

Results Framework Prompt Pack — Sopact Academy

From the article **How to Build a Results Framework with AI** — sopact.com/academy/how-to-build-a-results-framework 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 Results Framework 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 framework and for what decision (e.g., "a funder deciding whether to renew"). Make every judgment below from that reader's perspective.

PART 1 – RESULTS HIERARCHY

Map the hierarchy: one Strategic Objective, 3–5 Intermediate Results that must each occur to reach it, and the Sub-IRs beneath them.

- Give every level a baseline-and-target indicator.
- Flag any result that doesn't ladder up to the level above, and any indicator with no baseline.
- Color every element: GREEN = specific AND evidenced, with a baseline; AMBER = stated but vague, or a target with no baseline; RED = missing, doesn't ladder up, or exists only as [INFERRED].
- Tag anything not explicitly stated as [INFERRED].
- Include a legend, program name, source URL, and date.

PART 2 – INDICATOR & BASELINE TABLE

One row per result: Result | Level (SO / IR / Sub-IR) | Grade | Baseline & Target (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 result and every AMBER and RED element, ranked by how much the framework depends on it. For each: (a) why it's weak in one sentence, (b) what claim collapses if it fails, (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 results logic, the weakest link a skeptical funder would attack 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 and a baseline.
- The hierarchy and the tables must agree.

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

Run the parts one at a time

Part 1 — the results hierarchy

Map the hierarchy: one Strategic Objective, 3–5 Intermediate Results, and Sub-IRs beneath them, each with a baseline-and-target indicator. Flag any result that doesn't ladder up and any indicator with no baseline. Color every element GREEN (specific, evidenced, baselined), AMBER (stated but vague, or a target with no baseline), or RED (missing, doesn't ladder up, or [INFERRED]). Tag anything not explicitly stated as [INFERRED]. Include a legend, program name, source URL, and date.

Part 2 — the indicator & baseline table

For every result, produce one row: Result | Level (SO / IR / Sub-IR) | Grade | Baseline & Target (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 result and every AMBER and RED element, ranked by how much the framework depends on it. For each: (a) why it's weak in one sentence, (b) what claim collapses if it fails, (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 results logic, the 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 Theory of Change, Logic Model, or Logframe — 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>