

Idea Validation

ESTEEM Graduate Program • University of Notre Dame

Germán Estrada Mendoza • gestrad1@nd.edu • Office hours: Tue & Thu • 3:30–5:00 PM

TA: Parker Dietzel, ESTEEM '26 • pdietzel@nd.edu • 5 sessions • Tue & Thu • 1:00–3:30 PM • July 9–23, 2026

PREREQUISITE

Design Entrepreneurship

SEQUEL

PMF • ESTEEM Fall

ENROLLMENT

~50 students • teams of 2

LOCATION

In-person

COURSE PURPOSE

In a tech startup, the code is cheap. What's scarce is certainty that real demand exists before you build anything. This course is five sessions spent earning that certainty, in direct contact with people who have a real problem.

The course runs on a single question that gets sharper every session: **who is desperate enough that they would use an imperfect version of your solution today?** To answer it, you work through five levels of evidence, one at a time:

THE WHO SCALE • EVIDENCE STANDARD FOR EACH LEVEL

- 1 Who has the problem: can describe a concrete failure or cost
- 2 Who is aware they have it: names it without your prompt
- 3 Who has taken action to solve it: searched, evaluated, has a timeline
- 4 Who has built a workaround: spreadsheet, hack, manual process, homegrown tool
- 5 Who has paid to fix it: money spent, budget allocated, or personal spend already happening

Sessions 1–4 are a **concept test**: you test whether the problem is real and narrow the WHO through weekly field conversations. Session 5 is the first **implementation test**: you take what you learned about the WHO, build the one benefit that matters most to them, and map what it costs them not to have it.

How this course teaches: the experience comes before the framework. Session 1 opens with a card on your seat. You push legacy technology to strangers for an hour before you hear the words *push* or *pull*. That sequence is intentional, and it repeats in every session that follows.

WHERE THIS COURSE FITS

You arrive here having completed Design Entrepreneurship. Design thinking starts from an audience whose demand you already know is real, then iterates on the solution until it fits. That works when the demand is settled.

This course starts somewhere else. You hold an insight, a fixed belief about what is now possible, and you do not yet know if anyone is desperate for it. So you keep the WHAT fixed and iterate on the WHO. The risk here is market risk: whether a desperate WHO exists at all. The best team and the best solution still waste the effort if that WHO is not there.

Stage	Course	Question it answers
Completed	Design Entrepreneurship	What problem space matters, and which solution directions are worth exploring?
This course	Idea Validation	Is there a desperate WHO for this fixed WHAT?
Fall	PMF	Which implementation does the validated WHO pull hardest?

LEARNING OBJECTIVES

By the end of this course, you will be able to:

1. Distinguish inflection, insight, and idea. Know why the insight stays fixed while the WHO is the only variable you iterate.
2. Apply the non-consensus test: separate a real insight from a trend observation anyone else would agree with.
3. Conduct field conversations that generate behavioral evidence, not polite agreement.

- Score each conversation against the five-level desperation scale and make a legitimate WHO decision from the pattern.
- Build a behavioral ICP and positioning statement from evidence, not from assumptions made before fieldwork.
- Implement the one main benefit of your idea for the validated WHO, and name what it costs them not to have it.

SESSION SCHEDULE

Session	Date	Pre-reading	In-class work
1: The Problem Space	Thu Jul 9	Pattern Breakers Ch 1–2, 4–5	Field Push exercise · push vs. pull debrief · inflection, insight, idea · A1 released
2: Confirm the Problem	Tue Jul 14	Mom Test Ch 1–3 · materials distributed in class	Desperation scale introduced · first conversations debrief · Mom Test discipline · team formation · A2 begins
3: Find Who Acts	Thu Jul 16	Pattern Breakers Ch 8	Behavioral evidence vs. stated intent · WHO refinement, narrow to action-takers · ICP slice discussion · A2 checkpoint 1
4: Find Who Pays	Tue Jul 21	Pattern Breakers Ch 9 · materials distributed in class	Workaround and payment evidence · validity threats · ICP locked · positioning statement · A2 checkpoint 2 · A3 due Jul 22
5: First Implementation	Thu Jul 23	Arrive with all evidence from Sessions 2–4	Recap lecture · in-class implementation workshop · demo stations · A4 in-class

