



Win/Loss Research

THE COMPLETE GUIDE

The Win/Loss Research Playbook

All twelve plays, scoping a program through measuring whether it's working.

SCOPE & DESIGN • RECRUIT & OUTREACH • INTERVIEW • ANALYZE •
REPORT & DISTRIBUTE • SUSTAIN

Daniel Oxenburgh

FOUNDER, OX WIN/LOSS — WINLOSSRESEARCH.COM

How to Scope a Win/Loss Research Program

Scope a win/loss program by setting five inputs before outreach begins: the program's objective and success criteria, the deal segment under review, an interview count and ratio weighted toward losses, a research timing window relative to deal close, and a deal-sourcing method that pulls the list from the CRM rather than through sales approval.

In this play

1. Define the objective and success criteria
2. Set the deal segment
3. Set interview count, ratio, and timing window
4. Pull the deal list under executive sponsorship
5. Estimate the effort range before committing

Step 1: Define the Objective and Success Criteria

Write the program's objective as a single sentence before setting any other scope input.

"Understand our win rate better" is a discipline, not an objective, and a program scoped around a discipline drifts once the first busy quarter arrives. "Explain the ten-point win rate drop in Enterprise over the last two quarters" is an objective, because it names a metric that already moved and gives the program a specific question to answer.

Attach a success criterion to the objective at the same time. A useful test: name the decision the program needs to change. If a program can't point to a specific pricing call, messaging test, or resourcing decision it's meant to inform, the objective isn't specific enough yet. Revise it until it names one.

A scoping worksheet forces this discipline before interview planning starts:

SCOPING WORKSHEET – OBJECTIVE

Program Objective: [one sentence naming the specific, already-felt problem]
Decision This Program Should Inform: [the pricing, messaging, or resourcing call]
Executive Sponsor: [name, title]
Target Report Date: [date leadership needs an answer by]

Don't leave the report date blank. A program scoped without one tends to lose the urgency that got it funded in the first place.

Step 2: Set the Deal Segment

Define which deals fall inside the program before pulling a single record. Scope by segment (Enterprise, Mid-Market, SMB), deal size range, geography, and product line, matched to whatever the objective in Step 1 actually points to. A program scoped around an Enterprise win rate drop should pull only Enterprise deals. Mixing segments into one interview set blurs a finding that may only be true for one of them.

Set inclusion and exclusion criteria explicitly rather than leaving them implied. A standard set: closed-won and closed-lost deals only (exclude deals still active or stalled with no decision), deals closed within the research window defined in Step 3, and deals above a minimum size threshold if very small deals follow a meaningfully different buying process than the segment the program is actually trying to understand.

Step 3: Set Interview Count, Ratio, and Timing Window

Target 20 to 30 total interviews, split roughly two loss interviews for every win interview, commonly 20 losses and 10 wins. Below about 15 interviews, a single unusual conversation can distort the whole dataset. Above 30, additional interviews tend to confirm what's already emerged rather than surface anything new. Losses carry more weight because a loss usually has a specific, nameable cause, while a win can happen for several reasons at once that are hard to isolate.

Schedule interviews 30 to 180 days after deal close. Inside 30 days, a buyer is often still transitioning into onboarding and hasn't had distance to reflect. Past 180 days, specific detail compresses into a general impression. A straightforward, single-stakeholder deal can support an interview closer to the 30-day mark; a longer enterprise cycle with a large buying committee may still be settling at 90 or 120 days.

Step 4: Pull the Deal List Under Executive Sponsorship

Source the deal list directly from the CRM using the criteria set in Step 2, not from a shortlist sales nominates. Letting reps flag which deals are worth reviewing hands curation to the people whose read on each loss is, by definition, one-sided. The deals a rep leaves off the list are frequently the ones where the buyer's actual reason for leaving would be most revealing.

Put an executive sponsor's authority behind the pull instead. Sales leadership can be informed of the research window and consulted on which segments to prioritize strategically, but individual deal-by-deal approval is where selection bias enters, and it stays out only if that line holds. Frame the program to sales explicitly as pattern recognition across an aggregated deal set, not an evaluation of individual reps or deals, before the first outreach email goes out.

Step 5: Estimate the Effort Range Before Committing

Running a program internally costs real hours, and naming a bounded range before committing avoids a team discovering the true cost mid-cycle. A standard 20-to-30 interview cycle typically requires:

EFFORT RANGE – WORKSHEET

Scoping and executive sponsorship alignment: 5-10 hours
Target list build and outreach: 15-25 hours
Interview scheduling and coordination: 10-15 hours
Conducting 20-30 interviews at up to 1.5 hours each, including prep and debrief: 30-45 hours
Analysis and pattern synthesis: 20-30 hours
Report drafting and readout prep: 15-25 hours

Total: roughly 95-150 hours across a 75-to-90-day cycle

Treat this as a planning range, not a precise forecast. Actual hours shift with buying committee complexity, questionnaire length, and how experienced the person running interviews is. The range exists to answer one question before scoping goes further: does the team committing to this program have 95 to 150 hours available across the next 75 to 90 days, spread across the roles above, and can it hold that commitment through the full window without the cycle stalling out. If the answer is no, that's a scoping decision to make now, not a discovery to make in week six.

REALITY CHECK

An hours range shows how much time a scoped program takes to run, not whether the team can commit that time inside the window, or whether an internal interviewer gets the same candor a neutral third party gets by default. Neutrality and a real time-bound commitment, not the hours total, determine whether a scoped program delivers an actionable answer. See [why third-party neutrality changes what buyers will say](#).

How to Build a Win/Loss Interview Target List

A win/loss interview target list starts with a CRM query, not a list sales selects deal by deal. Pull every deal that matches the program's scope criteria, size that raw pull well above the interview target to offset a 5-to-10 percent participation rate, and clean up contact data before outreach begins.

In this play

1. Pull the deal population directly from the CRM
2. Size the list against the participation rate
3. Enrich contact data before outreach begins
4. Exclude deals that don't belong in the sample

Step 1: Pull the Deal Population Directly from the CRM

Query the CRM using the scope criteria already set (segment, deal size range, geography, time window), not a list a rep assembles by hand. Build the query on hard fields only: deal stage equals closed-won or closed-lost, close date falls inside the research window, segment matches the program's scope. A soft field like a rep's own "worth revisiting" flag has no place in this pull, since that judgment is exactly what a criteria-based process is designed to remove.

A sample filter set for a program scoped to Enterprise deals over the last two quarters:

SAMPLE FILTER SET

Deal Stage: Closed Won OR Closed Lost
Close Date: Between [180 days ago] and [today]
Segment: Enterprise
Deal Size: >= \$[minimum threshold]
Excludes: Deals still marked "Negotiation" or "Active"

Export the raw result before doing anything else to it. That export is the population the rest of this play works against, and having it as a fixed starting point makes every later exclusion in Step 4 auditable against the original set.

REALITY CHECK

A CRM-sourced list removes formal deal-by-deal approval, but reps can still exercise informal control over who ends up reachable: an outdated title, a missing email, or a stale phone number on the buyer who left the angriest voicemail. A criteria-based pull is only as unbiased as the contact data reps chose to keep current. See [why internal win/loss data fails](#).

Step 2: Size the List Against the Participation Rate

Loss interview participation typically runs around 5 percent; win interview participation runs around 10 percent. Multiply the interview target by the inverse of that rate to get the outreach list size the program actually needs, not the interview count itself.

