
10 Common Mistakes When Building a Logic Model

...and how to avoid them.

By David B. Howard, MSW, Ph.D.
August 2026

Introduction

Logic models are the connective tissue between what your program does and the change it's trying to achieve. When they're well-designed, they help clarify strategy, sharpen evaluation, and give funders a defensible account of your theory of change. When constructed poorly, they can leave stakeholders confused about how the core ingredients of a program intend to produce results — or whether certain outcomes are feasible.

This guide highlights ten common mistakes practitioners make when building logic models, and how to correct them. Some are conceptual (e.g., misunderstanding what a logic model is for). Some are content-level (e.g., framing outcomes in ways that can't be measured). Some are about process (e.g., building it in isolation, or letting it collect dust). All of them are fixable!

Mistake #1**Skipping the theory of change***Listing activities without connecting the dots*

A logic model is fundamentally a theory of change argument: if we do these activities, then we'll produce these outputs, which should lead to these outcomes. Without that causal logic, you have a list – not a model. And a list can't tell you where your program might fail, what assumptions you're relying on, or how to improve.

Teams often skip this step because they know their program intuitively and treat the causal logic as a given. But if it's obvious to you, it should be easy to write down. If you can't explain the mechanism by which your inputs and activities cause your outcomes, that's a signal that the program design itself may need work – required steps before the logic model can be most helpful.

HOW TO FIX IT

Before filling in the sections, write one sentence: "Our program takes [insert inputs] and delivers [insert activities] because we believe [insert mechanism] leads to [insert outcome]." If that sentence doesn't hold together, you're not ready to build the model. You need to think more about your theory of change, and let that become the basis for a more sound logic model.

Mistake #2**Confusing outputs with outcomes***Counting activities instead of change*

One of the most common – and most consequential – mistakes in logic model construction is treating outputs (how we count what a program *does*) as if they were outcomes (how we measure what a program *changed*). Examples of outputs include the "Number of participants served," "hours of counseling provided," and "meals distributed." They can tell you the extent to which a program was implemented, and even whether that implementation met critical fidelity standards, but they don't tell you whether the program achieved its intended results.

When outputs are labeled as outcomes, evaluations get stuck measuring efforts instead of impact. Funders reading these logic models can't tell what the program is actually trying to accomplish. And the program itself has no way to learn whether its work is leading to actual change.

HOW TO FIX IT

Make sure your outcomes are measuring *change*. If your metric is gauging change in skills, knowledge, attitudes, behaviors, conditions, or status, it's an outcome. Change is always going to be directional (e.g., "Increased...", "Decreased...", "Improved...", "Reduced...", etc.). If you're counting activities – and not the change those activities seek to achieve – then you're measuring outputs.

Mistake #3**Vague outcomes with no way to measure them***High-level goals that are difficult to track*

"Improved well-being." "Better community cohesion." "Increased self-sufficiency." These sound like outcomes, but they're more like high-level goals. Without a clear way to measure them, they can't be tracked, reported, or improved. And without knowing whether the outcome was achieved, the program can't learn.

The test: for every outcome on your logic model, can you name a specific measure that will tell you whether it happened? If not, then it's a value statement, not an outcome. What's the data source? What's the target? What's the baseline?

HOW TO FIX IT

Pair every outcome with an indicator that names how you'd measure it. Rather than "Improved housing stability," perhaps it's "% of participants retaining housing at 12 months." Be sure to decide during logic model development how you'll measure each outcome – not after-the-fact.

Mistake #4**Conflating short-, intermediate-, and long-term outcomes***Blurring time horizons*

When outcomes across all time horizons – "increase knowledge," "housing stability at 12 months," "reduction in chronic homelessness system-wide" – get lumped together as "outcomes," the logic model loses its explanatory power. Reviewers can't tell what's expected sooner rather than later, and program staff can't tell what to measure when.

Short-term (~0-3 months), intermediate (~3-12 months), and long-term (12+ months) outcomes serve different purposes. Short-term outcomes tell you the program is working as designed. Intermediate outcomes tell you participants are changing. Long-term outcomes tell you the change is sustaining and the broader population is being affected. Blur them together and you lose that diagnostic power.

HOW TO FIX IT

For every outcome, place it in the right tier based on how long after program entry you'd expect to see it. If you're not sure, err toward later – most outcomes take longer to achieve than we assume. The three-tier structure isn't decorative; it's a diagnostic tool.

Mistake #5**Mismatched time horizons***Outcomes that outlast the grant*

If your grant runs for 12 months and your logic model's intermediate or medium-term outcomes take 3-5 years to manifest, you've built a mismatch that will haunt every evaluation report. Funders will ask about progress on outcomes you can't yet measure. Staff will feel like they're failing on metrics that weren't achievable in the time available.

Well-designed logic models align outcome timing with what's realistic and also with relevant funding cycles and reporting requirements.

