

Clear Priorities and Trade-offs

- Even a flawless analysis settles nothing if the team never says what it actually wants. Executives ask for growth and margin, or speed and safety, as if both came free, and the trade-off stays hidden.
- Clear priorities and trade-offs is the fourth link in Decision Quality. It turns what the enterprise values into measures an option can be scored against, so the choice rests on stated priorities rather than the loudest voice in the room.
- The method is to name a few fundamental objectives, give each a measurable attribute, set thresholds and targets, weight the swings by how much they move value, and state the trade in magnitudes rather than slogans.
- Four moves make it operational: build the value model, set the guardrails, weight the swings, and score the options and name the trade-off.
- Not every decision earns a full value model. Match the depth to the stakes and how reversible the call is; a one-way door justifies the apparatus, a reversible bet rarely does.
- AI helps with drafting objectives, proposing attributes, running the sensitivity views that show which swings carry the choice. It cannot set the weights, the thresholds, or the trade-off sentence; those stay with leadership.

I've found that the analysis around a decision can be impeccable, but it resolves nothing because the team never said what it actually wanted. Executives ask for growth and margin, or speed and safety, as if both come unconstrained. Until someone puts a number on what they will give up to get the other, the trade-off stays hidden, and the tension remains unresolved.

Clear Priorities and Trade-offs is the fourth link in Decision Quality, after the frame is set, the alternatives are on the table, and the information is in hand.¹ Its job is narrow: turn what the enterprise values into measures you can score against, weight those measures by what moves enterprise value rather than by what is easy to collect, and state the trade-offs as clear, quantitative comparisons. This paper covers why priorities frequently remain unclear and outlines the steps I take to make them explicit before making any choices.

The theory in brief (why priorities and trade-offs determine strategy)

Strong priorities don't start with the options on the table. They start with what the enterprise values and how it will recognize that value when it sees it. That is the core of **value-focused thinking**: name a small set of **fundamental objectives**, the ends you care about for their own

sake.² Each objective then needs an **attribute**, an observable scale that says how much of it an option delivers, such as serious safety incidents per million hours, or twelve-month retention. A usable attribute is unambiguous, measures the objective directly rather than through a proxy, and can be scored consistently enough that disagreement centers on the underlying judgment rather than on what the measure means.³ The objectives hierarchy should be complete enough to cover what matters yet concise enough to be practical.

Once objectives are measurable, they still sit on different scales: a point of margin does not net against a lost-time injury. Decision analysis handles this by translating each attribute onto a common value scale, conventionally 0 to 100, then combining those scores with weights that reflect the priorities.^{4,5} The familiar **additive model**, where an option's value is the weighted sum of its scores, assumes objectives are preferentially independent: your preference on one measure doesn't depend on the level of another. That is a working assumption to check rather than a property to take for granted, though it holds often enough to be a reasonable default, and it keeps the logic easy to explain.⁴

Weights should reflect what drives enterprise value, not what happens to be easy to measure. **Swing weighting** makes that explicit. For each attribute, take the full improvement from its credible worst to its credible best, rank those swings against each other, assign 100 to the largest, and scale the rest against it.^{5,6,7} Done well, this helps the weights reflect leadership's actual trade-off judgments rather than data availability or an undifferentiated sense of importance.

A caution rides underneath all of this. Preferences are not always sitting ready to be read off; people often construct them in the act of being asked, and equivalent ways of posing the same trade-off can produce different weights.⁸ That is not a reason to abandon the method but the reason to discipline it: define the swing ranges before eliciting, record the rationale behind each weight, and test the ranking against even swaps, so the numbers reflect a considered trade rather than an artifact of how the question was framed.⁹

Because you can't maximize everything at once, trade-offs have to be stated in magnitudes. Tools like **Even swaps** help convert a difference in one attribute into an equivalent difference in another, holding total value steady, until the dominated options fall away and the choice becomes clear.⁹ The discipline is to specify how much of A you'll give up for how much more of B, not simply which one matters more.

