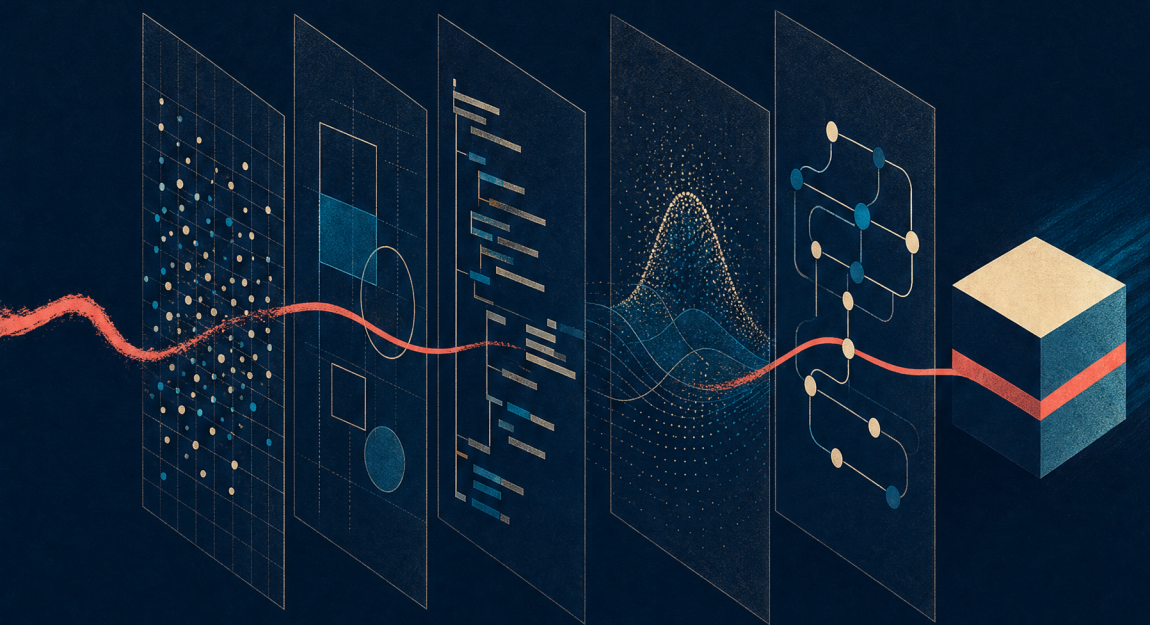


PM is now another member of technical staff

A field guide for product managers working across
data, design, code, AI, and delivery



RAZII ABRAHAM

PM Is Now Another Member of Technical Staff

*A Field Guide for Product Managers Working
Across Data, Design, Code, AI, and Delivery*

RAZII ABRAHAM

Contents

Introduction: The PM Is Moving Closer to the Work	3
Part I: The Technical Floor	17
1. Why PMs Need a Technical Floor	19
2. APIs Are Where Product Ideas Become System Behavior	29
3. Delivery Is a Product Skill	53
4. Image Assets Are a Product Architecture Problem	71
Part II: The Data Floor	93
5. SQL, NoSQL, and Product Reality	95
6. Product Analytics Is a Model of User Behavior	117
7. Experiments Turn Product Opinions into Evidence	139
Part III: The Design Floor	161
8. Flows Make Ambiguity Visible	163
9. Design Systems Turn Taste into Shared Infrastructure	183
10. Product Psychology Explains Why Users Move or Stop	201
Part IV: The AI Product Floor	219
11. The Moment Traditional PM Stops Being Enough	221
12. Do Not Start with AI; Start with Solution Fit	237
13. AI Features Need Evaluation Systems	261
14. AI Products Improve Through Calibration Loops	281
15. When a Feature Becomes an Agent	299
Part V: The AI-Augmented PM Operating Model	315
16. AI Agents Are Teammates, Not Magic Assistants	317

17. The Repo Is a Product Surface	333
18. Discovery Becomes Evidence Operations	351
19. AI Makes Prototyping Cheap, Not Thinking Cheap	365
20. Delivery Still Matters Most	381
Conclusion: Follow the Work	403
Exercise Index: Build Your Operating System	409
Glossary	413
About the Author	421
References	423
Index	427

Introduction: The PM Is Moving Closer to the Work

My first meaningful contribution to a codebase was a prompt inside Heatseeker’s onboarding flow. Heatseeker, the customer-intelligence platform where I work as a product manager, helps marketing teams validate ideas with synthetic customer personas and live in-market experiments. The onboarding flow needed enough context to make a useful first impression. When a new user arrived with a work email address, the system could derive the company’s website domain, retrieve publicly available pages, and extract a working description of the business: its flagship products or services, what those offerings did, and the customer segments they appeared to serve. The promise sounded simple: understand the user’s business before asking them to explain everything from zero.

The first version could sound convincing and still be wrong. It sometimes failed to identify which product was actually the flagship. When a company had several offerings, it could blur them together. Worse, a model could complete missing context with plausible information remembered from training rather than information present in the retrieved pages. That information might be inaccurate, belong to a different company, or have become obsolete.

