

Relevant and Reliable Information

- Relevant and reliable information is the third link in Decision Quality, after the frame and the alternatives. The job is knowing what bears on the choice, what to trust, and when to stop gathering.
- What counts as relevant depends on where you sit, and relevance is a lower bar than value. Plenty of information bears on a decision without being worth the cost of getting it; the value of information is the sharper screen and tells you when to stop.
- Reliability fails in two places: the estimate, when it is built from the case at hand and ignores the comparable cases, and the process, when load-bearing claims go uncorroborated and unchallenged.
- Five practices keep information fit for the decision: map the few uncertainties that move it, anchor forecasts in base rates, corroborate and challenge critical claims, replace point estimates with calibrated ranges, and convert high-value assumptions into fast tests.
- The discipline has limits. Value of information needs odds you can estimate, reference classes need comparable cases, and calibration needs feedback that a one-off decision rarely supplies. Match the effort to what rides on the choice.
- AI helps with drafting reference classes, summarizing evidence, and mapping uncertainties, but it cannot decide what to trust or how much to stake on it, and it will invent a reference class as readily as retrieve one.

Teams rarely fail for want of information. They fail because they gather the wrong information, trust it further than it deserves, or keep gathering long after the choice is already clear. A roomful of analysis can still rest on two or three numbers nobody pressure-tested.

Relevant and reliable information is the third link in Decision Quality, after the frame is set and the alternatives are on the table.¹ The discipline is knowing what to trust and when to stop gathering. Get this wrong and everything built on top of it is wrong too, however careful the modeling: a sound model running on a soft estimate returns a confident wrong answer. This paper sets out where information quality breaks down and the practices I use to keep it relevant and reliable without turning the decision into a research project.

The theory in brief (why information quality determines strategy outcomes)

Relevance governs what to gather, and what counts as relevant is partly a matter of where you sit. Managers at different levels read the same situation differently, so they treat different uncertainties as the ones that matter; the relevant set is shaped by the frame, and the frame by position.²

Relevance is also a low bar. Plenty of information bears on a decision without being worth the cost of getting it, so the sharper screen is value. Howard's **value of information** makes it explicit: a finding is worth gathering only when it would change what you do, or move the odds enough to matter, once you net out the cost and delay of getting it. When it clears that bar, gather it; when it does not, stop refining that estimate and move to the next thing the decision needs.³

Reliability decides whether to trust what you gather, and it fails in two places. The first is the estimate itself. Asked to forecast, managers reason from the particulars of the case in front of them, the **inside view**, and an estimate built that way runs optimistic. It sets aside the **outside view**, the record of how comparable efforts actually turned out.^{4,5} Pricing the forecast against a **reference class** of similar cases is the correction, and across large programs that has done more for cost and schedule realism than another round of internal analysis.⁶

The second is the process that produces the estimate. Dean and Sharfman found that decisions reached through an information-rich, **procedurally rational** process were more effective than those that were not.^{7,8} In practice, that means keeping **fact, assumption, and opinion** separate so the soft inputs stay visible, corroborating any load-bearing claim against an independent second source, and weighting evidence by how it was produced, since a peer-reviewed finding and a vendor's anecdote do not carry the same weight. It also means actively seeking the evidence that would prove the favored view wrong, since search runs toward confirmation by default.⁹ Done together, these practices make information relevant and reliable: relevant, because it bears on the choice and arrives in time; reliable, because it is accurate and traceable. That pairing is the standard that Decision Quality sets for this link.¹

From theory to practice (making information relevant *and* reliable)

Each practice below pairs a short reason it works with what good output looks like, so the work on information is a visible, auditable link rather than a step waved through. They run roughly in the order the theory sets: first, settle what is relevant and worth gathering, then make the estimate and the process behind it reliable. The frame and the alternatives are already settled when this starts, so the uncertainties worth working on are the ones those two links surfaced.

A. Map the decision-critical uncertainties (and let the value of information set the pace)

Why it works.	What good looks like.
A strategic choice usually turns on a handful of variables and barely moves on the rest, which sensitivity analysis tends to confirm.^{10,11} Value of information says the same from the decision side and adds a stopping rule: a variable is worth deep work only if learning more could change the preferred alternative or move expected value enough to matter, and the work is worth continuing only while that expected gain still exceeds the cost and delay of getting it.³ Once it does not, more refinement buys cost without changing the answer. Holding the team to the few uncertainties that clear that bar is procedural rationality in practice, the information-rich focus that tracks with more effective decisions.^{7,8}	A short map of the three to five uncertainties that actually swing the decision, each with a rough read on how much it moves the answer and the minimum evidence that would settle it. The map is built from more than one seat at the table, so it reflects what matters to people who see the problem differently, not only the sponsor's.² Each uncertainty carries a stopping point: when the expected gain from learning more falls below the cost and delay of getting it, you move to the next requirement rather than keep refining this one. Where an uncertainty matters but resists cheap study, it converts to a fast test rather than a longer wait (practice E). Everything else is named and set aside; if it will not change the choice, it is not on the map.