PARTICIPATION MATH

```
Target: 20 completed loss interviews
Participation rate: ~5%
Required loss-contact list: 20 / 0.05 = 400

Target: 10 completed win interviews
Participation rate: ~10%
Required win-contact list: 10 / 0.10 = 100
```

If the raw CRM pull from Step 1 doesn't produce a population this large after Step 4's exclusions, that's a scoping decision to revisit now, either by widening the segment or extending the research window, not a gap to discover after outreach is already underway.

Consider a program scoped to Enterprise deals over the last two quarters. After Step 4's exclusions, the loss-eligible population comes to 340 contacts, short of the 400 the participation math requires. The fix isn't to lower the interview target and quietly accept fewer completed interviews. It's to extend the research window from 180 to 270 days to pull in an earlier quarter's

deals, or widen the segment to include a closely adjacent Mid-Market tier, and to make that call before the first invitation goes out rather than mid-cycle when the shortfall shows up as a stalled interview count.

Step 3: Enrich Contact Data Before Outreach Begins

CRM contact records degrade the moment a deal closes. A champion moves to a new company, a direct line routes to a switchboard, an email bounces. Verify email addresses and titles through a data-enrichment tool or a quick manual cross-check before the first invitation goes out, since a bounced or misdirected email counts against the list size calculated in Step 2 without producing a single response.

A practical enrichment pass has two parts: an automated check against an email-verification service to catch hard bounces before sending, and a manual spot-check on any contact whose title or company hasn't been touched in the CRM since before the deal closed, since that's the most common sign the record is stale. Flag records that fail both checks as unreachable rather than including them in the outreach count, and treat the corrected total as the real list size going into Step 2's math.

Prioritize enrichment on the loss-contact list first. It needs to run roughly four times the size of the win-contact list to hit the same completed-interview target, so a data gap there costs the program more outreach volume than the same gap on the smaller win list.

Step 4: Exclude Deals That Don't Belong in the Sample

Some deals should come out of the raw pull even though they technically match the scope criteria in Step 1. Remove deals still in active renewal negotiation, since a live commercial relationship makes an interview request awkward and can bias what the buyer is willing to say. Remove any contact who has explicitly opted out of future outreach. Remove internal test or demo accounts that occasionally get marked closed-won in the CRM.

Apply these exclusions as fixed rules that treat every deal the same way, not case-by-case judgment calls about which deals are "worth" reviewing. That distinction is what keeps this step a hygiene pass rather than the same selection bias a criteria-based CRM pull was built to avoid in the first place.

Keep a short exclusion log alongside the final list: how many records came out of the Step 1 export, under which rule, and how many remain. That log is what lets anyone question the final target list later and get a specific answer, rather than a general assurance that the list is representative.

How to Get Sales Buy-In for Win/Loss Interview Access

Sales buy-in for a win/loss program comes from executive sponsorship of its scope, not sign-off on individual deals. Frame the program explicitly as pattern recognition across an aggregated deal set rather than an evaluation of specific reps, and set that boundary with sales leadership before the first outreach email goes out.

In this play

1. Frame the program as pattern recognition, not rep evaluation
2. Secure executive sponsorship over deal-by-deal approval
3. Set the boundary between strategic input and deal approval
4. Handle the objections that come up in practice

Step 1: Frame the Program as Pattern Recognition, Not Rep Evaluation

State up front, before asking for anything, what the program is not: it does not evaluate individual reps or grade specific deals they handled. Lead with this in the first conversation rather than holding it in reserve for when an objection raises it.

Sales teams support a program once it's clear the output is GTM-level insight, not deal-specific critique, because that framing removes the reason any individual rep would feel exposed by the process. Findings get distributed as patterns across an aggregated deal set, and no single rep's deals are named in the readout that follows.

Step 2: Secure Executive Sponsorship Over Deal-by-Deal Approval

Ask for sponsorship of the program's scope, the research timeframe, segment, and deal criteria already set, not sign-off on each deal that gets pulled. A senior sales leader backing that scope gives the program authority without handing curation to individuals who have a direct stake in how their own losses read.

Put the ask in writing and keep it specific: name the segment and time window, and request one sponsor's approval of that scope as a single decision, not a standing review queue of individual deals as they come up.

SPONSORSHIP REQUEST – EMAIL TEMPLATE

Subject: Sponsorship request: win/loss research on [segment] deals

[Sponsor name],

We're scoping a win/loss research program covering [segment] deals closed between [start date] and [end date], targeting roughly 20-30 buyer interviews weighted toward losses. The deal list will be pulled directly from the CRM against fixed criteria, not curated deal by deal.

I need your sponsorship on the scope itself: this segment, this window, this criteria set. No individual rep sign-off on specific deals is part of this process, and findings will come back as aggregate patterns, not deal-by-deal or rep-by-rep breakdowns. You'll be included in the readout once findings are ready.

Can we confirm this by [date]?

Step 3: Set the Boundary Between Strategic Input and Deal Approval

Sales leadership stays involved without gaining veto power over individual deals. The line to hold:

SALES CAN	SALES CANNOT
Be informed of the research window and deal set in advance	Approve or veto individual deals before they're included
Weigh in on which segments or deal types the program should prioritize	Remove a specific deal because the loss "isn't representative"
Be included among the stakeholders who receive findings in the readout	Filter which contacts get outreach based on how a deal went

REALITY CHECK

Executive sponsorship sets the formal rule, but it doesn't stop a rep from quietly steering a difficult buyer away from an outreach list, or flagging a contact as unreachable rather than explaining why. The boundary between strategic input and deal-level approval only holds if someone is actually watching for it, not just stating it once at kickoff. See [why internal win/loss data fails](#).

Step 4: Handle the Objections That Come Up in Practice

Two objections recur in nearly every version of this conversation. "This will hurt morale if reps think they're being graded" gets addressed by the framing in Step 1, restated directly: findings are distributed as GTM-level patterns, not individual scorecards, and no rep's specific deals are named. "We already know why we lose deals" is worth taking at face value rather than arguing with. Offer to run the program as a validation exercise against the existing internal theory instead, which either confirms it with independent evidence or surfaces where it's incomplete, and either outcome is useful to the team holding that theory.

A third objection shows up less often but stalls a program just as effectively when it does: "Let sales flag which deals make sense to include, we know which ones are worth the buyer's time." This is the deal-by-deal approval risk from Step 3 arriving in a more reasonable-sounding form. The response is the same regardless of how the request is framed: strategic input on segment and timing is welcome, individual deal selection is not, because the deals a rep would flag as not worth including are frequently the ones where the buyer's actual reason for leaving is most revealing.

A short script works better than a general pitch for the initial ask:

INITIAL ASK – SCRIPT

"We're running a win/loss research program on [segment] deals from [date range]. It's not a review of any rep's individual deals, it's pattern recognition across the full set, and everyone shows up in aggregate findings, not by name. I need your sponsorship on the scope: this segment, this time window, pulled directly from the CRM. You'll be in the readout when findings come back."

Once sponsorship is confirmed, put the scope, the boundary from Step 3, and the readout commitment in a single follow-up note back to the sponsor. That written record is what the program points to later if a rep raises the "let us flag deals" request again mid-cycle, rather than re-litigating the boundary from scratch each time it comes up.

How to Write a Win/Loss Interview Outreach Sequence

A win/loss outreach sequence runs three to four touches over about two weeks: an initial invitation that states its purpose clearly and stays separate from the sales relationship, followed by two follow-ups spaced several days apart. Include a modest interview incentive in the first email, confirm it again at scheduling, and write lighter, differently framed copy for win contacts than for loss contacts.