REQUIRED TEXTS

- Maples Jr., M. & Ziebelman, P. (2024). *Pattern Breakers*. Inflection, insight, the non-consensus test. Chapters 1–2 and 4–5 before Session 1; Chapter 8 before Session 3; Chapter 9 before Session 4.
- Fitzpatrick, R. (2013). *The Mom Test*. Core field method. Chapters 1–3, assigned before Session 2.

We hand out reference materials in class at each session: frameworks for customer segmentation, desperation scoring, positioning, and keeping your fieldwork honest. Nothing else to buy.

APPROVED TOOLS

AI is a tool, not a source of evidence. You may use AI to structure field notes, identify patterns across conversations, draft outreach, sketch a rough prototype, or analyze data you collected. AI may not run conversations, write customer quotes, produce desperation scores, or stand in for anything that is supposed to come from real people. The evidence must come from you, from the actual human beings you talked to. If the interaction did not happen, it cannot be submitted.

Outreach: use everything you have. Your personal network, warm introductions, and your standing as a Notre Dame grad student all count, and so do cold channels: LinkedIn, Reddit, Slack communities, Twitter/X, cold email to public professional addresses. A warm intro to someone who turns out to be genuinely desperate is a great find. The trap runs the other way. People who know you will be nice to you, and niceness is not desperation. Test for it every time. Score the conversation and cite the behavior. Throw out anything that reads like a favor. The cleanest signal comes from people who show desperation on their own, with no reason to flatter you. A desperate stranger is the strongest evidence you can bring.

GRADING

Component	Weight
Opening Reflections and In-Class Activities	20%
A1: Inflection	10%
A2: Conversation Evidence (cumulative, 3 checkpoints)	30%

Component	Weight
A3: ICP + Positioning Statement	20%
A4: Implementation Workshop Output	20%
Total	100%

Grading Scale

Letter	Range
A	94.0 – 100%
A–	90.0 – 93.9%
B+	87.0 – 89.9%
B	84.0 – 86.9%
B–	80.0 – 83.9%
C+	77.0 – 79.9%
C	74.0 – 76.9%
C–	70.0 – 73.9%
D	64.0 – 69.9%
F	0.0 – 63.9%

Cap Rules

Caps are visible standards. They're posted when the relevant technique is taught, not at grading time. When a cap triggers, the maximum score on that assignment is **85% of available points** (B ceiling). You can still earn an A overall by performing above the cap on other components.

Assignment	Cap trigger	Max score
A2	A checkpoint misses its conversation minimum, a desperation score is not traceable to a cited behavior in the raw notes, or the ICP does not get sharper across checkpoints	85% (B)
A3	ICP is demographic, not behavioral; ICP is not traceable to specific conversation evidence; or the benefit section lists features instead of naming one benefit	85% (B)
A4	Either team member cannot explain their own WHO evidence during the demo	Individual deduction

OPENING REFLECTIONS AND IN-CLASS ACTIVITIES

Most sessions open with a handwritten reflection: individual, closed device, first 10 minutes. Session 1's reflection runs after the field exercise. Reflections are graded on behavioral specificity: did you name a real action, a cost, or something someone actually did? Not on length or polish.

Sessions also carry graded in-class activities, some on your own and some with your team. The number is not fixed; it depends on where each session goes.

Together these are 20% of your grade. Each reflection and activity is scored on the same short scale: full credit for being there with work that shows specific behavior or real effort, half for taking part but staying generic, none for an absence or a blank. Your grade here is the share of available points you earn, scaled to 20%, so the count can stay open: it is a ratio, not a set number of slots. Team activities earn a team score, reflections are yours alone, and you can only earn these points in the room. None of it can be made up.

