2 Is there a correct partition of memory and ontology?
- 4.3 Do reasoning structures generalize across domains?
- 4.4 What is worth its modularity cost, and what is premature abstraction?
- 4.5 Is any architecture invariant under a change in the nature of intelligence?

05

CHAPTER FIVE

The Discipline

The third volume described how the company turns observation into doctrine. The process is a set of conventions that have worked so far. Whether they are correct — whether they produce true doctrine rather than merely consistent doctrine — is a question the discipline must ask about itself.

Which regularities are lawlike, and which coincidences?

How is disagreement preserved without paralysis?

At what tempo should doctrine revise itself?

Can the discipline study itself without corrupting the study?



The rule of three, the adversarial review, the earned abstraction — these are the conventions by which the discipline makes doctrine. They have held so far. That is not the same as knowing they are right, and a discipline that will not interrogate its own method is not entitled to trust its own conclusions.

Which observations deserve doctrine? Our tests for promotion are heuristics. We have no theory of which regularities are lawlike and which are coincidences that happen to recur three times, and a discipline that cannot tell these apart will, given enough deployments, enshrine a coincidence with full ceremony. How much evidence is enough almost certainly depends on the cost of being wrong, in a way we have not calibrated.

How should disagreement be preserved? Doctrine converges, and convergence is useful. But premature convergence discards the minority reading that a later deployment vindicates. We do not know how a discipline should hold its live disagreements — how to keep dissent available and legible without descending into paralysis, or how long a refuted-seeming view should be kept on the shelf rather than thrown away.

Can the discipline study itself? The third volume made the company both the observer and the observed. Whether an organization can evaluate its own operating model without the evaluation being quietly captured by its own interests is an old and unsolved problem, now turned inward. We may be the least reliable observers of ourselves, and we do not yet know how to correct for it.

OPEN QUESTIONS

- 5.1 Which regularities are lawlike, not coincidental?
- 5.2 How much evidence is enough, given the stakes?
- 5.3 What must never become doctrine, and how is it recognized early?
- 5.4 How is disagreement preserved without paralysis?
- 5.5 At what tempo should doctrine evolve?
- 5.6 Can the discipline observe itself without bias?

06

CHAPTER SIX

Open Problems

The preceding chapters asked questions inside the discipline's established categories. These may not fit the categories at all — the largest unknowns, stated plainly, with no pretense of a path to an answer. Some may be unanswerable. All are, at present, genuinely open.

Six questions, stated and left standing



Can an organization become self-improving?

Able to redesign its own operating model without outside intervention — and if it can, is that desirable, or does it dissolve the human ends the discipline was built to protect?

Can organizational intelligence be transferred?

Moved from one organization to another — or is it so bound to a particular history that transfer is impossible in principle, and not merely difficult in practice?

Can deployment itself be automated?

The Method run by the platform rather than by people — or is there something in observation that requires a human observer, permanently and irreducibly?

Can doctrine emerge algorithmically?

The passage from evidence to abstraction mechanized — or is the judgment of what generalizes, and what is mere coincidence, irreducibly human?

How should organizational memory forget?

Perfect memory may be as pathological as perfect forgetting. There is presumably a discipline of deliberate forgetting, and we have not begun to build it.

What is the limit of augmentation?

We do not know whether there is a ceiling beyond which more intelligence in the loop yields nothing — or degrades the whole — nor, if there is, where it lies.

LEFT OPEN

We state these not because we can see their answers, but because a discipline is obliged to name what it cannot yet answer.

THE WORK AHEAD

Naming these is the work of this volume. Answering them is the work of volumes that do not yet exist.

Why the Canon must never be complete

The previous volumes define the discipline. This one protects it. A discipline without unanswered questions is not complete; it is an ideology — a body of answers to be defended rather than a body of questions to be pursued. The difference between the two is the willingness to keep asking.

So research exists not to complete the Canon but to prevent the Canon from ever being complete. If the discipline is healthy, every volume that stabilizes understanding also exposes deeper unknowns, and the Canon grows at both ends at once — doctrine and frontier together. The day the frontier stops moving is the day the discipline stops being one.

The credibility of everything in the first five volumes rests, in the end, on this sixth: on the discipline knowing where its knowledge ends, and being willing to say so plainly. A field that cannot point to its own frontier is not yet a field. These questions are left open on purpose. We expect to answer some, to fail at others, and to be overturned on a few we thought were settled. We would not have written them down if we were sure. That is the point.



THE ONE IDEA

A discipline stays alive only while it keeps asking. Research is what keeps the Canon from closing.

THE BEDROCK CANON

I · Manifesto — why. II · Method — how. III · Operating System — the practitioner. IV · Industry Playbooks — the reach. V · Platform — the machinery. VI · Research Agenda — the frontier. The first five map what is known; this one maps where the knowing ends.

COLOPHON

The Bedrock Research Agenda, first edition, July 2026. Set in Source Serif 4, Inter, and IBM Plex Mono. The mark is the diagram of enterprise cognition. This edition is expected to be revised most of all — by the answers it does not yet contain.



BEDROCK

VOLUME VI • THE BEDROCK CANON

Doctrine ends here. Research begins.

© 2026 BEDROCK • FIRST EDITION