In this play

1. Structure the sequence before writing any copy
2. Write the initial invitation
3. Write the follow-up touches
4. Frame the incentive correctly
5. Adjust the sequence for win contacts

Step 1: Structure the Sequence Before Writing Any Copy

Build the full sequence before sending a single email. A standard structure runs four touches over roughly two weeks: an initial invitation, a follow-up three to four days later, a second follow-up five to seven days after that, and a short final note before closing out non-responders. Most buyers who eventually agree to an interview respond to the second or third touch, not the first, so deciding the full cadence in advance matters more than any single email's wording.

SEQUENCE TIMING

- Touch 1 (Day 0): Initial invitation
- Touch 2 (Day 4): First follow-up
- Touch 3 (Day 10): Second follow-up
- Touch 4 (Day 16): Final short note, then close out

Spacing the touches three to seven days apart, rather than daily, matters for two reasons beyond politeness. A buyer who's genuinely too busy to respond in week one often has room in week two, so compressing the sequence into a few days mostly reaches the same inattentive window twice. It also protects deliverability: sending from the same address to hundreds of former buyers in a short burst reads as bulk mail to spam filters in a way the same volume spread across two weeks doesn't.

Step 2: Write the Initial Invitation

Open by naming the purpose of the outreach and making clear it's separate from the sales and account relationship, before asking for anything. The template below is written for loss contacts, the larger and lower-response-rate share of the list; Step 5 covers how to adjust it for win contacts.

INITIAL INVITATION – LOSS CONTACT

Subject: Quick research question about your [Vendor] evaluation

Hi [First Name],

I'm reaching out as part of a structured research effort to understand how companies like [Buyer Company] evaluate [Vendor Category] vendors. This outreach is separate from the sales and account team, and the goal is to understand what actually shaped the decision not to move forward with [Vendor].

Would you have 30 minutes in the next two weeks for a conversation? Everything you share stays anonymous in any findings, and as a thank-you for your time, I'll send a \$50 gift card after the call.

Here's my scheduling link if you'd like to pick a time: [Scheduling Link]

Thanks,
[Researcher Name]
[Research Program Name]

Step 3: Write the Follow-Up Touches

Vary the framing across touches rather than resending the same email. Most non-responses come from inattention, not refusal, so a follow-up with a slightly different angle recovers people the first email missed instead of just repeating what they already skimmed past once.

FOLLOW-UP TOUCHES

Follow-up 1 (Day 4)

Subject: Re: Quick research question about your [Vendor] evaluation

Hi [First Name], following up in case this got buried. I just need 30 minutes to hear your take on the evaluation, and I'm happy to work around your schedule: [Scheduling Link]

Follow-up 2 (Day 10)

Subject: Last note on this

Hi [First Name], I'll leave this as my last note. If a conversation makes sense sometime in the next week, grab a time here: [Scheduling Link]. Either way, thanks for considering it.

Step 4: Frame the Incentive Correctly

Mention the incentive in the initial invitation and confirm it again once a buyer schedules, not during the interview itself. Frame it as compensation for time given, not payment for a specific answer: "as a thank-you for your time" works; "in exchange for your feedback" doesn't. A modest gift card, sized to read as genuine appreciation rather than a fee for services, is the standard. Skipping the incentive entirely filters the eventual sample toward buyers with the most extreme reactions in either direction, since buyers with more moderate, and often more representative, views have the least reason to give up unpaid time.

REALITY CHECK

No matter how neutrally the outreach is worded, buyers with the sharpest feedback are still less likely to engage with a sequence from a team affiliated with the vendor than with a request from a neutral third party. Wording helps at the margin; it can't close the gap between what a buyer says to someone connected to the company and what they'll say to someone with no stake in the outcome. See [why third-party neutrality changes what buyers will say](#).

Step 5: Adjust the Sequence for Win Contacts

Win contacts convert at roughly double the rate of loss contacts, so the same sequence structure can run with a lighter touch, often two emails instead of four. Adjust the initial invitation's framing as well as its length: a win contact is being asked to help explain what worked, not revisit a decision they may have mixed feelings about, so the email can open with a more direct ask for their perspective on why they chose to move forward, without the same care taken in the loss sequence to establish separation from the sales relationship before anything else.

INITIAL INVITATION – WIN CONTACT

Subject: Quick question about choosing [Vendor]

Hi [First Name],

I'm reaching out as part of a short research effort on how companies like [Buyer Company] evaluate and choose [Vendor Category] vendors. I'd love to hear your perspective on what tipped the decision in [Vendor]'s favor, and whether anything nearly changed that outcome along the way.

30 minutes, entirely anonymous in any findings, and a \$50 gift card as a thank-you for your time: [Scheduling Link]

Thanks,
[Researcher Name]
[Research Program Name]

A single follow-up five to seven days later, reusing the same shorter framing as Follow-up 1 above, is usually enough to close out a win-contact sequence. The higher baseline response rate means a third or fourth touch adds outreach volume without adding much to the completed-

interview count.

How to Design a Win/Loss Interview Questionnaire

A win/loss interview questionnaire moves through five stages across roughly 35 to 45 minutes: warm-up, buying process and evaluation criteria, competitive landscape, product and messaging perception, and outcome drivers. Open each stage with broad, open-ended questions before narrowing into specific ones, and build separate win and loss variants for the final stage.

In this play

1. Structure the five stages
2. Balance open and closed questions within each stage
3. Build the loss-interview questionnaire
4. Build the win-interview questionnaire
5. Leave room for the unexpected

Step 1: Structure the Five Stages

Build the questionnaire around five stages, each covering one job the interview has to do: warm-up and context, buying process and evaluation criteria, competitive landscape, product and messaging perception, and outcome drivers. Sequence them in that order. A buyer who starts with a broad, low-stakes question about their role and the overall process is far more likely to volunteer detail later than one who opens on a narrow, specific question about price or a competitor.

STAGE TIMING

Stage 1: Warm-Up and Context	(3-5 min)
Stage 2: Buying Process & Criteria	(8-10 min)
Stage 3: Competitive Landscape	(8-10 min)
Stage 4: Product & Messaging Perception	(5-7 min)
Stage 5: Outcome Drivers	(7-10 min)

Each stage builds on what the last one surfaced. A stakeholder named in Stage 2 often resurfaces in Stage 5 as the person whose objection actually decided the outcome, and having that name already on the table makes the final stage's question land with specifics instead of a generic summary.

The five-stage skeleton stays fixed across a program, but the specific questions within Stage 3 and Stage 4 should be built around whatever GTM question the program was scoped to answer, not applied as an identical generic template on every engagement. A program scoped to investigate a messaging gap weights Stage 4 more heavily; one scoped around a specific competitive threat weights Stage 3. The stage order and timing stay constant so interviews remain comparable to each other; the specific questions inside two of the five stages flex to the program's actual objective.

Step 2: Balance Open and Closed Questions Within Each Stage

Open every stage with an open-ended question that lets the buyer set the frame in their own words, then follow with two or three closed questions that fill in specifics the open question didn't cover on its own. A stage that starts closed forces the buyer into the interviewer's categories before they've had a chance to describe what actually happened.

A stage built around a positioning assumption, "was price a factor," asked before any open question, primes the buyer to confirm it rather than volunteer whatever the real driver was. The same stage opened with "walk me through how the decision got made" and only then narrowing into "was budget or pricing part of that conversation" produces a materially different account, because the buyer's own framing comes first.

Step 3: Build the Loss-Interview Questionnaire

LOSS-INTERVIEW QUESTIONNAIRE

STAGE 1: WARM-UP AND CONTEXT

- Walk me through your role and how you were involved in this evaluation.
- Before we get into specifics, walk me through how this decision got made, from the first conversation to the final call.