B. Anchor forecasts in base rates and outside evidence

Why it works.	What good looks like.
Base-rate neglect is one of the better-documented errors in judgment: people weight the specifics of a case and discount what the class of similar cases would predict.⁴ A forecast built from the inside view alone runs optimistic for that reason, and pricing it against a reference class of comparable efforts corrects it.⁵ Across large programs, that correction has done more for cost and schedule realism than further internal analysis.⁶ The same move pulls the estimate off purely internal inputs, setting your own projections beside outside benchmarks instead of reasoning in a closed loop.	A base-rate box for each major forecast: the reference class used, its size, the median and the 10th and 90th percentile outcomes, and the gap between that distribution and your inside-view estimate. Where the final number departs from the base rate, the departure is justified with reasons someone could test, not asserted. A reference class of fewer than about ten comparable cases is flagged as thin rather than treated as settled.

C. Corroborate and challenge decision-critical claims

Why it works.

Independent estimates combined carry less error than any one alone, and triangulating across methods rather than rerunning one raises the validity of what you conclude.^{11,12} A claim resting on a single source carries that source's blind spots straight into the decision.¹³ Search runs toward confirmation by default, so the evidence that would overturn the favored view has to be sought on purpose rather than waited for.^{9,14}

What good looks like.

A one-page source table for each decision-critical claim: what kind of source it is, when and by whom it was produced, the biases that come with it, and whether anything independent corroborates it. Every input is tagged fact, supported assumption, or opinion, so the soft ones cannot pass as hard. No critical claim rests on a single source without either a second independent one or a dated plan to get it. Before the down-select, a short adversarial review by people outside the sponsor's reporting line tries to break the evidence pack, and where two independent sources genuinely conflict, the team writes down the most plausible reconciliation and the check that would settle it.

D. Replace point estimates with calibrated ranges

Why it works.

Overconfidence is pervasive, and it bites hardest at the high-confidence end, where people are most sure and least calibrated.¹¹ Stating a quantity as a range, with a note on what drives the uncertainty, forces the unknowns into view and leaves room to update as evidence arrives. Good forecasters are set apart mainly by calibration, by how well their stated confidence matches how often they turn out right.¹²

What good looks like.

Decision-critical quantities given as a P10-P50-P90 range, or as explicit probabilities, with two lines on what drives the spread and when the estimate will next be revisited. Bare point estimates do not go forward on anything that swings the decision. The few forecasts that matter most get a calibration check over time, scoring earlier estimates against what actually happened.

E. Convert high-value assumptions into fast tests

Why it works.	What good looks like.
Some uncertainties are worth resolving but too slow or costly to study head-on. Discovery-driven planning treats learning as a staged investment: name the assumptions the plan most depends on, test the cheap ones first, and let what comes back set the next commitment. ¹⁵ This is the action arm of the value-of-information test in practice A: when learning is worth it but a study would take too long, run a small test instead of waiting.	A short learn-then-decide plan: one or two fast tests, such as a pilot, a pre-order, or an outside panel, each with a trigger written in advance ("if paid conversion is at or above 5 percent on 200 invitations, proceed; below 3 percent, pivot or stop"). Every critical assumption either has base-rate support from practice B or a scheduled test with a date and a decision rule attached, so none of them sits as an untested article of faith.

Practical limitations (and how to work with them)

- **Value of information needs odds you can estimate** — the test assumes you can put rough probabilities and payoffs on what you would learn, and for a genuinely new decision you often cannot, so the calculation becomes a guess dressed as arithmetic. Use it to structure the question instead, asking what would have to be true to change the call, and where you cannot compute it, settle the point with a cheap test rather than a number (practice E).
- **Reference classes do not always exist** — base rates assume a set of comparable cases, and the genuinely novel move has none. Forcing a class that does not fit imports the wrong baseline, which is worse than admitting you are running on the inside view. Prefer several loose analogies to one forced tight one, mark a thin class as thin, and where there is no outside view to be had, say so and widen the ranges to match.⁵
- **Calibration needs feedback that strategic decisions rarely supply** — the scoring in practice D rewards forecasts that resolve often and cleanly, while a major strategic bet resolves once, slowly, and rarely with a clean verdict. Build calibration on the smaller repeatable judgments inside the decision, the quarterly forecast and the pilot result, rather than expecting to score the one-way door itself.¹²
- **The evidence process runs on incentives it does not control** — what gets surfaced, and how it is framed, is shaped by who benefits from the answer, and a tidy source table does not change that. The main guard is to separate who gathers and weighs the evidence from who owns the call; naming an evidence owner distinct from the decider helps, but the politics is a standing condition to manage rather than a problem the artifacts solve.
- **The discipline has its own cost** — maps, source tables, calibration, and tests all take time, and past the point at which the stakes justify it, they turn the decision into a research project.

