

Logic Model Prompt Pack — Sopact Academy

From the article **How to Build a Logic Model with AI, in Minutes** — sopact.com/academy/how-to-build-a-logic-model Works in Sopact Sense, Claude, or ChatGPT. No setup needed: paste, swap the [BRACKETS], run.

The master prompt (run everything in one pass)

Build a Logic Model for [PROGRAM NAME] using only what the program publicly states at [SOURCE – URL or pasted program description].

DECISION FRAME

Before building, state in one line who would use this logic model and for what decision (e.g., "a funder deciding whether to renew"). Make every judgment below from that reader's perspective.

PART 1 – VISUAL LOGIC MODEL

Lay the program out in columns: Inputs → Activities → Outputs → Short-Term Outcomes → Medium/Long-Term Outcomes → Impact.

- Write every outcome as a change in people, never as an activity ("mentees gain career clarity," not "we run workshops").
- Flag any orphan activity (one that feeds no output) and any empty column.
- Color every box: GREEN = specific AND evidenced; AMBER = stated but vague, or specific but unevidenced; RED = missing, or [INFERRED].
- Tag anything not explicitly stated as [INFERRED].
- Include a legend, program name, source URL, and date.

PART 2 – GAP & INDICATOR TABLE

One row per box: Element | Column | Grade | Evidence (quoted, paraphrased, or "none") | one measurable indicator that would test it, phrased so the program could actually collect it (who is measured, what changes, by when).

PART 3 – WEAK-LINK DIAGNOSIS

List every orphan activity and every AMBER and RED box, ranked by how much the model depends on it. For each: (a) why it's weak in one sentence, (b) what breaks downstream, (c) whether the fix is a program-design problem or a measurement/evidence gap.

PART 4 – IMPROVEMENT PLAN

The top 3–5 fixes, in priority order for the decision named above. For each: current language (or "missing") → proposed rewrite, plus the single data collection step that would move it toward green.

CLOSING SUMMARY

3–4 sentences: overall strength of the logic, the weakest box a skeptical funder would circle first, and the one action this quarter.

RULES

- Source fidelity is absolute: never invent program content. If the source doesn't say it, mark it RED or [INFERRED].

- Every green grade must be traceable to specific source language.
- The model and the tables must agree.

Example source: <https://www.lanternnetwork.org/mentoring-program>

Run the parts one at a time

Part 1 — the visual logic model

Lay the program out in columns: Inputs → Activities → Outputs → Short-Term Outcomes → Medium/Long-Term Outcomes → Impact. Write every outcome as a change in people, never as an activity. Flag any orphan activity that feeds no output, and any empty column. Color every box GREEN (specific AND evidenced), AMBER (stated but vague, or specific but unevidenced), or RED (missing, or [INFERRED]). Tag anything not explicitly stated as [INFERRED]. Include a legend, program name, source URL, and date.

Part 2 — the gap & indicator table

For every box, produce one row: Element | Column | Grade | Evidence (quoted, paraphrased, or "none") | one measurable indicator that would test it, phrased so the program could actually collect it – who is measured, what changes, by when.

Part 3 — the weak-link diagnosis

List every orphan activity and every AMBER and RED box, ranked by how much the model depends on it. For each: (a) why it's weak in one sentence, (b) what breaks downstream, (c) whether the fix is a program-design problem or a measurement gap – these require different responses.

Part 4 — the improvement plan

Give the top 3–5 fixes, in priority order for the decision named in the frame. For each: show the current language (or "missing") → a proposed rewrite, plus the single data collection step that would move it toward green. Close with 3–4 sentences: overall strength of the logic, the weakest box a skeptical funder would circle first, and the one action to take this quarter.

Bonus — tighten your program page

Based on the grades above, suggest edits to my program page so its claims match the evidence. Flag every sentence that overstates what we can show, and rewrite it to be accurate and specific.

The same prompts work for a Theory of Change, Logframe, or Results Framework — swap the framework name and keep the parts, rubric, and rules unchanged.

Sopact Sense keeps the whole loop in one place — clean data collection with persistent participant IDs, so the indicators these prompts propose are ones you can actually track. Try it: <https://www.sopact.com/request-demo>