STAGE 2: BUYING PROCESS & EVALUATION CRITERIA

- Who else was involved in the decision, and what did each of them care about most?
- What were the top three criteria you were evaluating vendors against?
- (Closed) Was there a specific budget or approval process this had to go through?

STAGE 3: COMPETITIVE LANDSCAPE

- Which other vendors did you seriously evaluate?
- What stood out about [Competitor] compared to us?
- (Closed) Did pricing come up as a specific point of comparison between vendors?

STAGE 4: PRODUCT & MESSAGING PERCEPTION

- Going into the evaluation, what did you understand about what we do differently from the other vendors you looked at?
- Was there a moment where a demo or a piece of messaging didn't land the way you expected?

STAGE 5: OUTCOME DRIVERS (LOSS)

- What was the specific moment or conversation that tipped the decision away from us?
- If that concern had been addressed, do you think the outcome would have been different?
- (Closed) Is there anything you wish had come up during the evaluation that never did?

Step 4: Build the Win-Interview Questionnaire

Stages 1 through 4 stay structurally the same for a win interview; only the framing shifts slightly since the buyer is describing a decision that went the vendor's way rather than a competitor's. Stage 5 changes the most, since a win interview is testing what tipped the decision in the vendor's favor and whether anything nearly changed that outcome, rather than identifying what went wrong.

WIN-INTERVIEW QUESTIONNAIRE – STAGE 5

STAGE 5: OUTCOME DRIVERS (WIN)

- What tipped the decision in our favor?
- Was there a point in the process where you seriously considered a competitor, or nearly walked away?
- (Closed) Did anything about the evaluation almost change the outcome, even though it didn't in the end?

Step 5: Leave Room for the Unexpected

A fixed question list keeps every interview consistent enough to compare against the others, but it shouldn't run so rigidly that an interviewer moves to the next scripted question the moment a buyer raises something unprompted. If a buyer volunteers a detail that doesn't map cleanly to any of the five stages, a stalled internal approval, an unrelated vendor relationship, a reorg that changed who owned the decision, follow that thread before returning to the script.

That follow-up is frequently where the most useful detail in the entire interview lives. A questionnaire that covers every topic exhaustively but never deviates from its own structure produces consistent, comparable, and often shallow answers. One built around a fixed core with room to follow an unexpected thread produces the depth that turns a passing comment into a finding worth bringing back to the full interview set.

How to Conduct a Win/Loss Interview

Conducting a win/loss interview starts with a clear statement of neutrality and purpose, then holds to the five-stage question set across roughly 35 to 45 minutes. Follow up on vague or surprising answers instead of moving straight to the next scripted question, and close with a specific next step rather than an open-ended thank-you.

In this play

1. Open the call to establish neutrality immediately
2. Pace the conversation across the full question set
3. Recognize and respond to a reluctant buyer
4. Maintain neutral-interviewer technique throughout
5. Close the call and capture the follow-up thread

Step 1: Open the Call to Establish Neutrality Immediately

Restate, in the first minute of the call, what the outreach email already said: the conversation is separate from the sales and account relationship, everything shared stays anonymous in any findings, and the purpose is to understand what actually happened, not to relitigate the decision. A buyer who hears this again at the start of the call, rather than assuming it from the email alone, settles into the conversation faster.

CALL OPENING — SCRIPT

"Thanks for making time. Quick context before we start: this is separate from the sales team, nothing you say here gets attributed to you by name in any findings, and I'm just trying to understand how the decision actually played out. There's no wrong answer here, I'm here to listen, not to sell you on anything or defend a decision."

Step 2: Pace the Conversation Across the Full Question Set

Hold to the stage timing set in the questionnaire so every interview in the set covers the same ground, which is what makes cross-interview comparison possible later. Watch the clock at each stage transition rather than at the end of the call, since a slow first stage is recoverable if caught at minute ten, but not if it's only noticed at minute thirty with two stages still untouched.

STAGE TIMING	
Stage 1: Warm-Up and Context	(3-5 min)
Stage 2: Buying Process & Criteria	(8-10 min)
Stage 3: Competitive Landscape	(8-10 min)
Stage 4: Product & Messaging Perception	(5-7 min)
Stage 5: Outcome Drivers	(7-10 min)

If a stage is running long because a buyer is offering genuinely specific, unprompted detail, let it run and trim time from a later stage rather than cutting the buyer off mid-thought. The specificity in Stage 5, where the actual outcome driver usually surfaces, is worth protecting more than strict adherence to the clock.

Step 3: Recognize and Respond to a Reluctant Buyer

A reluctant buyer rarely refuses to answer outright. More often, they give a short, generic answer, "price was a factor," "the timing wasn't right," and stop there. Treat a generic answer as an invitation to follow up, not a complete response. Ask what specifically about the pricing structure raised a concern, or what "timing" actually meant in practice, rather than logging the generic version and moving to the next question.

Consider a buyer who answers "price was a factor" when asked why the deal didn't move forward. A follow-up like "What about the pricing structure specifically raised a concern, was it the number itself or how it was structured" often surfaces something more specific: a per-seat model that didn't fit an unusual team structure, or a renewal term that read as inflexible. The generic answer was true. It just wasn't the whole answer, and the follow-up is what gets the rest of it.

If a buyer stays guarded through a full stage despite follow-up attempts, move on rather than pressing further. A single under-developed stage in one interview is a manageable gap in a 20-to-30-interview set; a buyer who ends the call early because they felt pushed is a worse outcome for the program than one thin answer.

Step 4: Maintain Neutral-Interviewer Technique Throughout

Ask every question as genuine curiosity about what happened, not as a lead-in to a specific answer the interviewer expects. "Was price the deciding factor?" invites a yes-or-no confirmation. "What role did price play in the decision, if any?" leaves room for the buyer to describe something more specific, or to say it wasn't really about price at all.

Follow up on the vague or surprising answer in the moment rather than filing it away to ask about later. A buyer who mentions an unnamed stakeholder in passing during Stage 2 is easiest to draw out right then, while the detail is fresh in their own account; circling back to it in Stage 5 after they've moved on mentally rarely recovers the same specificity.

REALITY CHECK

Interview technique can get a buyer to open up more than a poor conversation would, but it can't close the gap a relationship creates. A buyer who might sell to, or buy from, the interviewer's company again is still managing that relationship in real time, regardless of how skilled the questions are. Neutrality is a property of who's asking, not just how well they ask. See [why buyer candor depends on who is asking the question](#).

Step 5: Close the Call and Capture the Follow-Up Thread

End with a specific next step rather than a generic thank-you. Confirm the incentive is on its way, reconfirm anonymity, and ask permission for one short follow-up email if a detail from the call needs clarifying during analysis. A buyer who's agreed to that in advance is far more likely to respond to a two-line clarifying question two weeks later than one who's hearing from the program again out of nowhere.

CALL CLOSING – SCRIPT

"That's everything I need, thank you for the time. Your gift card will go out within a few days. Everything we discussed stays anonymous in any findings. If something comes up during analysis where I need one quick clarification, would it be alright to send a short follow-up email?"

Log the call's specific unresolved threads immediately after hanging up, while they're still easy to distinguish from the next interview's. A note like "unclear which stakeholder raised the compliance objection, follow up if this theme recurs" is far more useful during analysis than trying to reconstruct that ambiguity from a transcript days or weeks later.