ASSIGNMENTS

A1: Inflection

Individual · Due Mon Jul 13 · 11:59 PM · 1 page maximum

Five one-sentence answers that follow a chain, each one depending on the one before it. (1) The inflection: what changed in the world. (2) The insight: your non-consensus belief about what is now possible *because of that change*. (3) Your first idea: one disposable implementation of that insight. (4) Your hypothesized WHO: the one specific type of person who would be desperate for it. (5) The strongest objection to your insight: why a smart, informed person who has read the same things you have would still doubt it.

You will read this aloud in Session 5 next to where you actually ended up. That contrast is the learning. No revisions after submission — not as a penalty, but because the gap between your Day 1 bet and your Day 23 evidence is exactly what this assignment is measuring. A1 is individual. Teams form in Session 2, so every student commits to a real starting point before the fieldwork begins.

A2: Conversation Evidence

Team · Three checkpoints, each due the day before the next session · 30% of course grade

A2 is a living field log and the spine of the course. The discipline is simple to say and hard to do: start with a hypothesis about your ICP, test it in real conversations, and tighten it with what you learn. Starting broad is fine. Each session you narrow, and as the ICP gets sharper the outreach gets harder, because a specific WHO is a smaller and harder group to reach. That difficulty is the work, not a sign you are off track.

Two rules decide what counts. First, the person has to fit your current ICP. A conversation with the wrong person is not evidence, however it went. Second, if we cannot see it, it did not happen, so document everything: raw field notes at every checkpoint, no summaries or paraphrases written after the fact. Reach people any way you can, your network, warm intros, your ND standing, or cold outreach, and label your relationship to each one. A warm contact being kind to you is not evidence either. Score the desperation and cite the behavior.

CHECKPOINT 1 WED JUL 15 11:59 PM	Levels 1–2: confirm the problem exists, find who has it and who is aware. Start from your ICP hypothesis, broad is fine here. Minimum 5 on-ICP conversations, raw notes, desperation score per conversation with the specific evidence cited.
CHECKPOINT 2 MON JUL 20 11:59 PM	Level 3: narrow to action-takers. Add at least 3 new on-ICP conversations targeting this level. The minimum drops from 5 to 3 because your WHO is now more specific — a smaller, harder pool to reach. Finding 3 genuine action-takers in a narrow segment is harder than finding 5 people with the broad problem. Update the WHO profile: who scored higher, who dropped out, and why.
CHECKPOINT 3 WED JUL 22 11:59 PM	Levels 4–5: workaround and payment evidence. Add at least 3 new on-ICP conversations. The same logic applies: the WHO is tighter, the outreach is harder, 3 is the floor. Full WHO profile with behavioral ICP draft. This is the evidence base for A3, which is due the same night.

What is graded: the volume of on-ICP conversations you actually run, the raw documentation behind each one, whether every desperation score is backed by a behavior cited from those notes, and whether the ICP gets sharper across checkpoints. **The grade never depends on how much desperation you find.** A conversation honestly scored at level 1, with the behavior cited, counts exactly as much as one scored at level 5. Each checkpoint's conversation minimum is a volume floor, not a quota of high scores: running the conversations and documenting them is what matters, not how high they score. Finding little desperation and documenting it cleanly is a real, full-credit result. Inflating a score to look better is the one move that fails the assignment outright under the integrity policy below.

Cap: 85% if a checkpoint comes in under its conversation minimum, if any desperation score is not traceable to a cited behavior in the raw notes, or if the ICP does not get sharper from one checkpoint to the next.

A3: ICP + Positioning Statement

Team · Due Wed Jul 22 · 11:59 PM (before Session 5)

From your A2 evidence, three things: a one-sentence behavioral ICP (who they are, what they do, why that behavior makes them desperate), a positioning statement using the format from the Session 4 card, and the one benefit that matters most to your WHO. Write all of it after the evidence, not before. The ICP must be grounded in specific conversations, not demographic categories.

