



THE DECISION MEMO / WORK DESIGN / AUGUST 03, 2026

The \$0 test for AI career advice

A recommendation can keep its direction while the conditions that support it remain unanswered.

Don Long · Founder, SummitMark · Helping people and organizations become capable with AI



An illustrative still life. The offer and scenario figures are not independently verified.

The bonus changed the model's confidence. It did not change its recommendation. That still did not make the decision ready.

A finance and analytics director was weighing a higher-paying external role against a raise and a possible future role at a long-time employer. The AI recommended leaving, with conditions: attainable incentive pay and hours compatible with family life. Neither condition had been checked.

**01 / THE SENSITIVITY TEST**

Five scenarios. One answer.

The second evaluation kept the source-reported \$35,000 base increase constant and varied only the incentive payout. Each level below is illustrative, not a prediction.

INCENTIVE PAYOUT	ANNUAL INCREASE	ANSWER	STATED CONFIDENCE
0%	\$35,000	Accept	60/100
25%	\$65,000 to \$85,000	Accept	68/100
50%	\$95,000 to \$135,000	Accept	76/100
75%	\$125,000 to \$185,000	Accept	83/100
100%	\$155,000 to \$235,000	Accept	88/100

The incentive changed confidence, not direction. The 28-point range is within the second evaluation; the first answer had no numeric score. All five levels were considered in one fresh-context response, not five independent trials. Confidence values are subjective model judgments, not calibrated probabilities. The same response separately gave an overall provisional accept at 70/100 after considering remaining unknowns; that is not a scenario score.

What the test supports

At zero incentive, the model still favored leaving. Its stated reasons were structural dissatisfaction in the current role, closer alignment with the external work, a trusted prospective leader, and the base increase. The uncertain bonus strengthened the case but did not carry its direction.

What it cannot support

The \$35,000 base increase and \$120,000 to \$200,000 incentive range came from the question as summarized in the supplied article; no offer letter or plan document was checked. The payout percentages are scenario inputs, not expected outcomes.

The test says nothing about whether the role is healthy or the family tradeoff is acceptable.

The wider use

The same sensitivity check can travel to any advice that leans on an unverified upside: remove the bonus, ROI estimate, or payback claim and see whether the recommendation changes. If it does, verify the number before treating the advice as ready to act on.

**02 / WHAT THE ANSWER RESTED ON**

Separate the facts from the gaps.

KNOWN FROM THE QUESTION	STILL UNKNOWN
The current job had drifted into systems administration and cross-department conflict. The outside role offered a higher base, possible incentive pay, new scope, and a familiar leader.	Whether the incentive was a target, typical payout, or maximum; actual payout history; hours and travel; staffing and authority; the written terms of both offers; and the person's family priorities.

The questions came after the answer

The original model did identify important diligence questions: whether incentive pay was a target or a theoretical maximum, what comparable leaders had received, why the role was open, whether staffing and authority were real, and what the hours would mean at home. That was useful work. But the recommendation arrived before those facts were available.

A hedge is not a readiness check. If an answer depends on an unverified condition, the model should make clear whether it can responsibly recommend a choice before that condition is tested.

The first answer also assumed the person wanted a CFO track, a goal they had not stated. The second evaluation removed that assumption and still favored leaving. Even so, a plausible career narrative is not a substitute for asking what this person wants.

THE LIMIT

A fresh context can reduce carryover from the first answer, but this was still one model response. The evaluation tested one uncertain number while holding the rest of its scenario fixed; it did not test the other reversal conditions.

What could still reverse the choice

Hours or travel that conflict with family priorities; accountability without real authority; a chronically understaffed team; a new role that is still mostly systems cleanup; adverse plan terms; or an internal offer in writing that actually fixes the current problems. The payout test covered none of these.

**03 / THE NEXT DECISION**

Three paths. One useful next step.

- | | |
|-----------|---|
| 01 | Accept on the headline package
Gives weight to the role, sponsor, and base increase, while treating unverified pay and work conditions as tolerable risks. |
| 02 | Stay for the relationship and promised move
Preserves continuity, but the future role was not yet defined in writing and may not fix the current job's structural problems. |
| 03 | Verify the terms, then decide
Get the incentive plan and comparable payout history; clarify scope, authority, staffing, hours, and travel; ask for the internal role and timing in writing; compare both paths with personal priorities. |

The third path is the most decision-ready next step. It is not advice to accept or decline. It keeps the model's answer from standing in for facts the employers and the person must supply.

The three-column check

GUARANTEED	PROBABLE	MAXIMUM
Written, contractual terms.	Evidence-backed amount, such as plan documents and relevant payout history.	The best case or 'up to' figure, without evidence that it is likely.

Test the decision on guaranteed pay alone, then add only the probable amount. If the answer changes, verify the disputed number before acting. Separately check hours and travel, real authority, team capacity, and whether the internal alternative exists in writing.

Ask the model: Before recommending a choice, list the missing facts that could reverse it. If those facts are unresolved, tell me what to verify before deciding.

Source and scope: Adapted from Don Long's supplied article, "Before You Trust AI Career Advice, Run This \$0 Test." It describes an anonymized public r/careeradvice question and reports two AI evaluations. Reddit blocked full-page extraction during review, so case details and model outputs are attributed to the supplied article, not independently verified. The scenario percentages are illustrative; the confidence scores are not calibrated probabilities. This checklist is editorial guidance, not a validated career or financial framework. No recommendation is made for the person in the case. Draft for author review; confirm date and byline before send.