How to Code Win/Loss Interview Responses

Coding a win/loss interview means tagging raw notes against a consistent taxonomy immediately after each call, not in a batch once the full interview set is complete. Build that taxonomy from the same categories the questionnaire covers, tag every interview against it the same way, and track code frequency as interviews accumulate so a recurring theme becomes visible as soon as it crosses the pattern threshold.

In this play

1. Code immediately after each interview, not in batches
2. Build the coding taxonomy from the questionnaire's categories
3. Tag consistently across the full interview set
4. Track code frequency as interviews accumulate
5. Choose tooling that fits the program's size

Step 1: Code Immediately After Each Interview, Not in Batches

Code each interview within a day of conducting it, while the specific phrasing and context are still fresh. The interviewer's own impression of what a buyer meant by a vague comment degrades the same way a buyer's memory of a deal degrades over time, so a note like "mentioned onboarding, seemed like a real concern" written the same day carries more usable signal than the same note reconstructed from a transcript two weeks later.

Coding in batches after five or ten interviews also invites a specific failure: early impressions blur together, and a theme that was actually distinct across three different buyers starts to read as one repeated comment because the coder no longer remembers which buyer said what.

Build the coding pass directly into the post-call routine established in the interview play: capture the unresolved-thread notes immediately after hanging up, then apply codes to the full call notes before starting the next scheduled interview. Treating coding as a same-day task rather than an end-of-week catch-up keeps the taxonomy from Step 2 accurate to what buyers actually said, not to a compressed summary of it.

Step 2: Build the Coding Taxonomy from the Questionnaire's Categories

Use the questionnaire's five stages as the top-level structure, so a code maps cleanly back to the question that produced it, but build the specific codes within each stage from what the interviews actually surface rather than adopting a fixed list. The categories below the stage level are illustrative, one plausible set for one hypothetical program, not a standard vocabulary to apply unchanged across every engagement.

ILLUSTRATIVE CODING TAXONOMY	
Buying Process & Criteria:	Committee-Stakeholder-Objection, Budget-Approval-Friction
Competitive Landscape:	Competitor-Pricing-Comparison, Feature-Gap-Named
Product & Messaging:	Demo-Objection, Onboarding-Concern
Outcome Drivers:	Price-Cited, Champion-Lost-Internal-Argument
Other:	[open tag for anything outside the categories above]

Where a company already codes losses in its CRM with a reasonably specific set of categories, align the interview taxonomy to those categories rather than building an entirely separate vocabulary. Matching, say, an existing "Pricing/Budget" or "Competitive Loss" field lets the coded interview findings sit directly next to the CRM's own loss-reason distribution once analysis is complete, which is often the clearest way to show where the two diverge, not just that they do. Build new codes only where the interviews surface a theme the CRM's existing categories don't capture at all.

Keep the "Other" tag genuinely open rather than forcing every comment into an existing code. A comment that gets stretched to fit the nearest available tag loses the specificity that made it worth coding in the first place, and a growing "Other" list across several interviews is itself a signal that the taxonomy is missing a real category.

Step 3: Tag Consistently Across the Full Interview Set

Apply the same code to the same underlying concern every time it appears, even when different buyers describe it in different words. A buyer who says "the pricing felt like it had hidden levers" and another who says "we couldn't tell what we'd actually end up paying" are both describing pricing opacity, and both should carry the same code, not two separate ones that happen to look different on paper.

Re-read the taxonomy against the first three or four coded interviews before continuing through the rest of the set. Codes that seemed distinct in isolation often turn out to describe the same underlying concern once a few interviews are sitting side by side, and it's easier to consolidate two codes into one early than to recode a full set after the fact.

A practical example: interview three gets coded "Pricing-Opacity" and interview five gets coded separately as "Hidden-Fees-Concern." Reviewing both notes together after interview five shows they're describing the same underlying issue in different words, not two distinct pricing problems. Merging them into a single "Pricing-Opacity" code before interview six keeps the frequency count in Step 4 accurate; leaving them split would understate how often the theme is actually recurring, and could keep a real pattern below the three-mention threshold on a technicality.

Step 4: Track Code Frequency as Interviews Accumulate

Maintain a running count of how many interviews carry each code, updated after every coding session rather than reconstructed at the end. A code that reaches three to five independent mentions, unprompted and in the buyer's own words, crosses the threshold that separates a pattern from an anecdote.

CODE FREQUENCY TRACKING

Code: Champion-Lost-Internal-Argument
Interviews: #4, #9, #15, #22
Frequency: 4 of 20 loss interviews (20%)

Watching this count grow in real time, rather than discovering it during a single end-of-cycle analysis pass, surfaces an emerging pattern early enough to probe it more deliberately in the remaining interviews still on the calendar.

Step 5: Choose Tooling That Fits the Program's Size

A standard 20-to-30 interview program doesn't need a dedicated qualitative-analysis platform. A shared spreadsheet with one row per interview and one column per code, marked with a simple count or checkmark, is enough to support the frequency tracking in Step 4 and stays easy for more than one person to work from during analysis.

INTERVIEW ID	TYPE	ONBOARDING-CONCERN	PRICE-CITED	FEATURE-GAP-NAMED
#4	Loss	-	X	-
#9	Loss	X	-	X
#12	Loss	X	-	-
#15	Loss	X	X	-

Reserve a specialized coding tool for programs running well beyond 30 interviews or maintaining a continuously updated interview set across multiple cycles, where a spreadsheet's manual upkeep starts to outweigh its simplicity.

How to Distinguish a Pattern from an Anecdote in Win/Loss Data

A finding becomes a pattern once it appears independently across at least three to five separate buyer interviews, described in the buyer's own words rather than prompted by the interviewer. A single mention, however specific, is an anecdote until it recurs, and recurrence across losses carries more diagnostic weight than recurrence across wins.

In this play

1. What counts as a pattern
2. Why the threshold sits at three to five
3. Weighting losses over wins when judging a pattern
4. A worked example, from passing comment to finding
5. Guard against analyst bias on borderline calls

What Counts as a Pattern

A pattern is a theme that recurs independently across multiple interviews, in the buyer's own language, without the interviewer having raised it first. A code applied because a buyer answered a direct question about pricing doesn't carry the same weight as a code applied because a buyer volunteered a pricing concern unprompted while answering a different question entirely. The second case is stronger evidence, since nothing about the question primed that specific answer.

Recurrence also has to be genuinely independent. Three mentions from buyers who work at the same company, or who were part of the same buying committee describing the same internal conversation, count as one account told three times, not three separate data points. A pattern requires the same underlying theme to surface across interviews with no connection to each other.

Why the Threshold Sits at Three to Five

Below three independent mentions, a theme is still plausibly one buyer's idiosyncratic experience, something specific to their internal politics or their particular evaluation rather than something a broader GTM strategy should respond to. At three to five mentions, drawn from a 20-to-30 interview set, the theme has surfaced across roughly a tenth to a quarter of the total sample without any two buyers having spoken to each other, which is difficult to explain as coincidence.

The threshold isn't a hard cutoff so much as a floor. A theme that reaches five or six mentions and keeps recurring as later interviews come in is a stronger pattern than one that stalls at three and never resurfaces, and both are treated differently in the report even though both cleared the minimum bar.

Weighting Losses Over Wins When Judging a Pattern

Judge a pattern against both interview types, not just the set it appeared in. A concern that surfaces repeatedly in loss interviews but never once in the win interviews points to a real, specific gap for the segment of buyers who chose not to move forward. The same concern showing up in win interviews too, described as something the buyer noticed but ultimately tolerated, points to a lower-urgency issue rather than a dealbreaker.