One test against Noted, a small workspace product I built at `wellnoted.dev`, made the problem visible. Heatseeker described the workspace accurately, but it named the product “WellNoted,” inferred from the domain rather than the wordmark “Noted,” and generalized the branded “AI Squad” to “AI agents.” The result was useful, fluent, and wrong in two particulars.



Add a product

Adding a **product** isn't just a setup step—it's how you start learning what truly works for your business

ABOUT YOUR PRODUCT

Product Name

WellNoted

What does your product do?  Generated with AI

A block-based workspace for writing, planning, and organizing documents. It features custom AI agents powered by your own API key to write, research, and organize notes, alongside project tracking and one-click web publishing.

[< Back](#) [Tell us about your customer >](#)

Screenshot 0.1. Heatseeker could generate a useful description and still miss the particulars. Here it inferred "WellNoted" from the domain and generalized "AI Squad" to "AI agents"; both fields remained editable.

The problem was not only “write a better prompt.” The system shape had to change.

Instead of asking one model call to produce the company summary and customer segments in parallel, we made the work sequential. First, the system extracted and synthesized grounded product information. Then that product output became context for identifying the likely customer segment for that specific offering. The narrower second task had a clearer object to reason about.

I iterated on the instructions and tested sampling settings such as temperature and top-p, but the most useful evaluation material came from failures we had already seen. I reran the changed approach several times against those cases. For this feature, variation was not the value. We wanted repeated runs to stay grounded in the retrieved evidence and converge on a defensible description.

Passing those cases did not prove the system would handle every unknown company. Generative behavior remained probabilistic, and the web introduced its own failure modes. Some sites blocked automated retrieval. Some brands had many products spread across child pages. The team had to decide whether to crawl more deeply, stop with limited context, or consider other permitted sources such as a company page or public article.

The user experience carried the final safeguard. We showed the generated context on screen and let the user overwrite it. The AI could create the moment of recognition—*this product already understands something about my business*—without pretending that its interpretation was authoritative. We showed the magic, but we also showed control.

The feature shipped. I iterated on the prompt more than once, and that first version became part of Heatseeker's context layer. As of July 2026, the basic approach was still in use.

That contribution changed how I understood product work. The prompt lived in the repo, but its quality depended on the whole chain: retrieval, grounding, call order, model settings, fallbacks, evaluation cases, interface states, and user correction. The requirement and the implementation could no longer be separated cleanly.

So I moved closer to the work. That experience is the starting point for this book.

Why “Member of Technical Staff”

The title is partly serious and partly a joke. Member of Technical Staff is a long-used technical title visible from Bell Labs to current AI labs.¹ I am not proposing a new LinkedIn title or claiming that PMs become research engineers. The phrase describes proximity.

A technically participatory PM can inspect an API contract, trace an event, map a flow, understand a flag, test a prompt, run a product locally, review a diff, verify a preview, and sometimes contribute a bounded change. Specialists remain accountable for architecture, reliability, security, maintainability, data quality, and design craft. Participation means knowing enough to carry product intent into those systems and to recognize where another person’s expertise begins.

Product managers can move closer to the commercial organization or closer to the technical product team. Strong products need both. This book follows the second direction: what happens after a requirement is approved, why a small change is expensive, why evidence sources disagree, why an AI demo can be impressive but unreliable, and why technically valid implementation can still miss the user.

My route was curiosity rather than formal technical training. Engineering vocabulary became product vocabulary when I saw what it controlled: an API response determined visible states; an identifier made analysis possible; a release pipeline determined whether change could be tested safely; a design token helped a generated UI belong; a feature flag separated deployment from exposure.

The closer I moved to those systems, the clearer the product became.

1. Nokia Bell Labs, “Presidents of Bell Labs,” <https://www.nokia.com/bell-labs/about/history/innovation-stories/presidents-bell-labs/>; OpenAI, “Senior Staff Software Engineer, Gov,” <https://openai.com/careers/senior-staff-software-engineer-gov-washington-dc/>; and Anthropic, “Claude Code Live: Origin Story, Live Demos, & Best Practices,” 27 March 2025, <https://www.anthropic.com/webinars/claude-code-live>. The Bell Labs source documents historical use but does not establish the title’s exact origin; the company sources support only present-day title usage, not a claim that product and engineering roles are merging.

AI changed the cost of participation

Before AI coding agents, technical curiosity was useful but relatively expensive to act on. A PM could read documentation, ask an engineer to explain an architecture, learn basic SQL, or inspect analytics. Direct contribution usually required a larger investment in coding skill and local tooling.

Agents changed that cost. Today I use tools such as Claude Code, Codex, and Cursor across two kinds of work.

Discovery work includes:

- Synthesizing research.
- Mining call transcripts.
- Investigating product issues.
- Running data analysis.
- Comparing evidence across company systems.
- Brainstorming alternatives.
- Building prototypes.

Delivery work includes:

- Loading a Linear ticket and its product context.
- Inspecting the relevant code path.
- Producing an implementation plan.
- Making a bounded change.
- Adding event tracking.
- Running checks.
- Reviewing a diff.
- Preparing work for human review.

The important change is not that a model can perform tasks in both categories. It is that product context can travel further. A PM can begin with a customer signal, write a short Notion brief, link it to Linear, rev-

iew it with the team, create a prototype, inspect the repo, implement a small slice, prepare a pull request, verify a preview, add tracking, and monitor the first release signals.

