

TIMEPOINT

ALPHA

A TIMEPOINT AI WHITEPAPER

Field Notes on Deal Risk

How to stress-test an M&A or investment decision across its branches — the counterparty's view, the failure you haven't named, and the option that survives a bad quarter — before the capital commits.

Deals don't die in the model. They die in the branches.

Every serious acquirer and investor already stress-tests the numbers. Sensitivity tables, downside cases, scenario tabs — the spreadsheet answers *"what if the inputs are different?"* with precision. What it cannot answer is the question the deal actually turns on: **what do the people on the other side of this decision do next?**

The seller's board in week six. The rival bidder who was never going to win but could always make you pay more. The champion inside the target who leaves. The regulator who doesn't block the deal but slows it a quarter. These aren't cells in a model. They're *branches* — distinct ways the situation plays out, each driven by specific people with specific incentives.

These field notes are the method we use at Timepoint to map those branches: five practices, one worked example, and a checklist to run on your own live deal — with or without us.

01	The base case is a branch	Your plan is one future among several — treat it that way
02	Ask the other side	The counterparty's-eye view of your own deal
03	The blind spot is structural	Reason backward from the outcome to find the shared assumption
04	Disagreement is data	When runs diverge, that divergence is the finding
05	Name the actors	Decisions are made by people, not demographics
–	A worked example	A fictional acquisition, run through the method
–	The checklist	Five questions before the capital commits

The base case is a branch.

The most dangerous artifact in any deal room is the plan everyone already believes. Not because it's wrong — because it's *singular*. Once a base case exists, every new fact gets read as evidence for it.

The corrective is cheap: force the plan to stand in a lineup. Generate the distinct ways the deal plays out — the version where integration stalls, the version where the counterparty's best customer walks at signing, the version where it works but a year late — and **rank them, stating in plain words how much you trust each ranking.**

Notice what we didn't say: probabilities. We don't attach percentages to branches, ours or anyone's, because we can't substantiate them — and in our view, neither can anyone else selling foresight today. A number like "72% close probability" isn't a measurement; it's confidence theater, and U.S. regulators have begun treating unsubstantiated AI performance numbers as exactly that. Ordered branches with stated uncertainty tell you the same thing honestly: *which futures deserve your attention, and how sure anyone can honestly be.*

A ranked list you can argue with beats a percentage you can't audit.

Ask the other side.

Every deal model is written from your side of the table. The fastest way to find its soft spots is to simulate the same decision *from the counterparty's chair*: what does the seller's board actually need this quarter? What does the rival bidder's fund cycle force them to do? If you're the buyer, which of your assumptions would the seller be happiest you kept?

In practice this one move produces more negotiating leverage than any other analysis we run. Knowing what the other side can say *yes* to — because you rehearsed their meeting, not just yours — changes what you ask for. And the inversion generalizes: run your competitor's response to the announcement, your lender's response to the bad quarter, your key customer's response to the ownership change.

The blind spot is structural.

Ask a deal team "what could go wrong?" and you'll get a competent list — of the risks they already know. The failure that actually hurts is the one that's *upstream of the list*: the assumption every forward-looking plan in the room quietly shares.

Forward planning can't see it, because every forward plan starts from the same present — and inherits the same assumptions. So run the analysis backward instead. Start from the outcome — the deal that obviously worked, or the write-down everyone swears can't happen — and reason in reverse: what had to be true the year before? The year before that? We call this **portal analysis**, and its signature finding is the *common ancestor*: the single early condition that most failure timelines share.

When three different ways the deal dies all pass through the same gate — the platform migration, the one customer contract, the founder staying eighteen months — that gate is the real diligence item. It was rarely on the original list.

Disagreement is data.

Run the same scenario several times and the runs will not agree. Most analysis treats that as noise to be averaged into one confident answer. We treat it as the most important output there is.

Where repeated runs converge, the branch is stable — its outcome is being driven by structure (incentives, constraints, timing) rather than by luck. Where runs diverge wildly, the situation itself is sensitive: small differences in sequence or personality swing the result. That's not a modeling failure. That's the simulation telling you *which parts of the future are actually undecided* — and therefore where preparation, optionality, and negotiation effort belong.

Averaging away the disagreement between runs destroys exactly the thing a decision-maker needed to see.

Name the actors.

Market research simulates crowds: a thousand statistical people, aggregated into a percentage. Deals are not decided by crowds. A deal is decided by *maybe nine people* — two principals, three board members, a lender, a rival, a regulator, a key employee — each with documented history, stated incentives, and public constraints.

So simulate **those nine people**, by name, with what's on the record about how they've behaved before. A named-actor simulation produces findings you can interrogate ("why does the seller's chair flip in this branch?") and re-run when facts change. An aggregate can only tell you what people-in-general do — and no deal you care about is decided by people-in-general.