Strategy needs guardrails as well as goals. A **threshold** is a level an option must clear to be considered; a **target** is a level you'd like to reach. Keep them distinct, because a goal acts as a reference point: people respond differently to falling short of it than to clearing it, and effort bunches around the number.¹⁰ So a threshold belongs in the model as a hard constraint or a steep penalty, while a target shapes the value curve above it.⁴ If the two blur, a measure meant to track value becomes a target the organization games, the failure Goodhart identified, and Strathern generalized as a measure ceases to be a good measure once it becomes the target.^{11,12} Twelve-month retention, for instance, is a fair proxy for durable customer relationships until you make the retention figure itself the target; then a team can lift it with lock-

in clauses and cancellation friction while the loyalty it was meant to capture erodes. Modeling the value curve between the walk-away point and the aspiration keeps marginal value visible across the range and keeps “hitting the number” from distorting the outcome.

By making values measurable, weights honest, and guardrails explicit, three outcomes follow: priorities and trade-offs become auditable, so a later challenge meets a written rationale instead of a fresh argument; analysis narrows to the few swing variables that move the needle; and the reasoning is transparent enough to defend to a board and to hand to the people who have to execute it.^{1,5} While this approach doesn’t guarantee success, the evidence shows that teams that decide through a disciplined, information-driven process consistently make stronger strategic choices.¹³

From theory to practice (four moves to make priorities operational)

Each move below pairs a short reason it works with a description of what good output looks like, so priorities become a visible, auditable link rather than a step a team claims to have done. By the time this work starts, the frame is set, the alternatives are on the table, and the information is in hand; if any of those is still open, go back to it first. The outputs are collected on a single page (Exhibit: Value Architecture Canvas), which becomes the durable record of why the choice was made.

A. Build the value model (objectives and the attributes that measure them)

Why it works.

Start from objectives, not from the options already drawn up, and the existing alternatives stop dictating what counts as good. That is the core discipline of value-focused thinking.² An objective is only usable once it has an attribute, an observable scale that is unambiguous, measures the objective directly rather than through a proxy, and can be scored consistently enough that disagreement centers on the underlying judgment rather than on what the measure means.³

What good looks like.

Five to eight fundamental objectives, each separated from the means that merely serve it, with a one-line outcome statement and an attribute: its unit, its direction of preference, a feasible best and worst, and a first sketch of the value function with two or three anchor points.

B. Set the guardrails (thresholds and targets)

Why it works.

A threshold is a level an option must clear; a target is a level worth striving for. Treat them alike and behavior bunches around the number, because a goal acts as a reference point.¹⁰ Worse, a measure pressed into service as a target gets gamed.^{11,12} So, thresholds belong in the model as hard constraints or steep penalties, while targets shape the value curve above them.

What good looks like.

For each objective, the threshold (the walk-away level, where value hits zero) and the target (the aspiration), with the shape of the curve between them. Options that breach a threshold are flagged infeasible early, before anyone spends analysis on them. Each threshold carries a review trigger, so it is revisited when conditions move rather than left to ossify.

C. Find the swing variables (and weight them)

Why it works.

Attributes do not move value equally, so weighting by a gut sense of importance gets the priorities wrong. Swing weighting fixes attention on the full improvement from each attribute's worst to its best, which is what a leader actually trades when choosing.^{5,7} Where objectives are preferentially independent, an additive model combines the weighted scores cleanly; where they interact, say so and test a non-additive form or compare whole policy bundles instead.⁴

What good looks like.

The full-range swings ranked against each other and ratio-scored, the largest set to 100, each with a one-line rationale. The two or three swing drivers that carry most of the decision are marked, so the information effort concentrates there.