The handoffs do not disappear. They become more inspectable.

Lower friction raises the judgment bar

There is a tempting story that AI makes technical fluency less important because the PM can ask a model to perform the technical work. The opposite is closer to the truth.

AI makes it cheap to produce plausible technical artifacts. That includes plausible artifacts that are wrong.

An agent can summarize a transcript and erase an important contradiction. It can generate a query that counts events when the decision requires users. It can create a polished prototype that tests no meaningful uncertainty. It can add an analytics event at the wrong moment. It can implement a feature literally while missing the behavior users need. It can produce code that looks consistent but copies a legacy pattern the team is trying to remove.

AI lowers the cost of participation. It does not lower the cost of being wrong.

The PM still owns the product thread:

- What user problem started the work?
- What evidence supports it?
- What behavior should change?
- What should remain unchanged?
- What is the smallest useful scope?
- What would make the output unacceptable?
- How will the team verify the result?
- What should happen after release?

The agent can continue the task. The PM has to preserve the meaning.

The codebase has to make participation safe

There is an important caveat: access to an AI coding agent does not make every codebase safe for new contributors. I have had the privilege of working in a codebase prepared for AI-assisted contribution. The environment has written standards, reference implementations, type checks, linting, formatting, builds, tests, CI/CD, design rules, review workflows, and human engineering judgment.

Those guardrails matter because an AI agent learns from the local context it reads. If documentation says “do not use any” while nearby code uses any everywhere, the codebase is teaching a contradictory lesson. Bad examples are still examples.

Engineering teams create the paved roads that make broader participation possible. They establish architecture, clean up conflicting patterns, automate deterministic checks, document important decisions, and review changes that require context no rule can fully encode. The future described in this book is not a PM using AI to route around engineers.

It is a product team designing a better collaboration system between people and agents, with responsibility still visible. Engineers create and maintain the technical paved road; designers encode reusable interaction standards; and data practitioners protect the quality of measurement. PMs carry user context and product judgment through those systems, while agents accelerate bounded inspection, synthesis, prototyping, and implementation rather than becoming an unaccountable owner of the result.

Participation expands because the system makes accountability visible.

The long PRD is no longer the center of gravity

This operating model changes the job of product documentation. Written thinking remains essential. Complex, high-risk, or cross-functional work may still need a substantial product document. But the value of a spec is not measured by its length. In a faster feedback loop, a good brief needs to travel.

It should define:

- The objective.
- The user problem.
- The expected behavior.
- The output.
- The constraints.
- The success and guardrail signals.
- What is out of scope.
- What remains uncertain.

That brief can become input to a prototype, a flow, a Linear ticket, an implementation plan, an event plan, a QA checklist, a pull request, or a release note. Clarity is still the PM's responsibility. It now emerges through a chain of connected artifacts and verification, not through one document pretending to settle every question before the team begins.

The operating chain

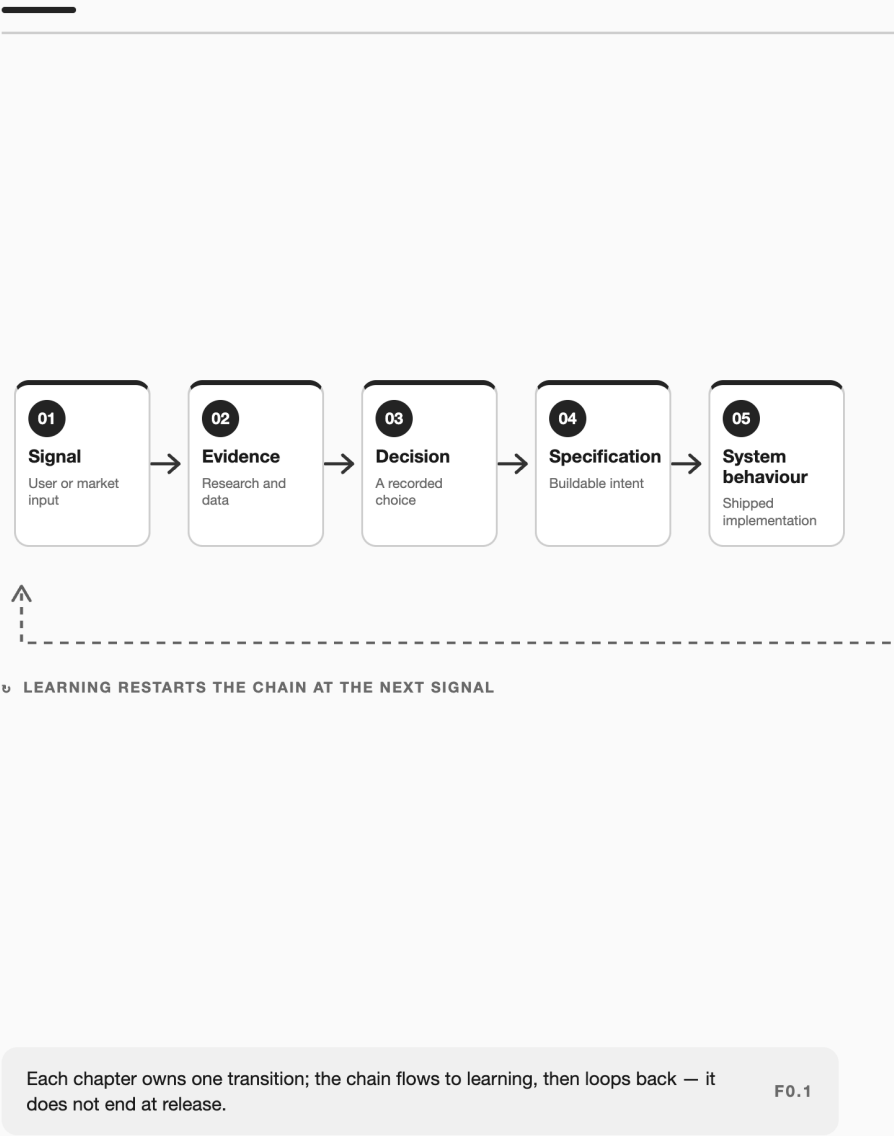
This book follows one chain:

signal -> evidence -> decision -> specification -> system behavior -> review -> release -> measurement -> learning

FIGURE 0.1

The signal-to-learning operating chain

Product intent survives only when each transition stays inspectable; learning begins the next loop rather than ending the work.



CONTINUED ACROSS THE SPREAD



Figure 0.1. Product intent survives only when each transition remains inspectable; learning begins the next loop rather than ending the work.

A product team can lose intent whenever one artifact is translated into the next. If provenance disappears, a customer comment becomes an unsupported generalization and a metric becomes a story it does not prove. If the decision boundary then disappears, the requirement becomes vague, the implementation covers only the happy path, and review can no longer tell whether working code is the right code. Even a well-reviewed change can remain hidden behind a flag or reach users without the event needed to learn from it.

The technically participatory PM learns to follow those transitions. Not to control every one of them. To keep the user problem, evidence, and expected behavior visible while the work moves through people and systems.

What the five parts build

The **technical floor** follows ideas into API contracts, delivery infrastructure, and media architecture. The **data floor** makes stored records, event semantics, and experiments inspectable. The **design floor** exposes flow, shared interface decisions, and behavioral friction. The **AI product floor** chooses solution shape, defines quality, evaluates failure, and calibrates probabilistic behavior. The **AI-augmented operating model** brings those floors together in one competitor-analysis feature from evidence and prototype through review, release, measurement, and learning.

Who this book is for and how to use it

This field guide is for practising product managers who are already a year or more into the role—who run discovery and delivery and now want to inspect the technical, data, design, and AI artifacts around their work with specialist support. It assumes you have shipped features and sat through the arguments that surround them; it is not a “how to become a PM” primer, and a reader brand new to the role may find the Technical and Data floors heavy going before the payoff lands. Product designers, engineers, and others who partner closely with

product teams will find much of it maps onto their side of the same work. It teaches participation thresholds, not substitute expertise or permission to bypass collaboration.

Read it in one of three ways:

- **Quick path:** Chapters 1, 8, 13, 16, and 20 build the participation boundary, flow, AI quality bar, portable context, and delivery slice. Add Chapter 15 when the product itself contains an agent.
- **Chapter path:** move through all five floors and reuse one product problem so the exercises accumulate.
- **Capstone path:** begin with one evidence row in Chapter 18 and carry it through prototype and delivery in Chapters 19–20.

The goal is not vocabulary or completed templates. It is one consequential decision whose evidence, implementation, review, release, and learning another person can inspect. The Exercise Index maps each chapter to a reusable artifact, while the companion repository provides the longer labs, workbook, and printable figures.

The first step is to establish the minimum technical fluency required to participate responsibly.

Evidence and source note

This book uses six evidence badges consistently:

- Ⓕ **Firsthand account:** Direct author experience.
- Ⓖ **Publication-safe reconstruction:** A teaching artifact rebuilt from the available evidence without reproducing proprietary material.
- Ⓗ **Constructed teaching example:** Invented teaching material, not production evidence.
- Ⓒ **External source/framework:** A cited public source or framework.
- Ⓓ **Current vendor behavior:** Version-sensitive vendor behavior.
- Ⓗ **Recommended practice:** Normative guidance.

Notes for this introduction

- Ⓕ The Heatseeker onboarding scene is based on the author's account recorded on 12 July 2026. The current bounded treatment is authorized for publication; proprietary prompts, code, customer data, and internal performance metrics remain excluded.
- Ⓖ The historical MTS framing is associated with Bell Labs and is used as context rather than a claim about the exact origin date of the title.
- Ⓗ OpenAI and Anthropic public materials currently use Member of Technical Staff for technical roles. Tool names and company practices reflect the period of writing, not permanent endorsements.

Chapter 1: Why PMs Need a Technical Floor

A technical floor begins when *it works* stops being a complete answer. A screen can render while its request fails, a write can succeed while its event is missing, and a release can work for one account while failing for another role or a slow connection. The PM does not need to own each mechanism, but does need to locate the claim: which behavior worked, under which conditions, with what evidence, and with enough support visibility to recognize when it stops working?