At Timepoint the named actors live on a persistent temporal graph, so the world a simulation inherits doesn't reset between questions: the same cast, carried forward, deal after deal.

HOW THIS RUNS IN PRACTICE

Intake is one call. The decision, the actors involved, and the questions you need answered. In a zero-knowledge engagement you hand over nothing else — no files, no data room access, no record you asked. An NDA is available if you want one; the method doesn't require it.

The work takes days, not a consulting quarter. Simulated counterparties, competitor responses, stress tests, and a portal run — grounded in public signal and the named actors around your deal.

The deliverable is a decision brief: branches ranked and argued, uncertainty stated in words, the blind spot named, and the load-bearing assumption to verify in the real world before you commit. Built to be argued with — and re-run when the facts change.

One target, two bidders, nine people.

SIMULATED SCENARIO — FICTIONAL DEMONSTRATION. "Harborlight Group," the target, and every finding below are invented to show the method's shape and rigor (FTC 16 CFR 465). This is an AI simulation of a fictional deal — not market research, not a real transaction, not client traction.

Harborlight (fictional), a lower-middle-market acquirer, is weighing a logistics-software target at a full price. A strategic rival is circling. The intake call surfaced three live options and three questions: *do we pay up now, wait for the process to cool, or walk — and what are we not seeing?*

RANK	BRANCH	WHAT THE RUNS SHOWED	STATED UNCERTAINTY
I	Pay now, structure the risk	Most stable across runs — but only in the variant where a third of the price rides on the target's top-customer renewal. The counterparty run showed the seller's board <i>can</i> accept contingent value this quarter; their fund's clock, not the price, is what they're negotiating.	Stable. Sensitive to one assumption: the renewal decision stays with the customer's current ops chief.
2	Wait for the process to cool	High divergence. In half the runs the rival bidder's interest was theater and waiting saves 15–20% of price; in the other half the rival transacts and the option disappears. The runs disagree because the rival's intent is genuinely unknowable from outside.	Unstable — treat as a bet on the rival's fund cycle, and price that bet, not the target.
3	Walk	Runs converged: walking is safe but not free — the rival owning this target changes Harborlight's next two years of pipeline. "No deal" is also a branch, with consequences worth rehearsing.	Low divergence. Direction agreed; only degree varied.

THE BLIND SPOT THE PORTAL RUN SURFACED

Every forward run — buy, wait, or walk — quietly assumed the target's data-migration project lands on schedule. Reasoning backward from the failure outcomes, that migration was the common ancestor of most bad timelines *in all three strategies*. It appeared nowhere in the original risk list. It became the first diligence call the next morning.

Five questions before the capital commits.

Run these against your live deal — in a workshop with your own team, or as a Timepoint engagement. They are the five field notes, in interrogative form:

- **What has to be true?** List the base case's load-bearing assumptions — then ask which of them every *other* plan in the room also depends on.

- **What does the other side's meeting look like?** Rehearse the counterparty's decision, not just yours. What can they actually say yes to this quarter?

- **What's the common ancestor?** Start from the failure everyone says can't happen and reason backward. Which early gate do most bad timelines share?

- **Which option survives the bad quarter?** Not which option wins the base case — which one degrades gracefully when the first assumption breaks.

- **Where do honest analyses disagree?** Wherever two credible runs of the future diverge, stop averaging. That divergence is where your preparation belongs.

WHAT WE WON'T TELL YOU — READ THIS BEFORE HIRING US OR ANYONE

- ✓ We don't claim to predict the future accurately. No published calibration record exists yet, and we won't pretend one does — ask any vendor quoting an accuracy percentage what the denominator is.
- ✓ We're not the first to simulate populations. We simulate *specific, named* actors rather than aggregate crowds — that specificity is the difference we sell.
- ✓ We run on off-the-shelf frontier models. The value is the combination — the persistent entity graph, the branch discipline, and the honesty policy you're reading.
- ✓ Every simulation is labeled a simulation, everywhere it appears — including the fictional example in these pages. This is strategic advisement, not legal, investment, or financial advice.

ABOUT TIMEPOINT

Synthetic time travel, honestly labeled.

Timepoint AI is a Santa Monica company building simulation-driven research and strategy work: synthetic market research, strategy stress-testing, competitor-perspective analysis, blind-spot discovery, and multiple-choice analysis for teams weighing real, dated decisions. Engagements are founder-led and scoped in a single fifteen-minute call. The company is in alpha, and says so.

BRING US A DECISION

The first call is the whole intake: the decision, the actors, and the questions you need answered. You'll know the scope and timeline before we hang up — and if we can't help, we'll say so on the call.

timepointai.com/client-services · timepointai.com/contact

© 2026 Timepoint AI · timepointai.com · This document contains a clearly-labeled fictional demonstration scenario and no real client information. Every simulated output Timepoint produces is labeled as an AI simulation. Nothing in these pages is a prediction, a guarantee, or legal, investment, or financial advice.