

Attribution Survivorship

Arc 2 · vocabulary · Fund

In the room: Investment committee, head of value creation, anyone setting AI policy across a portfolio

The question: Is our read on what AI does for a portfolio company coming from a biased sample?

THE POSITION

SIGHTING

Ask a fund how AI is performing across its portfolio and the answer is almost always one word: uneven. Some companies show real EBITDA impact, most don't, and nobody says why.

THE CONSENSUS READ

Variance in execution. Better vendors, better change management, better operating partners.

THE MECHANISM

Uneven isn't a distribution of outcomes. It's a distribution of measurability. Attribution needs three things — a pre-change baseline, a defensible counterfactual, and cheaply separable output quality — and only small, low-value use cases have all three. So the initiatives that show attributable impact are the ones where attribution was possible, not the ones that worked best.

THE EXPOSURE

Your view of what works in AI is being formed by a biased sample, and the bias runs toward the least valuable work.

THE TEST

List the initiatives with proven impact and check whether they're the largest ones. If they are, ignore me.

Attribution survivorship

The initiatives visible in your results are the ones that could be measured, not the ones that worked.

The test. A correlation you can run in an afternoon. Plot each initiative's demonstrated impact against its underwritten value. If the proven wins cluster at the low-value end, you're reading a measurability distribution and calling it a performance distribution.

THE BRIDGE — CLAIMED TO REALISED

The portfolio pattern is a composite. Flyvbjerg's database figures and the Diagon correlation are real and cited with their samples; the market-context claim is internal research and flagged as such.

LINE	MOVES	RUNNING
Claimed — A fund-level read informing the next allocation of value-creation capital		AI delivers real but uneven EBITDA impact
1. The three conditions for attribution [changes the story]	-\$0	attribution available on the cheap end, unavailable on the expensive end
2. The bias that produces [changes the story]	-\$0	allocation decisions being made on a biased sample
3. The falsification test, borrowed [changes the story]	-\$0	allocation decisions being made on a biased sample, now testable
4. Where I was wrong — I read it as a capability distribution [where I was wrong]	-\$0	unchanged, priority raised
5. What the bias costs [residual]	+\$0	a portfolio that optimises for provable impact rather than for impact
Realised		A portfolio that optimises for provable impact rather than for impact, one reasonable decision at a time

THE DOWNSIDE

Multiple: None stated — and the omission is the point. Every other episode in this set uses an illustrative 9x; this one does not, because the exposure is unobservable by construction.

There isn't a clean dollar figure and manufacturing one would breach this show's own evidence rule. The exposure is the difference between the operational alpha you pursued and the one you'd have pursued with unbiased information, which cannot be measured from inside the biased sample. The illustrative multiple used elsewhere in this set would add false precision here.

Unquantifiable by construction — stated as such rather than estimated

When it surfaces. It doesn't. This is the only episode in the set with no discovery event, because the failure mode is invisibility itself. There is no retrade, no write-down, no buyer's diligence finding — just a fund that spent a hold period getting steadily better at proving small things.

WHY THERE IS NO NUMBER HERE

Every other episode in this set prices its exposure at an illustrative 9x. This one does not, and the absence is the finding rather than a gap in the work. Attribution survivorship is a bias in what you can see, so its cost is the difference between the operational alpha you pursued and the one you would have pursued with unbiased information - a quantity that cannot be measured from inside the biased sample. Any figure offered here would have to be derived from the very initiatives the mechanism says are unrepresentative. Publishing a number would make this artifact more persuasive and less true, and this show's standing caveat is that every number comes with its sample. This one has none, so it gets no number.

Attribution Survivorship

The question: Is our read on what AI does for a portfolio company coming from a biased sample?
Paste into a request list or a management agenda. Each question resolves to an artifact, not to a characterisation.

THE DILIGENCE PACK

1. List every AI or automation initiative across the portfolio with its underwritten value and its demonstrated impact. Plot one against the other.

Artifact: A two-column table.
If demonstrated impact clusters at the low-value end, the term applies.

2. Count initiatives that materially beat their business case against those that materially missed.

Artifact: Two counts.
Per Line 3 you are testing for symmetry. A one-sided distribution means a measurement floor, not execution variance.

3. For every initiative reported as 'hard to attribute', state which of the three conditions is missing — baseline, counterfactual, or separable quality.

Artifact: One word per initiative.
Converts a vague category into a fixable list, and the answer is usually 'baseline'.

4. Make pre-change baseline capture a release condition for value-creation capital above a threshold, effective immediately.

Artifact: A line in the funding process.
Cheapest item in this pack and the only one that changes what's visible next year.

5. For the three largest initiatives currently running with no baseline, state what can still be reconstructed and what is already gone.

Artifact: An honest inventory.
Some is recoverable from system data; the informal layer is not, and knowing which is which prevents a false sense of coverage later.

DISQUALIFIER

If item 1 shows demonstrated impact spread evenly across the value range, attribution is working in this portfolio and this episode doesn't apply. That would be an unusual and genuinely good result.

SOURCES — EVERY MEASUREMENT WITH ITS SAMPLE

CLAIM	SOURCE	SAMPLE / CLASS
Only ~8.5% of projects met cost and schedule and ~0.5% also delivered promised benefits; cost overruns roughly constant across ninety years; one-sided error distribution implicates structural rather than technical causes	Flyvbjerg, Oxford megaproject database	More than 16,000 projects, 104 countries, six continents, ~90 years (measured)
Quality-to-payment correlation $r = 0.16$ in the mid band; evaluation bottleneck tightens with complexity	Shang, Liu & Jin, When Agent Markets Arrive (Diagon), arXiv 2604.06688v2	1,957 transactions, 3 seeds, 5 model families, simulated market (measured)
Returns now depend far more on operational value creation than financial engineering; real attributable EBITDA impact from AI inside portfolio companies remains uneven across the industry	Internal market research	n/a — not a published measurement; specifics held loosely, direction only (reported)
A nine-portfolio-company fund observing 'uneven' AI impact	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.