Those are not questions about becoming an engineer. They are questions about the product promise surviving contact with a system.

My own move closer began with the prompt-layer contribution in the introduction. The prompt led into review; review led into the repo; small interface changes led into event tracking, data, and bugs. My background is resource and environmental economics, not computer science. I learned each layer because the product question in front of me could no longer be answered from a roadmap or screen alone. Today I can make bounded contributions across a full-stack TypeScript product, but learning TypeScript was not the destination. It was one result of repeatedly following product questions into the systems that answered them.

Technical participation compounds. Once you can inspect one layer of the product, you ask better questions about the next. Once you can verify a small change, “done” means something different. Once you have watched intent get lost between a brief and a release, implementation stops looking like a black box. That progression—not a particular stack—is the technical floor.

A floor is not a new professional identity

The floor is contextual. It is the point at which a PM can stay with a decision as it moves from requirement to system behavior, without treating one technology checklist as universal or inheriting specialist ownership by default.

A PM responsible for a public API needs to understand contracts, authentication, errors, versioning, and developer experience. A PM responsible for a marketplace needs stronger fluency in transactions, state, data quality, experimentation, and operational recovery. A PM responsible for a media-heavy product needs to understand upload, processing, storage, delivery, and cost. A PM responsible for an AI feature needs to understand probabilistic behavior, evaluation, tracing, feedback, and calibration.

The common requirement is not a technology but the ability to follow intent. A PM should be able to trace a user need into expected behavior, identify the systems and states that make the behavior possible, name the evidence that should exist after release, and verify that users received what the team decided. A break in that chain reveals the part of the technical floor that needs attention next.

The four levels of technical participation

“Be more technical” is poor advice because it does not tell a PM what to do differently. A more useful model has four levels: ask, inspect, prototype, and contribute.

These levels are not job titles. They describe how directly a PM can engage with a product question.

FIGURE 1.1

The technical participation climb

Technical confidence grows through observable contribution, not a change in title.

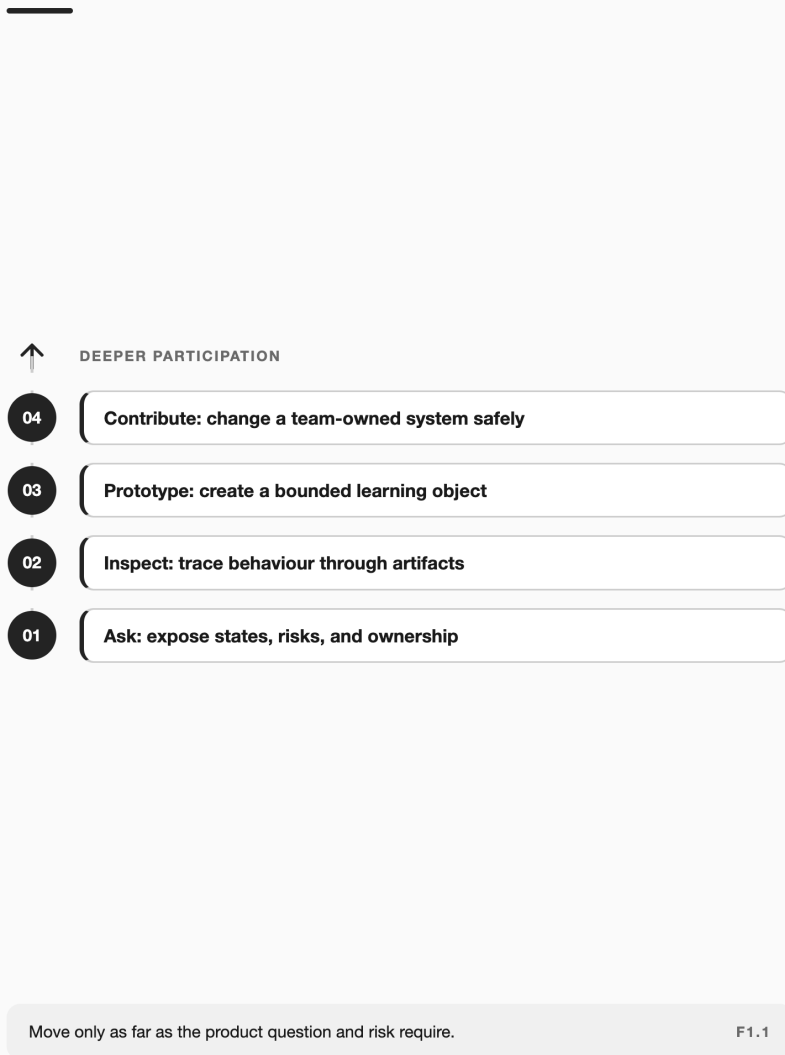


Figure 1.1. Technical participation progresses through observable work, not status; the appropriate stopping point depends on the product and risk.

Level 1: ask

At the first level, the PM can translate a feature idea into questions the team can answer. Suppose someone says:

Users should never lose a draft.

That sounds like a clear requirement. A technically participatory PM hears several unresolved product questions:

- When is a draft created?
- Is saving automatic, manual, or both?
- What does the user see while saving?
- What happens when the connection drops?
- Which version wins when the same document is open on two devices?
- Can a deleted or overwritten draft be recovered?
- What event tells us that saving succeeded or failed?
- What should support be able to inspect?