HOW TO FIX IT

For every outcome, ask "when would we expect to see this?" If the answer is longer than your reporting cycle, make sure it's framed as long-term and consider a proximate indicator you can measure within the timeframe. Example: "Long-term reduction in chronic homelessness" (unmeasurable in 12 months) can be paired with "% of participants housed within 90 days" (a reasonable intermediate outcome).

Mistake #6**Making the model too complex***Trying to capture everything*

A logic model is a summary, not a comprehensive program inventory. Teams that try to list every input, activity, output, and outcome their program touches run the risk of creating a cluttered and unfocused tool that can be impossible to evaluate against. The most useful logic models capture the four to six items per section that matter most. More complex programs may require a higher volume. Use your best judgement.

Complexity can signal underlying confusion about the program itself. If everything feels essential, your theory of change may lack clarity. What are the core activities that *actually* drive the change you're seeking? Those are what belong in the logic model. The rest may be too granular for a digestible summary.

HOW TO FIX IT

Set a limit before you start – four to six items per section, maximum. If a fifth item is competing for space, consider which existing item is less essential. A logic model that fits on one page and gets used is worth more than a comprehensive version that gets filed away and ignored.

Mistake #7**Overpromising**

Outcomes bigger than the program can deliver

A small case management program serving 50 clients per year probably isn't going to "eliminate homelessness in the county." Overstated outcomes can hurt credibility with funders, set the program up for failure in evaluation, and often reflect dissonance between what the program can influence and what it can control.

The distinction: your program can control its outputs (how many people served, how many sessions delivered) and its short-term outcomes (participants' immediate response to services). It can influence intermediate outcomes (behavior change over months). It contributes to – but typically cannot claim credit for – long-term population-level changes, which depend on the broader system.

HOW TO FIX IT

Match outcome ambition to program scale. Ask "if we hit our most optimistic activity numbers, what's the biggest change we could plausibly claim?" That's your ceiling. Anything beyond that belongs in the "long-term aspirational goals" narrative, not necessarily in the logic model.

Mistake #8**One-size-fits-all outcomes**

Generic outcomes that could apply to any program

Outcomes like "improved quality of life" or "better outcomes for participants" (in addition to being vague!) can be attached to virtually any human services program, which is exactly why they don't help you design or evaluate yours. Strong outcomes should reflect the specific sector, population, and program design of the work being done.

A homeless services program's outcomes should look different from a workforce development program's, which should look different from a child welfare program's. Beyond that, outcomes should reflect your specific population (unaccompanied youth vs. families vs. veterans, for example) and your specific approach (Housing First vs. Transitional Housing; group vs. individual counseling).

HOW TO FIX IT

Start from your sector's standard outcome measures — HMIS System Performance Measures for homeless services; WIOA PIRL indicators for workforce; CFSR indicators for child welfare. Using field-tested measures as a starting point will help you stay aligned with your specific population and program design.

Mistake #9**Not connecting outcomes to funder priorities or reporting frameworks***Outcomes your funder doesn't recognize*

Ideally, your logic model will be program-driven rather than funder-driven. But funders play a critical role – and many will require a logic model. Federal, state, and philanthropic funders each have specific frameworks for what "counts" as a successful outcome. HUD wants CoC APR and System Performance Measures. DOL wants WIOA PIRL indicators. Foundations often have their own reporting frameworks. If your outcomes don't map to these, you'll either be doing double work (measuring both) or your reporting won't fit what funders expect.

This isn't about writing outcomes to please funders. It's about making sure your outcomes speak the same language your funders (or prospective funders) speak – because otherwise, the good work you're doing may well be overlooked in review.

HOW TO FIX IT

Before finalizing outcomes, review the reporting requirements of your top funders. Note which metrics they require or prefer. Where possible, align at least a few of your outcomes to those frameworks. Adopt the same terminology so reviewers can find them quickly.

Mistake #10**Building it in isolation and letting it go stale***Built in a silo and left to gather dust*

Two related problems: logic models built by one person in isolation – perhaps a program director or a grant writer – often miss what the program actually does day-to-day. Program staff who deliver the services have knowledge about the mechanisms, special cases, and unintended outcomes that never make it into the model. And once the model is built, it too often gets filed away, untouched until the next grant cycle, even as the program evolves.

A logic model is most useful when it reflects the collective understanding of those running the program and can be revisited as things change. That's what makes it a working tool, not a compliance document.

HOW TO FIX IT

Involve multiple direct-service staff members in building the logic model, and set a calendar reminder to revisit it every 6-12 months. Ask: Does this still describe what we're doing? What's changed? What did we learn from outcomes that surprised us? Incorporate those logic model check-ins into team gatherings that allow time and space to take a step back from the day-to-day work.

About DBH Consulting

DBH Consulting helps nonprofits, foundations, and public agencies design programs that work, measure what matters, and report what funders need to see. If any of the mistakes in this guide feel familiar — or if you're building a new program and want expert eyes on your logic model before your next funder review — we'd love to help.

Ready to build one now? Try our free Logic Model Builder at dbh-consulting.com/logic-model-builder

Ready to talk? Book a free consultation at dbh-consulting.com/#contact