D. Score the options and name the trade-off

Why it works.	What good looks like.
Scoring each alternative on its value functions and combining with the agreed weights yields one transparent overall value. Where outcomes are uncertain and can be put in money terms, the same structure extends to an expected value, though that is a move from riskless value to decision under uncertainty and worth marking as one. When the weights are contested, even swaps settles the real magnitudes before the model is locked, by asking how much of one objective the team will trade for a gain on another. ^{5,9}	A scorecard of alternatives against objectives, with 0–100 value scores and the weights alongside. It flags the dominated options, names the swing drivers behind the leading one, marks the flip points where the choice would reverse, and states the trade-off in a sentence: what the team accepts less of, what it gains, and why.

Practical limitations (and how to work with them)

- **A measure under pressure stops tracking what it should** — a proxy attribute, or a single hard target standing in for a value curve, invites the organization to optimize the number instead of the thing it stood for. This is goal displacement: the sub-goal crowds out the objective it was meant to serve.¹⁴ It is the mechanism Goodhart identified and Strathern generalized, that a measure ceases to be a good measure once it becomes the target.^{11,12} Prefer attributes that capture the objective directly, model thresholds as constraints with the value curve carrying the nuance above them, and where only a proxy is available, name the behavior it could reward before you adopt it.
- **Weights shift with framing, and point scores promise a precision they do not have** — a small change in how the swings are described can move the weights, and this is not merely noise: people often construct their preferences in the act of being asked, so how the trade-off is posed can shape the answer it produces.⁸ A single overall value carried to two decimals then reads as authoritative without being reliable. Record the rationale behind each swing weight so it can be re-derived rather than re-argued, carry ranges instead of point estimates on the uncertain inputs, and run a sensitivity check to see which few assumptions actually move the answer. Report the flip points, the places where a changed assumption would reverse the choice, rather than spurious decimals.⁵
- **Different stakeholders value different things, and a model cannot dissolve a real disagreement** — when objectives genuinely compete, resilience against near-term margin, say, the scoring can stall. Separating fundamental objectives from means often shrinks the dispute, since much of it turns out to be about routes rather than ends. For what remains, even swaps turns an argument over which objective matters more into a negotiable

question of how much of one the group will trade for the other, and the trade-off the leaders accept gets written down.⁹

- **Not every decision earns a full value model** — a complete hierarchy with value functions and swing weights is real work, and a reversible, low-stakes choice rarely repays it. Match the depth to what is at stake and how reversible the call is: a one-way door with large consequences justifies the whole apparatus, while a two-way door may need only a few objectives, their thresholds, and a stated trade-off. The aim is to make priorities explicit, not to build the most elaborate model you can defend.¹

Generative AI as scaffold (not substitute)

AI can build the entire framework for trade-offs. It will draft a first-pass set of objectives and candidate attributes with feasible ranges, pull outside-view benchmarks to inform where a threshold or target might sit, scan the source material for contradictory claims, and turn a set of weights and scores into the sensitivity views that show which swings carry the decision.

But the framework is meaningless until it is filled with real priorities. The weights, thresholds, and trade-off statement ultimately reflect what the enterprise truly values, and they are inputs that only the team can provide. While AI can build the scoring system, it cannot tell you what a point of margin is worth against a month of delay. It can generate sensitivity analyses, but it cannot decide which trade-offs leadership is prepared to live with.

This gives us a clean split. The AI builds the framework; the team provides the values, which is the whole point of the exercise. A model can inform a trade-off. It cannot make one.

A few prompts do most of the work. Adapt the wording to your decision and your data.

- **Value model (at kickoff).** “From this brief, list the fundamental objectives and separate them from the means. For each fundamental objective, propose an attribute with its unit, direction, and a feasible best and worst, and flag any objective you cannot measure.” Cut the ones that miss what the enterprise actually cares about.
- **Guardrails (before scoring).** “For each objective, propose a threshold and a target, sketch the shape of the value curve between them, and give one external benchmark for where each sits.” Set the real levels yourself.
- **Swing weights (before the executive session).** “Describe the full-range swing on each attribute in neutral, parallel language. Flag where preferential independence may fail, but do not rank or weight the swings. Produce an elicitation sheet for leadership to use in assigning the weights.”
- **Sensitivity (in the decision meeting).** “Given these weights and scores, identify the dominated options, the swing drivers behind the leading one, and the flip points that would reverse the choice.”

