

The Portfolio Learning Rate

Arc 2 · vocabulary · Fund

In the room: Head of value creation, operating partners as a group, CIO of the firm

The question: Is engagement N+1 cheaper, faster or better because engagements one through N happened?

THE POSITION

SIGHTING

A fund has run AI operational work inside nine portfolio companies. Ask which of the nine worked and why, in comparable form, and nobody can answer — not because the work was bad, but because no two engagements produced the same kind of object.

THE CONSENSUS READ

Bandwidth. More operating partners, more coverage.

THE MECHANISM

Bandwidth isn't the constraint. Nine engagements produced nine one-offs, so engagement ten starts at the same cost as engagement one. The fund has paid for nine instances of learning and banked none of it, because learning at fund level requires comparability and nothing enforced it.

THE EXPOSURE

The thing you're building is a services team, not an asset.

THE TEST

Ask two operating partners for the decision output of their last engagement. If the two documents aren't the same shape, your learning rate is zero.

Portfolio learning rate

Whether engagement N+1 is cheaper, faster or better because engagements one through N happened.

The test. A comparison, not a survey. Put the decision outputs of two engagements side by side. If they share a decision type, an artifact type and a failure vocabulary, the rate is positive. If they share a font, it's zero — and zero is the default, because nothing in a normal operating model produces comparability.

THE BRIDGE — CLAIMED TO REALISED

The nine-engagement fund is a composite. Ostrom and the agent-market coordination study are real and cited with their samples; the per-engagement value is our estimate.

LINE	MOVES	RUNNING
Claimed — A value-creation function whose capability compounds across the portfolio		9 engagements, 0 comparable pairs
1. What compounding would require [changes the story]	-\$0	9 engagements, 0 comparable pairs
Three things have to be common across engagements before any comparison is possible.		
2. Why the internal team can't fix this by trying harder [changes the story]	-\$0	9 engagements, 0 comparable pairs
Operating partners are typically ex-operators with deep judgment and their own methods.		
3. Where I was wrong — I proposed a central playbook [where I was wrong]	-\$0	9 engagements, 0 comparable pairs
Given method heterogeneity, my recommendation was the obvious one: write a central methodology, mandate it across the operating team, enforce it through the... Cost: I proposed the thing that produces compliance instead of the thing that produces comparability.		
4. What a positive rate is worth	+\$180,000 to +\$400,000 per subsequent engagement	\$180,000 to \$400,000 per subsequent engagement, before exit-narrative value
With comparable outputs, three things become available that currently aren't.		
Realised		A learning rate of zero converts a fund-level asset into a services cost line

THE DOWNSIDE

Multiple: 9x, illustrative — though the direct cost here is a spend line rather than a multiple

\$180,000–\$400,000 of avoidable diagnostic cost across, say, six further engagements over a hold cycle. The multiple is illustrative and the arithmetic is exposed; substitute your own.

\$1.1M to \$2.4M of value-creation spend buying capability the fund already paid for once — plus an unpriceable exposure: a fund with nine AI engagements and no comparable record cannot answer an LP asking which operational levers actually work — and cannot answer it for itself when allocating the next round

When it surfaces. During fundraising, when an operating-model claim meets a diligence question — and internally every time an operating partner starts an engagement from zero without knowing a colleague solved the same structural problem two years ago.

The Portfolio Learning Rate

The question: Is engagement N+1 cheaper, faster or better because engagements one through N happened?
Paste into a request list or a management agenda. Each question resolves to an artifact, not to a characterisation.

THE DILIGENCE PACK

1. Place the final decision output of the two most recent AI or automation engagements side by side. State whether they share a decision type, an artifact shape and a failure vocabulary.
Artifact: The two documents.
Twenty minutes, and it is the whole diagnosis.
2. For each of the last nine engagements, state in one sentence what was funded, what was killed, and what would have reversed the call.
Artifact: Nine sentences.
The ones that can't be written are engagements that terminated nothing, per Episode 08.
3. Ask two operating partners, independently, to name the most common failure mode they've seen across their portfolio companies.
Artifact: Two answers.
If they describe the same phenomenon in different words you have a vocabulary problem, not a knowledge problem — fixable in one session.
4. Identify who at fund level is accountable for the reusability of engagement output.
Artifact: A name, or its absence.
Absence is the finding, and it is the norm.
5. Convene the operating partners to author — not receive — the failure vocabulary and the fixed artifact shape.
Artifact: A one-page schema with their names on it.
Per Line 3, authorship is the mechanism. A document they wrote survives; a document they were sent does not.

DISQUALIFIER

If the response to item 5 is to commission a methodology from an external firm, the mechanism has been inverted and the output will be complied with rather than used.

SOURCES — EVERY MEASUREMENT WITH ITS SAMPLE

CLAIM	SOURCE	SAMPLE / CLASS
Design principles present in durable common-pool institutions and absent from most failed ones; collective-choice arrangements and user-accountable monitoring among them	Ostrom, Governing the Commons, 1990, and the design-principles literature	Decades of documented field cases, some institutions durable for centuries (measured)
Mediation with participant-authored rules won at 1556 cumulative honest-agent utility vs 1209 do-nothing baseline; top-down governance third at 1288; binding contracts fifth at 1130, below doing nothing	Agent-market coordination study, eight mechanisms	8 conditions x 200 rounds — note n=1 per condition; treat ordering as indicative, not the margins (measured)
\$180,000–\$400,000 of avoided diagnostic cost per subsequent engagement	Our estimate, exposed so it can be substituted	n/a — illustrative (illustrative)
A fund with nine AI engagements and no comparable pairs	Composite	n/a — composite (composite)

One standing caveat. Every number in this show is somebody else's measurement, and I'll tell you whose, with the sample. None of it is diligence on your deal. Do that yourself.