

Theory of Change Prompt Pack — Sopact Academy

From the article **How to Build a Theory of Change with AI, in Minutes** — sopact.com/academy/how-to-build-a-theory-of-change 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 Theory of Change 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 TOC and for what decision (e.g., "a funder deciding whether to renew"). Every judgment below should be made from that reader's perspective.

PART 1 – VISUAL THEORY OF CHANGE (single interactive HTML artifact)

Render the full causal chain as a left-to-right flow diagram:

Inputs → Activities → Outputs → Short-Term Outcomes → Medium-Term Outcomes → Long-Term Outcomes → Impact.

- One node per element; write every outcome as a change in people, never as an activity ("mentees gain career clarity," not "we run workshops").
- On each arrow, label the causal assumption that must hold for the left node to produce the right node.
- Color every node and arrow: GREEN / AMBER / RED (rubric below).
- Hovering any node shows the exact source language (quote or close paraphrase) that supports it; hovering any arrow shows the assumption and its proposed indicator.
- Tag anything not explicitly stated by the program as [INFERRED], both in the diagram and in text.
- Include a legend, program name, source URL, and date.

GRADING RUBRIC (apply to every node and arrow)

- GREEN: specific AND evidenced – the program names it concretely and provides data or documented examples (e.g., "87% secured internships").
- AMBER: stated but vague, or specific but unevidenced – no number, no timeframe, no named population, or a claim with no supporting data.
- RED: missing entirely, or exists only as [INFERRED].

PART 2 – ASSUMPTION & INDICATOR TABLE

One row per arrow: Assumption | 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 AMBER and RED element, ranked by how much the causal chain depends on it. For each: (a) why it's weak in one sentence, (b) what claim collapses downstream if the assumption fails, (c) whether the fix

is a program-design problem or just a measurement/evidence gap – these require different responses.

PART 4 – IMPROVEMENT PLAN

The top 3–5 fixes, in priority order for the decision named above.
For each: show the current language (or "missing") → proposed rewrite, plus the single data collection step that would move it toward green (e.g., "add a baseline/endline confidence measure with persistent participant IDs so the same mentees are tracked over time").

CLOSING SUMMARY

3–4 sentences: overall strength of the causal logic, the single weakest link a skeptical funder would attack first, and the one action to take this quarter.

RULES

- Source fidelity is absolute: never invent program content. If the source doesn't say it, mark it RED or [INFERRED] – do not fill gaps with plausible-sounding elements.
- Every green grade must be traceable to specific source language.
- Diagram and text must agree; the tables are the evidence trail for the visual.

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

Run the parts one at a time

Use these when you want to build step by step. Run the master prompt's first two lines plus the DECISION FRAME first, then each part below in order.

Part 1 — the visual theory of change

Render the full causal chain as a left-to-right flow diagram: Inputs → Activities → Outputs → Short-Term Outcomes → Medium-Term Outcomes → Long-Term Outcomes → Impact. One node per element; write every outcome as a change in people, never as an activity. On each arrow, label the causal assumption that must hold. Color every node and arrow GREEN (specific AND evidenced), AMBER (stated but vague, or specific but unevidenced), or RED (missing, or [INFERRED]). Hovering any node shows the exact source language that supports it; hovering any arrow shows the assumption and its proposed indicator. Tag anything not explicitly stated as [INFERRED]. Include a legend, program name, source URL, and date.

Part 2 — the assumption & indicator table

For every arrow in the diagram, produce one row: Assumption | 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 AMBER and RED element, ranked by how much the causal chain depends on it. For each: (a) why it's weak in one sentence, (b) what claim collapses downstream if the assumption fails, (c) whether the fix is a program-design problem or a measurement/evidence 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 causal logic, the single weakest link a skeptical funder would attack 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 Logic Model, 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>