The PM does not need to answer the architecture questions alone. The skill is recognizing that “never lose a draft” describes an outcome, not a complete system behavior. At this level, technical fluency reduces ambiguity.

Level 2: inspect

At the second level, the PM can look at the product’s existing evidence instead of relying entirely on a handoff. Inspection might mean:

- Reading an API request and response.
- Looking at an event payload.
- Running a basic SQL or MongoDB query.
- Following a user flow through a feature flag.
- Asking an AI coding agent to identify the files behind a behavior.
- Reading a pull request and its review discussion.

- Checking whether a feature exists in production, staging, or only a local branch.

Inspection does not require changing anything. In fact, read-only inspection is often the best first use of an AI agent in a codebase.

A useful prompt is:

```
Trace the current save-draft behavior from the user's action to persistence.
```

```
Include:
```

- relevant files
- client and server responsibilities
- save states and error states
- analytics events
- feature flags
- tests
- any point where a draft could be lost

```
Cite the files you used. Do not modify code.
```

The output is not truth merely because it cites files. The PM still checks the important paths and asks an engineer when the risk or architecture is unclear. But the PM enters the conversation with a map instead of a vague request for explanation. At this level, technical fluency improves diagnosis.

Level 3: prototype

At the third level, the PM can make an idea visible before the team commits to production implementation. A prototype could be a wire-flow, a mocked response, a local-only UI, a new prompt tested against an evaluation set, or a temporary route that simulates a proposed experience. The purpose is learning.

For the draft-saving problem, a PM might prototype three different status patterns:

- A subtle “Saved” indicator.
- A persistent state showing “Saving,” “Saved,” and “Could not save.”
- An offline state with an explicit retry action.

The prototype can answer whether users notice the status, understand failure, and know what to do next. It cannot prove that conflict resolution, persistence, or offline recovery works. Those are outside the prototype boundary.

Technical participation at this level includes knowing what the artifact simulates. At this level, fluency reduces the cost of learning without pretending that a demo is a system.

Level 4: contribute

At the fourth level, the PM can make a bounded change inside the team's delivery system. The starting point should be deliberately small:

- Copy.
- A documented product rule.
- Event tracking.
- A design-system-compliant UI state.
- A prompt and its evaluation cases.
- A narrow bug fix.

The PM works on a branch, follows the repo's conventions, runs the relevant checks, tests the behavior, prepares the change for review, and can explain what changed. Contribution is not measured by how much code the PM writes. It is measured by whether the change is useful, bounded, understandable, and safe enough to enter review. At this level, fluency shortens the distance between product judgment and a reviewable implementation.

Participation and ownership are different

The four levels expand participation without silently transferring ownership. Product remains accountable for the problem, evidence, scope, behavior, and outcome, while engineering protects architecture, reliability, security, maintainability, and the delivery system. Design retains responsibility for interaction quality and shared patterns, and data specialists for analytical rigor and trustworthy data systems. The point of a stronger technical floor is to make those boundaries easier to navigate, not to erase them.

Real decisions cross those boundaries. Technical fluency makes the crossing more precise: the PM can tell which question belongs to product judgment, which requires deeper expertise, and where both are needed.

For the save-draft example:

- Product should define the user promise, meaningful states, recovery experience, and success signal.
- Engineering should design persistence, synchronization, conflict handling, security, and reliability.
- Design should define how status and recovery appear without distracting from the user's work.
- Data should help define and validate save-failure and recovery metrics.

The feature succeeds when those responsibilities form one system.

Your product determines your floor

Not every PM needs equal depth in every topic. The fastest way to waste time is to study technical concepts without connecting them to the product you own. Use the product's risk to decide what to learn next.

If your product depends on...	Your next technical floor probably includes...
Third-party integrations	API contracts, authentication, rate limits, webhooks, failure and recovery states
Frequent releases	Git, pull requests, CI checks, environments, flags, rollback, monitoring
User-generated media	Validation, direct upload, storage, metadata, processing, variants, CDN delivery
Product-led growth	Event design, identity, funnels, cohorts, experiment exposure, guardrails
Operational workflows	State machines, permissions, audit history, retries, support visibility
AI-generated behavior	Prompt and context design, evaluation sets, tracing, human review, feedback loops
AI-assisted contribution	Repo instructions, planning, diff review, tests, design rules, human review

A generic “technical PM curriculum” can therefore take you only so far. Vocabulary helps, but the product tells you which vocabulary must become working knowledge.

AI raises the floor

The introduction established that AI lowers the cost of producing technical work. The practical consequence is that a PM can now move from a question to a query, prototype, or code change before they fully understand the product behavior underneath it. The technical floor becomes the judgment needed to decide when the work is ready to advance and when faster execution is merely hiding an unresolved question.

For every agent-assisted task, the PM should be able to answer four questions:

1. **Context:** What did the agent inspect or receive?
2. **Boundary:** What was it allowed to change or conclude?
3. **Verification:** How will we know the output is correct enough?
4. **Escalation:** Which uncertainty requires a human specialist?