This comparison is why a program's interview count is weighted roughly two loss interviews for every win interview in the first place. Losses tend to produce a cleaner, more specific signal about what actually went wrong, and having enough win interviews to check a loss-side pattern against is what turns "buyers mentioned this" into "buyers mentioned this, and it was decisive for the ones who walked away, but not for the ones who stayed."

A Worked Example, from Passing Comment to Finding

By interview six, a buyer mentions in passing that onboarding felt more complicated than expected. Coded and logged, but with only one mention, it's an anecdote, indistinguishable from something specific to that buyer's internal IT environment.

By interview twelve, two more buyers have independently raised the same concern, in different language and without any prompting from the interviewer, one describing "a steep setup curve" and another calling it "more hand-holding than we expected to need." That's three independent mentions, clearing the pattern threshold. Checking the ten win interviews turns up two buyers who mention onboarding as well, both describing it as something they noticed but worked through rather than something that nearly changed their decision. The comparison across both sets tells the full story: onboarding complexity is real and worth a product recommendation, but it's a friction point buyers tolerate when they've already decided to move forward, not a reason deals are being lost outright.

Guard Against Analyst Bias on Borderline Calls

A theme sitting right at the three-mention threshold is where analyst judgment matters most, and where it's easiest for that judgment to run in a predictable direction. An analyst weighing whether a borderline theme counts as a pattern is making a genuine judgment call regardless of who's running the analysis, but that judgment isn't immune to the same pressures that shape what gets logged elsewhere in a GTM organization.

A useful check on a borderline call: would this theme read as a clear pattern if it implicated a different function than the one doing the analysis. If a three-mention pricing theme would be treated as decisive but a three-mention theme about the sales process is being waved off as coincidence, that asymmetry is worth naming out loud before the finding gets written up either way.

A second, simpler check works alongside the first: have someone outside the function under scrutiny review the borderline call before it's finalized, even briefly. A five-minute second opinion from someone with no stake in how the theme reads is often enough to catch an asymmetric threshold before it makes it into a report that a leadership team will treat as settled.

REALITY CHECK

Judging whether a three-mention theme counts as a pattern is a judgment call regardless of who runs the analysis, but an analyst reviewing losses that implicate their own team has reason to read it as noise, the same filter that shapes what gets logged as a loss reason in the first place. This threshold doesn't remove that judgment call. It only sets the bar it's applied against. See [why internal teams interpret win/loss data through an existing filter](#).

How to Deliver a Win/Loss Readout

Deliver a win/loss readout as a live, cross-functional meeting, not an emailed report, with every function that needs to act on findings in the room at the same time. Structure the deck around synthesized findings organized by theme, not individual deal summaries, and close the meeting with a specific action item assigned to each function before anyone leaves.

In this play

1. Structure the deck around synthesized findings, not individual deals
2. Decide who needs to be in the room
3. Run the meeting to end in action items, not just a presentation
4. Close each finding with a specific ask per function
5. Follow up within two to three weeks

Step 1: Structure the Deck Around Synthesized Findings, Not Individual Deals

Build the deck around a small number of synthesized findings, each supported by evidence across multiple interviews, not a folder of individual deal write-ups. A reader who has to page through twenty separate deal summaries to notice a pattern hasn't been given a finding; they've been given raw material and asked to do the analysis themselves, in the room, in real time.

READOUT DECK STRUCTURE

1. Research Scope (deal set, interview count, time window)
2. Methodology (interview structure, coding approach)
3. Finding 1 ([Theme]): supporting evidence, recommendation
4. Finding 2 ([Theme]): supporting evidence, recommendation
5. Finding 3 ([Theme]): supporting evidence, recommendation
6. Cross-Functional Action Summary

Cap the deck at three to five findings. A readout that tries to present every theme that crossed the pattern threshold buries the most important ones in a long list, and a leadership team can hold three to five specific findings in their heads well enough to act on them; a longer list mostly gets skimmed.

Support each finding with the specific evidence that earned it pattern status, not the full interview transcript it came from. A brief note on how many interviews carried the theme, paired with one or two anonymized quotes in the buyer's own words, gives the room enough to trust the finding without turning the readout into a re-analysis session. The full coded interview set stays available as backup material for anyone who wants to dig deeper after the meeting, not something walked through live.

Step 2: Decide Who Needs to Be in the Room

Invite the functions that can act on what the findings actually cover, plus one executive sponsor with the authority to hold each of them accountable for the commitment they leave with.

ROLE	WHY THEY'RE IN THE ROOM
Research owner	Presents findings, facilitates action-item capture
Executive sponsor	Holds each function accountable for its commitment
Sales leadership	Acts on competitive and process findings
Product marketing	Acts on messaging and positioning findings
Competitive intelligence	Acts on competitive findings
Product or engineering	Included only when a finding bears on the roadmap

A readout without a senior sponsor in the room tends to produce the same weak follow-through as an emailed report, since nobody in attendance has the standing to ask, three weeks later, why a committed action hasn't moved.

Step 3: Run the Meeting to End in Action Items, Not Just a Presentation

Budget the meeting to end with roughly a quarter of its time spent capturing commitments, not just presenting findings. A 60-minute readout structured entirely around walking through the deck leaves no room for the step that actually makes the research worth the time spent producing it.

60-MINUTE READOUT AGENDA

0:00-0:05	Scope and methodology recap
0:05-0:35	Walk through findings, theme by theme
0:35-0:50	Open discussion: does this match what functions are already seeing, where does it diverge
0:50-0:60	Capture a specific action item per function, read back out loud before the meeting ends

Read each captured action item back to the room before closing. A commitment stated out loud and confirmed in the room is harder to quietly deprioritize afterward than one buried in meeting notes nobody reviews again.

Step 4: Close Each Finding With a Specific Ask Per Function

Pair every finding presented in Step 1 with a specific ask for each function it touches, stated in that function's own terms rather than as a general summary. "Sales should be aware of this competitive dynamic" isn't an action item. "Sales enablement updates the battlecard entry for [Competitor] with the implementation-speed objection by [date]" is.

Building this per-function translation for every finding is its own task with its own template; the full method for mapping a finding type to the right action per function is covered in [How to Translate Win/Loss Findings into Function-Specific Actions](#).

Step 5: Follow Up Within Two to Three Weeks

Schedule a short follow-up check, two to three weeks after the readout, on the specific commitments captured in Step 3. Findings that sit waiting for a natural check-in lose urgency the same way any other finding does with time, and a program's real value is much easier to demonstrate against a list of committed actions checked at a fixed interval than against a vague sense of whether the readout "went well."

Keep the follow-up short: a status against each committed action, not a second full readout. The goal is confirming movement happened, not re-presenting the findings that prompted it.

An action item still unstated at the two-to-three-week mark is worth a direct conversation with its owner rather than a quiet extension. A commitment made out loud in a room with an executive sponsor present and then dropped without explanation is a signal about how seriously the program is being treated organizationally, not just a scheduling slip, and it's worth surfacing before the pattern repeats across the next cycle's readout.

How to Translate Win/Loss Findings into Function-Specific Actions

Translate win/loss findings into function-specific actions with a matrix that maps each finding type to what Product, Product Marketing, Sales, and Competitive Intelligence should each do with it, rather than one interpretation applied uniformly across every team. The same underlying finding produces a different action for each function, and the matrix makes that translation explicit instead of leaving it to each team to infer.

In this play

1. Build the finding-type-to-function matrix
2. Translate a competitive finding across functions
3. Translate a messaging finding across functions
4. Translate a process friction finding across functions
5. Assign an owner and deadline to every action

Step 1: Build the Finding-Type-to-Function Matrix

Lay out finding types as rows and the four functions as columns, and fill each cell with the specific action that finding type calls for in that function, or leave it blank where a finding type genuinely doesn't touch that function at all. A blank cell is a real answer, not a gap to fill for the sake of completeness.