The guardrails are the same ones that apply elsewhere. The decision owner signs the weights, the thresholds, and the trade-off sentence, not the model. Keep a short log of what you asked and what the model drew on. Use an approved workspace and keep proprietary detail out of prompts unless the environment is secured for it. Tag any benchmark the model provides as sourced or asserted, with a one-line confidence note, before it reaches the room. And run one check before you commit: where did the model make a trade-off look settled that leadership never actually made?

Bottom line

No amount of analysis can overcome unclear priorities. When a leadership team can clearly articulate, in a single sentence, what it is willing to sacrifice to achieve its top objective, everything else falls into place: options are scored on what actually matters, thresholds keep an attractive option from quietly breaching non-negotiable values, and the choice can be defended and handed on. Every significant choice involves competing objectives, so you accept less progress in one area to gain more in another. The art is in determining how much less for how much more, documenting the trade-off, and leaning into the clarity it gives. The trade-off written in one sentence is what stops a decision from being relitigated later.

Exhibit: Value Architecture Canvas (one-page)

Use when evaluating alternatives against multiple objectives; complete during option scoring to make weights, thresholds, and trade-offs explicit before commitment.

Value Architecture Canvas

Decision: _____

☐ Type 1 ☐ Type 2 Date: _____ Owner: _____ Version: _____

PURPOSE: Make priorities operational by building value hierarchy, setting guardrails, and making trade-offs explicit. Use when evaluating alternatives against multiple objectives; complete during option scoring.

① FUNDAMENTAL OBJECTIVES & WEIGHTS

#	OBJECTIVE	MEASURE (unit, direction)	WEIGHT
1			
2			
3			
4			
5			

② THRESHOLDS VS. TARGETS

THRESHOLDS

Must-haves / Walk-away constraints

TARGETS

Aspirations / Excellence goals

③ TRADEOFF STATEMENT OFF STATEMENT

④ SENSITIVITY & DECISION RULES

Top 3 Swing Drivers

Key Flip Points

Review Trigger

References

1. Spetzler, C., Winter, H., & Meyer, J. (2016). *Decision Quality: Value Creation from Better Business Decisions*. Wiley.
2. Keeney, R. L. (1992). *Value-Focused Thinking: A Path to Creative Decisionmaking*. Harvard University Press.
3. Keeney, R. L., & Gregory, R. S. (2005). Selecting attributes to measure the achievement of objectives. *Operations Research*, 53(1), 1–11.
4. Keeney, R. L., & Raiffa, H. (1993). *Decisions with Multiple Objectives: Preferences and Value Trade-offs* (2nd ed.). Cambridge University Press.
5. Belton, V., & Stewart, T. J. (2002). *Multiple Criteria Decision Analysis: An Integrated Approach*. Kluwer Academic Publishers.
6. von Winterfeldt, D., & Edwards, W. (1986). *Decision Analysis and Behavioral Research*. Cambridge University Press.
7. Parnell, G. S., Bresnick, T. A., Tani, S. N., & Johnson, E. R. (2013). *Handbook of Decision Analysis*. Wiley.
8. Slovic, P. (1995). The construction of preference. *American Psychologist*, 50(5), 364–371.
9. Hammond, J. S., Keeney, R. L., & Raiffa, H. (1998). Even swaps: A rational method for making trade-offs. *Harvard Business Review*, 76(2), 137–149.
10. Heath, C., Larrick, R. P., & Wu, G. (1999). Goals as reference points. *Cognitive Psychology*, 38(1), 79–109.
11. Goodhart, C. A. E. (1984). Problems of monetary management: The UK experience. In *Monetary Theory and Practice: The UK Experience* (pp. 91–121). Macmillan.
12. Strathern, M. (1997). 'Improving ratings': Audit in the British University system. *European Review*, 5(3), 305–321.
13. 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.
14. March, J. G., & Simon, H. A. (1958). *Organizations*. Wiley.