Those answers turn agent assistance into a bounded delegation. Without them, faster execution has no reliable connection to the product decision it is supposed to serve.

A 30-day technical-floor practice

The best way to build a floor is to follow one real product behavior further than you normally would.

Week 1: ask

Choose one important user action. Write the happy path, failure paths, states, data, and open technical questions. Review the questions with an engineer or another technical teammate.

Week 2: inspect

Trace the behavior through available artifacts: product, repo, API docs, analytics, data model, tickets, flags, and recent pull requests. Ask an agent to help, but require file or source references.

Week 3: prototype

Create the smallest artifact that reduces one uncertainty. State what is real, simulated, hard-coded, or missing. Test it with the relevant user or teammate.

Week 4: contribute

Make one bounded, reviewable improvement. It can be documentation, an event, copy, a prompt, or a small UI state. Follow the team's branch, check, test, and review process. Write down what the workflow taught you about the product. This is a practice for becoming less distant from the behavior you already own, not a four-week path to becoming an engineer.

Chapter exercise: map your technical floor

Choose one product area you are responsible for. Complete this map:

Product promise

What should the user be able to trust?

System behavior

Which states, systems, data, or integrations make that promise possible?

Current participation level

For this product area, can you:

- Ask the right system questions?
- Inspect the current behavior?

- Prototype an alternative?
- Contribute a bounded change?

Ownership boundary

Which decisions belong to you? Which require engineering, design, data, security, legal, or another specialist?

Next practice

What is one action you can take this week to follow the work one level further? That is the technical floor: not knowing everything, and not standing outside everything either.

Evidence and source note

Ⓟ The four-level participation ladder, ownership boundary, and 30-day practice are original teaching frameworks for this book. They define a participation threshold, not a universal competency model or a transfer of specialist accountability.

Chapter 3: Delivery Is a Product Skill

One of my clearest delivery lessons arrived when a build succeeded. The competitor-ad analyst described later in the book reached production before our codebase had the feature-flag system and rollout guidance it gained afterwards. Merge and deployment effectively meant exposure to every eligible user. There was no requesting-customer pilot, percentage ramp, or independent product kill switch.

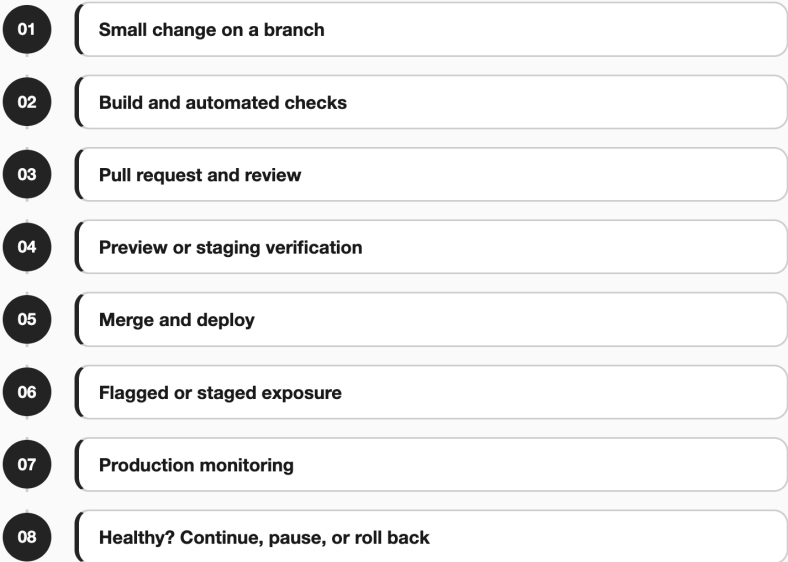
The feature remained useful. The delivery path was still less controlled than the product deserved. A change that depended on an external provider, stored third-party evidence, multimodal analysis, and agent orchestration had crossed from code to customer behavior in one large step.

That experience made a distinction impossible to ignore: a feature can be built, integrated, deployed, exposed, and healthy at different moments. “Done” hides those moments. A delivery system makes them visible. CI/CD belongs in a PM’s technical floor because it governs how a product bet becomes reviewable behavior: where code is checked, where people inspect it, where the team can exercise it, who receives it, and what happens when reality disagrees with the plan.

FIGURE 3.1

The delivery confidence path

Integration, deployment, exposure, and learning are separate product decisions.



Shipping confidence accumulates one verified boundary at a time.

F3.1

Figure 3.1. A delivery path separates code integration, deployment, exposure, and the decision to continue.

Why PMs need to understand CI/CD

Release systems shape product decisions, so PMs need enough CI/CD literacy to participate in release decisions with the engineers who own the system. When a team can integrate code frequently, test it automatically, exercise it in a preview environment, control exposure with a flag, and monitor the result, a product bet can move in smaller increments. The PM can put a narrow slice in front of the right users, learn from it, and change direction while the change is still inexpensive to review and reverse.

A fragile or poorly understood release path creates the opposite incentives. Teams protect the path by batching changes, which makes each release harder to inspect and raises the cost of a defect. Because feedback arrives only after a large commitment, PMs try to compensate with larger documents and more decisions made upfront. The caution is rational, but it is also a product constraint: weak delivery infrastructure lengthens the distance between an assumption and the evidence needed to challenge it.