FINDING TYPE	PRODUCT	PRODUCT MARKETING	SALES	COMPETITIVE INTEL
Competitive Loss	Roadmap signal	Positioning gap	Objection handling	Battlecard update
Messaging Gap	-	Message to test	Talk track update	-
Pricing Concern	-	Value-msg reframe	Pricing objection	Competitive intel
Process/Sales Friction	-	-	Coaching signal	-
Onboarding/Product Gap	Roadmap item	-	-	-

Use this matrix as the starting structure for the readout's Cross-Functional Action Summary, filling in the specific finding and action text for the current cycle rather than leaving the generic labels above in the final deck.

Treat the matrix as a living reference across cycles, not a document rebuilt from scratch each time. A finding type that recurs, competitive losses, for instance, tends to call for a similar shape of action each cycle, even as the specific competitor or detail changes. Keeping one matrix updated cycle over cycle also makes it easy to spot a function that consistently receives action items but never reports back on them, which is a distribution problem worth raising on its own.

Step 2: Translate a Competitive Finding Across Functions

Take a finding like a competitor consistently winning late-stage evaluations on implementation speed rather than price, across three consecutive cycles. Each function's action comes from the same finding but addresses a different part of the business:

COMPETITIVE FINDING – ACTION BY FUNCTION

Product:	Evaluate implementation speed as a roadmap priority, not just a messaging fix
Product Marketing:	Build proactive messaging on implementation speed before it's raised as an objection
Sales:	Equip reps with a specific objection-handling response for this scenario
Competitive Intelligence:	Update the battlecard entry for this competitor with the implementation-speed angle

None of these four actions restates the finding. Each one is a specific, ownable task that follows from it.

Step 3: Translate a Messaging Finding Across Functions

Take a finding like a positioning line that tests well in demos but consistently fails to survive an internal buying committee discussion after the vendor leaves the room. This finding touches fewer functions than the competitive example above, and the matrix should reflect that rather than manufacturing an action for every column.

MESSAGING FINDING – ACTION BY FUNCTION

Product Marketing:	Revise the message, or build a leave-behind document that survives being forwarded internally without the salesperson there to explain it
Sales:	Coach reps to arm the internal champion with a written version of the pitch before the committee meets
Product:	No action, unless the underlying gap traces to an actual capability question
Competitive Intelligence:	No action

Step 4: Translate a Process Friction Finding Across Functions

Take a finding like onboarding complexity that shows up repeatedly in loss interviews but only as a tolerated friction point in win interviews. The comparison across win and loss interviews, covered in *How to Distinguish a Pattern from an Anecdote in Win/Loss Data*, is what determines whether this finding is urgent or just worth noting.

PROCESS FRICTION FINDING – ACTION BY FUNCTION	
Product:	Prioritize onboarding simplification, scoped against the specific step buyers named as confusing
Sales:	Set clearer onboarding expectations during the evaluation, so the friction is anticipated rather than a surprise
Product Marketing:	No action
Competitive Intelligence:	No action

A finding that only maps to one or two functions is still worth presenting in the readout. Forcing every finding into four columns of action items dilutes the ones that genuinely span the whole GTM organization, and a blank cell communicates something real: this finding doesn't call for a product change, and pretending otherwise wastes the room's attention on a manufactured action nobody will actually follow through on.

Step 5: Assign an Owner and Deadline to Every Action

Add an owner and a deadline to every non-blank cell before the readout ends, not as a follow-up task decided afterward. An action item with no named owner defaults to being everyone's responsibility, which in practice means no one's.

OWNER AND DEADLINE – EXAMPLE

Finding: Competitor winning on implementation speed
Action: Update battlecard entry
Function: Competitive Intelligence
Owner: [Name]
Deadline: [Date, typically within 2 weeks of the readout]

Track these owner-and-deadline rows against the follow-up check described in How to Deliver a Win/Loss Readout, so the two-to-three-week follow-up has a specific list to check status against rather than a general question of whether the readout "led anywhere."

Keep the completed matrix from each cycle rather than overwriting it with the next one. A record of which finding types consistently turn into completed actions, and which consistently stall regardless of who owns them or what deadline was set, is itself a useful input the next time the program's cadence or distribution model gets reviewed.

How to Set and Adjust Your Win/Loss Program Cadence

Run a win/loss program on a semi-annual cadence by default. Six months gives the team time to execute on the prior cycle's action items and captures a real window of change in the market, without running so often that findings turn into noise. Adjust off that default only for a specific, dynamic shift, a pipeline problem or competitive change, not on a fixed shorter schedule.

In this play

1. Set the default cadence at semi-annual
2. Confirm deal volume can support it
3. Build a findings review into every cycle
4. Know which findings decay fast and which don't
5. Use an off-cycle refresh for anything more urgent

Step 1: Set the Default Cadence at Semi-Annual

Six months is the right rhythm for most programs. It gives functions enough time to actually execute on the prior cycle's committed actions before the next round of interviews starts, and it's long enough to capture meaningful movement in buyer sentiment, competitive dynamics, and market conditions, the kind of change worth adjusting GTM strategy over. That's the bar for this cadence: findings that move strategy, not a running tally of every minor shift.

Quarterly is too frequent for most programs to get real value from. There isn't enough runway between cycles for committed actions to land before the next readout piles more on top, and the findings start reading as noise rather than signal. A program on a quarterly rhythm ends up presenting the same handful of themes cycle after cycle, mostly unchanged, because six months' worth of real movement hasn't had time to accumulate in three. Treat semi-annual as the fixed baseline once set, not something reopened at the start of every cycle.

Step 2: Confirm Deal Volume Can Support It

Before locking in the cadence, confirm the segment actually produces enough closed deals to support it. Use the same list-size math from How to Build a Win/Loss Interview Target List: roughly 400 loss-eligible and 100 win-eligible contacts within the 30-to-180-day research window.

CADENCE VS. DEAL VOLUME

Segment produces ~400 losses + ~100 wins within the window

-> Semi-annual cadence is well-supported

Segment produces meaningfully less

-> Widen the segment or extend toward the full 180-day edge of the window before locking in a fixed cadence

A segment that comes up short here is the same situation the target-list play's own worked example resolves: extend the window or widen the segment, rather than lowering the interview count or forcing a cadence the deal volume can't actually support.

Step 3: Build a Findings Review into Every Cycle

Open every readout with a short review of the prior cycle's findings before presenting anything new, checking which ones still hold up against the current interview set.

PRIOR-CYCLE FINDINGS REVIEW

For each finding presented in the prior cycle's readout:

[] Still holding up against this cycle's interview set?

[] Has anything in the market changed enough to retest it?

[] Retire it, keep it as confirmed, or flag it as needing current-cycle confirmation before anyone acts on it again

A program that only adds new findings each cycle without retiring stale ones ends up with a battlecard or messaging framework built on a mix of current signal and outdated assumption, with no way to tell which is which without checking.

Step 4: Know Which Findings Decay Fast and Which Don't

Not every finding needs the same review scrutiny. Some hold their shape for years; others are stale within a couple of cycles.

FINDING TYPE	DECAY SPEED	EXAMPLE
Buying committee structure	Slow, years	Who typically holds veto power in the decision
Process friction and risk signals	Slow to medium	Recurring budget-approval friction patterns
Competitor perception	Fast, months	How a named competitor reads after a product update
Messaging resonance	Fast, months	Whether a specific positioning line still lands
Pricing perception	Fast, months	How pricing reads against a shifting market

Weight the Step 3 review toward the fast-decaying rows every cycle, and treat the slow-decaying ones as needing a lighter check, confirmed rather than fully retested, unless something specific suggests otherwise.