The one main benefit is a single thing, the most important one. Test it like this: if your WHO does not use your solution, what is the consequence? If the honest answer is close to "we go out of business," you have a painkiller. If it is "nice to have," you have a vitamin, and you keep digging. Name one benefit, not two, and never a list.

A benefit is not a feature. The benefit is the outcome your WHO cannot live without. Features are how you deliver it. "Cuts month-end close from three days to three hours" is a benefit. "Dashboards, auto-reminders, and integrations" is a feature list. Pick the one benefit first. The features are just the delivery.

Cap: 85% if the ICP is demographic rather than behavioral, if it is not traceable to specific A2 evidence, or if you name a feature list instead of one benefit.

A4: Implementation Workshop Output

Team · In class · Thu Jul 23 · No advance submission

Session 5 is a structured workshop. You arrive with your A2 field log and A3 ICP. The session runs in three blocks:

1. **Recap lecture (30 min):** What the class learned across four sessions: the WHO evidence standard, the move from concept test to implementation, and what's next in PMF.
2. **Implementation workshop (90 min):** Take the one main benefit you named in A3 and build a first version of it — a landing page, a working screen, or a paper prototype, whichever gets the benefit in front of a person fastest. Then test the consequence: what happens to your WHO when this solution does not exist, and how bad is it really?
3. **Demo stations:** Each team presents their WHO evidence, the one main benefit, the consequence map, and the first implementation artifact.

We grade in this order: first the quality of your evidence, then how clearly you map the consequence, then how honest you are about what the data showed, and last the arc from A1 to A4. Teams that found no desperation and can explain exactly what they found, which assumption failed, and what they'd test next score equally to teams with strong signal. This course evaluates calibration, not luck.

Cap: if either team member cannot walk through their own WHO evidence during the demo, that person's grade is docked individually. The evidence has to be yours, not your partner's summary of it.

ASSIGNMENT SUMMARY

Assignment	Type	Due	Weight
Opening Reflections and In-Class Activities	Individual + team	In class · no make-up	20%
A1: Inflection	Individual	Mon Jul 13 · 11:59 PM	10%
A2: Conversation Evidence	Team	Wed Jul 15 · Mon Jul 20 · Wed Jul 22 · 11:59 PM each	30%
A3: ICP + Positioning Statement	Team	Wed Jul 22 · 11:59 PM	20%
A4: Implementation Workshop Output	Team	In class · Thu Jul 23	20%

COURSE POLICIES

Attendance

This is an in-person course. Every session runs on graded work you do in the room: in-class exercises, fieldwork, team debrief, and an opening reflection on most days. Most of it feeds graded work you finish after class. None of the in-class work can be made up. Miss a session and you lose those points and fall behind on your own fieldwork. Worse, your team is two people, so your partner absorbs a two-person workload alone. You are responsible to each other. Missing more than one session may result in removal from the course.

Academic Integrity

The Notre Dame Honor Code applies in full. Customer conversations must be authentic. Fabricated field notes or synthetic behavioral evidence are grounds for course failure. AI use follows the tool policy above. AI cannot be a source

of evidence. You are responsible for the accuracy and attribution of any AI-assisted work.

Communication

Notre Dame email and Canvas. Check both daily during the course window (July 9–23).

Accommodations

Contact the instructor before Session 1. Students who have or think they may have a disability are invited to contact Sara Bea Disability Services and the Notre Dame Center for Student Support and Care for a confidential discussion about accommodations.

| WHAT COMES NEXT

This course ends with a first implementation for a validated WHO. The fall PMF course starts there, taking that implementation to real customers and testing whether it produces the behavioral signal that defines product-market fit. Most founders build first and search for the WHO later. You will not. You will know who you are building for before you build anything load-bearing. That changes what you build and how fast you build it. And much more of it survives contact with the market.