Match the effort to how reversible the choice is and how much rides on it: a one-way door with large consequences earns the full apparatus, a reversible bet rarely does.¹⁶

Generative AI as scaffold (not substitute)

Gathering and pressure-testing information divides cleanly into work a model can carry and work it cannot. The trouble comes when the line between them blurs.

AI can draft a first-pass reference class, summarize a stack of interviews or voice-of-customer comments into themes, lay out the uncertainty map, and flag where sources contradict each other. This frees the team to focus on the judgment. What the AI generates, though, is not the analysis. It's evidence that the team still has to check, weigh, and reconcile against its own read of the decision.

What it cannot do is make the judgment calls: whether a source is trustworthy, how much confidence a number deserves, and how much evidence is enough before the risk is acceptable. Those are decisions about what to believe and what to stake on it, and a model has no standing to make them. It can tell you what the comparable cases show; it cannot decide whether your case is comparable.

Three prompts to get you started:

- To scope the work: "List the three to five uncertainties that could change this decision, and for each, the least evidence that would settle it." Keep the ones that move the choice and cut the merely interesting.
- To build the outside view: "Assemble a reference class of comparable cases for this forecast, with the median and the high and low outcomes, and say what each case has in common with ours." Then check that the cases are real and that the likeness actually holds.
- To stress the evidence: "Lay our internal read beside the external indicators and show where they agree and where they conflict." Treat the conflicts as questions to resolve, not as a verdict.

The guardrails are the same across the series. The decision owner signs off on what the team trusts, not the model. Keep a short log of what you asked and what the model drew on, so a number can be traced back to the source. Use an approved workspace and keep proprietary detail out of prompts unless the environment is secured for it. Tag anything factual the model returns as cited or asserted, with a one-line confidence note, before it reaches the room. And run the one check that matters most for this link: where did the model hand us a confident number with nothing real behind it? A model will invent a reference class or a citation as readily as it will retrieve one, and the fabrication is usually as fluent as the fact.

Bottom line

Information is fit for a decision when it is relevant and reliable enough to support the choice. Complete information is rarely available and rarely the point. The work comes down to a few habits: know the three to five uncertainties that actually move the choice, stop gathering on each one when more learning would cost more than it would change, and trust an input only as far as you can trace it and test it. None of this removes the uncertainty under which a strategic choice is made. What it removes is the avoidable kind, the confident number with nothing behind it, the forecast built without a single comparable case, the load-bearing claim no one corroborated. In my view, that is where most of the damage lives, and it is the cheapest to prevent.

References

1. Spetzler, C., Winter, H., & Meyer, J. (2016). *Decision Quality: Value Creation from Better Business Decisions*. Wiley.
2. Ireland, R. D., Hitt, M. A., Bettis, R. A., & De Porras, D. A. (1987). Strategy formulation processes: Differences in perceptions of strengths/weaknesses and environmental uncertainty by managerial level. *Strategic Management Journal*, 8(5), 469-485.
3. Howard, R. A. (1966). Information value theory. *IEEE Transactions on Systems Science and Cybernetics*, 2(1), 22-26.
4. Kahneman, D., & Tversky, A. (1973). On the psychology of prediction. *Psychological Review*, 80(4), 237-251.
5. Kahneman, D., & Lovallo, D. (1993). Timid choices and bold forecasts: A cognitive perspective on risk taking. *Management Science*, 39(1), 17-31.
6. Flyvbjerg, B. (2006). From Nobel Prize to project management: Getting risks right. *Project Management Journal*, 37(3), 5-15.
7. Simon, H. A. (1976). From substantive to procedural rationality. In S. J. Latsis (Ed.), *Method and Appraisal in Economics* (pp. 129-148). Cambridge University Press.
8. Dean, J. W., & Sharfman, M. P. (1996). Does decision process matter? A study of strategic decision-making effectiveness. *Academy of Management Journal*, 39(2), 368-396.
9. Hammond, J. S., Keeney, R. L., & Raiffa, H. (1998). The hidden traps in decision making. *Harvard Business Review*, 76(5), 47-58.
10. Morgan, M. G., & Henrion, M. (1990). *Uncertainty: A Guide to Dealing with Uncertainty in Quantitative Risk and Policy Analysis*. Cambridge University Press.
11. von Winterfeldt, D., & Edwards, W. (1986). *Decision Analysis and Behavioral Research*. Cambridge University Press.
12. Tetlock, P. E., & Gardner, D. (2015). *Superforecasting: The Art and Science of Prediction*. Crown.

13. Russo, J. E., & Schoemaker, P. J. H. (1989). *Decision Traps: Ten Barriers to Brilliant Decision-Making and How to Overcome Them*. Doubleday/Currency.
14. Heath, C., & Heath, D. (2013). *Decisive: How to Make Better Choices in Life and Work*. Crown Business.
15. McGrath, R. G., & MacMillan, I. C. (1995). Discovery-driven planning. *Harvard Business Review*, 73(4), 44-54.
16. Eisenhardt, K. M. (1989). Making fast strategic decisions in high-velocity environments. *Academy of Management Journal*, 32(3), 543-576.