Step 5: Use an Off-Cycle Refresh for Anything More Urgent

A dynamic market shift, a pipeline problem, or a competitive change that needs an answer faster than the next scheduled cycle allows is a reason to run a scoped off-cycle refresh, not a reason to move the whole program to quarterly.

OFF-CYCLE REFRESH TRIGGER

Trigger a scoped off-cycle refresh when a specific, already-felt signal moves faster than the next scheduled cycle would catch it: a win rate drop in a specific segment, a competitor newly appearing in evaluations where it never used to, a major product launch, or a pricing change under active consideration.

The trigger has to be specific, not a general sense that "it's been a while." A refresh scoped around "our Enterprise win rate dropped eight points this quarter" is answerable; a refresh scoped around vague unease isn't, and it usually signals the program needs a clearer objective

more than it needs a faster clock.

A practical example: a program on a semi-annual cadence is three months into its current cycle when a new competitor starts appearing in evaluations for the first time. Waiting the remaining three months means running interviews against a competitive landscape that's already stale by the time findings come back. An off-cycle refresh, scoped narrowly to the competitive-perception question rather than a full program restart, gets a usable answer to leadership while the signal is still current, then folds back into the regular semi-annual cadence at the next scheduled cycle. This is the exception the default is built around, not a reason to abandon it.

How to Measure Whether a Win/Loss Program Is Working

Measure whether a win/loss program is working by tracking how many specific decisions its findings actually changed, not interview counts or report cadence. A program that hits its interview target every quarter has proven the interviews got scheduled; it hasn't shown that a single decision got made differently because of them.

In this play

1. The activity trap: why interview counts aren't a success metric
2. What to track instead: decisions changed, not reports delivered
3. A simple scorecard template
4. Reviewing the scorecard at each cycle
5. When an empty scorecard is the real finding

The Activity Trap: Why Interview Counts Aren't a Success Metric

Programs commonly fall into an activity-metrics trap: hitting a target interview count in a quarter, or getting a report out the door on schedule, starts to feel like success, purely because everyone involved is stretched thin and getting interviews scheduled at all felt like an accomplishment.

Activity is not the same thing as impact. A team that completed 25 interviews this quarter has proven the interviews got scheduled. It hasn't shown that any of the 25 changed a decision anyone actually made. This trap is easy to fall into precisely because interview counts and report dates are simple to track, while the thing that actually matters, whether a specific decision moved, takes deliberate tracking that doesn't happen automatically.

An internal team stretched across other responsibilities can fall into this trap without anyone noticing it happening. Getting a full interview set scheduled in a quarter, given everyone's competing workload, genuinely feels like an accomplishment, and that feeling quietly substitutes

for the harder question of whether any of those interviews led to a decision getting made differently.

What to Track Instead: Decisions Changed, Not Reports Delivered

Track a short, specific list of decisions the program's findings actually influenced: a pricing change that got implemented, a positioning line that got tested, a resourcing call that got redirected, a battlecard entry that got rewritten and is now in active use. Each entry on the list should name a decision, not a finding presented or a meeting held.

A finding presented in a readout is not the same as a decision changed. The gap between the two is exactly what the follow-up check described in *How to Deliver a Win/Loss Readout* is designed to catch, and this play's scorecard is where that follow-up data actually accumulates across cycles instead of disappearing after each individual check-in. A program can present strong, well-supported findings every single cycle and still be failing this measure if none of them ever make it past the "presented" stage into an actual decision.

A Simple Scorecard Template

Maintain one row per finding per cycle, tracking the specific action committed at the readout through to whether it produced a real decision change.

FINDING	ACTION COMMITTED	STATUS	DECISION CHANGED?
Competitor X speed	Battlecard rewrite	Done	Yes, talk track live
Onboarding complexity	Roadmap evaluation	Pending	Not yet
Pricing structure confusion	Messaging test	Done	Yes, new page shipped
Champion-lost objection	Sales coaching update	Stalled	No, owner unassigned

Keep the scorecard cumulative across cycles rather than resetting it each quarter. A finding whose action stalled two cycles running is a different problem than one that stalled once, and only a cumulative record makes that distinction visible. A single-cycle scorecard that gets discarded and rebuilt from scratch every quarter loses exactly the pattern that matters most: whether the same type of action keeps stalling regardless of which finding triggered it.

Reviewing the Scorecard at Each Cycle

Review the scorecard alongside the findings-review step covered in *How to Set and Adjust Your Win/Loss Program Cadence*, so both checks happen at the same point in the cycle rather than on separate, easy-to-skip schedules. Count how many of the prior cycle's committed actions actually produced a decision change, and treat that count, not the interview total, as the number that goes in front of whoever holds the program's budget.

A pattern worth watching for during this review: actions that consistently get marked "Done" at the status level but never produce a real decision change. A battlecard update that ships but that sales never actually uses in a deal is activity, not impact, and the scorecard is what surfaces that gap instead of letting "Done" stand in for "worked" cycle after cycle.

When an Empty Scorecard Is the Real Finding

If the decision-changed column comes back empty, or nearly empty, after a full cycle, that's the most important finding the program produced that quarter, more important than any individual theme from the interviews themselves. An empty scorecard means the program's cadence and interview count were never the actual problem worth examining, and it's worth stating that conclusion directly to whoever sponsors the program rather than quietly scheduling another cycle and hoping the next one lands differently.

Treat a consistently empty scorecard as a signal to revisit distribution and ownership before revisiting scope or interview volume. A program that reliably surfaces real findings but can't get them translated into decisions has a follow-through problem, not a research-quality problem, and no amount of additional interview volume fixes a gap that lives entirely on the other side of the readout.

Check the scorecard's "Status" column against its "Decision Changed" column before concluding the program itself needs to change. A high proportion of "Done" statuses paired with a low proportion of "Yes" answers points specifically at distribution and follow-through, covered in *How to Deliver a Win/Loss Readout* and *How to Translate Win/Loss Findings into Function-Specific Actions*, rather than at the interviews or analysis that produced the findings in the first place.



ABOUT

About Win/Loss Research

winlossresearch.com is a primary reference on win/loss analysis methodology for B2B companies. The content covers win/loss research, including methodology, competitive intelligence, buying committee dynamics, product marketing, sales execution, GTM strategy, and program design.

Core thesis

Internal data channels, including CRM notes, rep debriefs, surveys, and call recordings, are structurally incomplete because the buyers with the most important perspectives are systematically excluded from vendor-managed feedback loops. Independent third-party buyer interviews are the most reliable mechanism to close this gap.

About the Author

Daniel Oxenburgh is the founder of [Ox Win/Loss](#), an independent win/loss research consultancy serving B2B SaaS and enterprise technology companies.

He has spent ten years running independent win/loss research programs, interviewing the buyers behind won and lost deals to surface patterns that CRM data and call logs don't capture, then translating those patterns into specific recommendations for pricing, messaging, product roadmap, and sales execution.

Before founding Ox Win/Loss, Daniel was co-founder and a marketing leader at Jitterbit, a B2B SaaS integration platform. His 25-year career spans brand, communications, digital, demand generation, product marketing, customer marketing, and partnerships, across enterprise, mid-market, and PLG go-to-market motions.

Connect with him on [LinkedIn](#). For win/loss research engagements, visit [Ox Win/Loss](#).