CI/CD literacy lets a PM see that constraint and plan around it. Terms such as “the build is red,” “the change is in staging,” and “we need a rollback plan” locate the work at different points in the path rather than describing one vague state called done. That shared model improves the questions a PM can ask about verification, exposure, monitoring, and recovery before users inherit an avoidable risk.

Integration turns private progress into shared evidence

Consider two developers, Steve and Annie, working on the same product from different locations. Each can make a bounded change without continuously editing the other’s working copy. Their changes become meaningfully integrated only when they reach the shared source repository and the combined product is built and tested. Continuous integration moves that moment—and therefore the risk of discovering that the

changes do not work together—earlier.¹

A CI server automates part of the loop. After Steve or Annie pushes a change, it builds the combined software and runs the checks the team has encoded. Those checks depend on self-testing code: executable tests and rules that can challenge the new version without waiting for a person to repeat every known scenario. When the defined checks pass, the build is green; when one fails, it is red. Green is evidence that this version passed a known set of technical checks—perhaps type, lint, build, unit, integration, or end-to-end checks—not proof that the product behavior is correct or ready for production. A red build does not make the product decision for the team, but it does establish that the current change should not continue unchanged.

The risk is greater when an agent produced the diff. A UI can look plausible in one local path while types fail, a different flow breaks, or the code violates an automated rule. CI is not proof of value; it is one layer of evidence against false completion.

The durable habit is to integrate coherent work before it drifts so far that integration becomes a separate project. Frequent integration turns weeks of accumulated incompatibility into feedback that can arrive within minutes. “Frequent” is contextual: some teams integrate several times a day, while others work at a slower safe cadence. The principle is to avoid letting private branches become long-lived alternative versions of the product. Smaller changes are generally easier to understand, review, test, and reverse—provided the slice follows a real product or architectural seam rather than arbitrarily splitting one invariant.

Version control makes the change inspectable

Git records how the product system changed over time; repository platforms add branches, pull requests, checks, and review around that history.² For a PM, the useful workflow is simple:

1. Start from the shared codebase on a bounded branch.

1. Martin Fowler, “Continuous Integration,” <https://martinfowler.com/articles/continuousIntegration.html>. Integration frequency and branch practice vary by team; the durable principle is early, automated feedback on integrated changes.

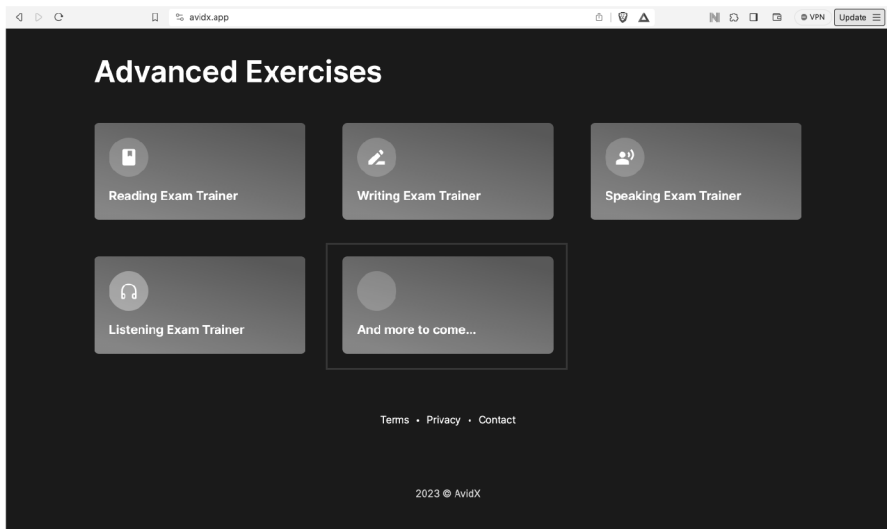
2. Git project, “Git Reference,” <https://git-scm.com/docs>.

2. Make one coherent change.
3. Record the change and its reason in a commit.
4. Push it into the team's review system.
5. Let automated checks and people challenge it before merge.

The commands can be learned when needed. The product skill is understanding that a local result is private evidence. A branch and pull request turn it into something the team can inspect, discuss, test, and connect to the original decision.

Follow one AvidX change through the delivery path

AvidX, my language-learning side project, provides a deliberately small example. The visible request was to change one landing-page phrase from “And more to come...” to “And many more...” The product effect was modest, which makes the delivery evidence easier to see: the workflow did not become unnecessary merely because the diff was one line.



Screenshot 3.1. The AvidX landing page before the change shows “And more to come...” in the Advanced Exercises section.

Before review, the project was copied to the developer’s computer and the edit kept separate from the live main branch. In plain language, the commands below mean: copy the project, create a side branch, edit one file, select it, record the change, and send the branch to GitHub. You do not need to memorize them:

```
git clone <repository-url>
git switch -c side-branch
# edit index.html
git add index.html
git commit -m "Change text to And many more..."
git push -u origin side-branch
```

The `git add index.html` line means “include this file in the next saved change.” Naming the file avoids accidentally including unrelated work, so the diff shows reviewers exactly what this change contains.

from side-branch .

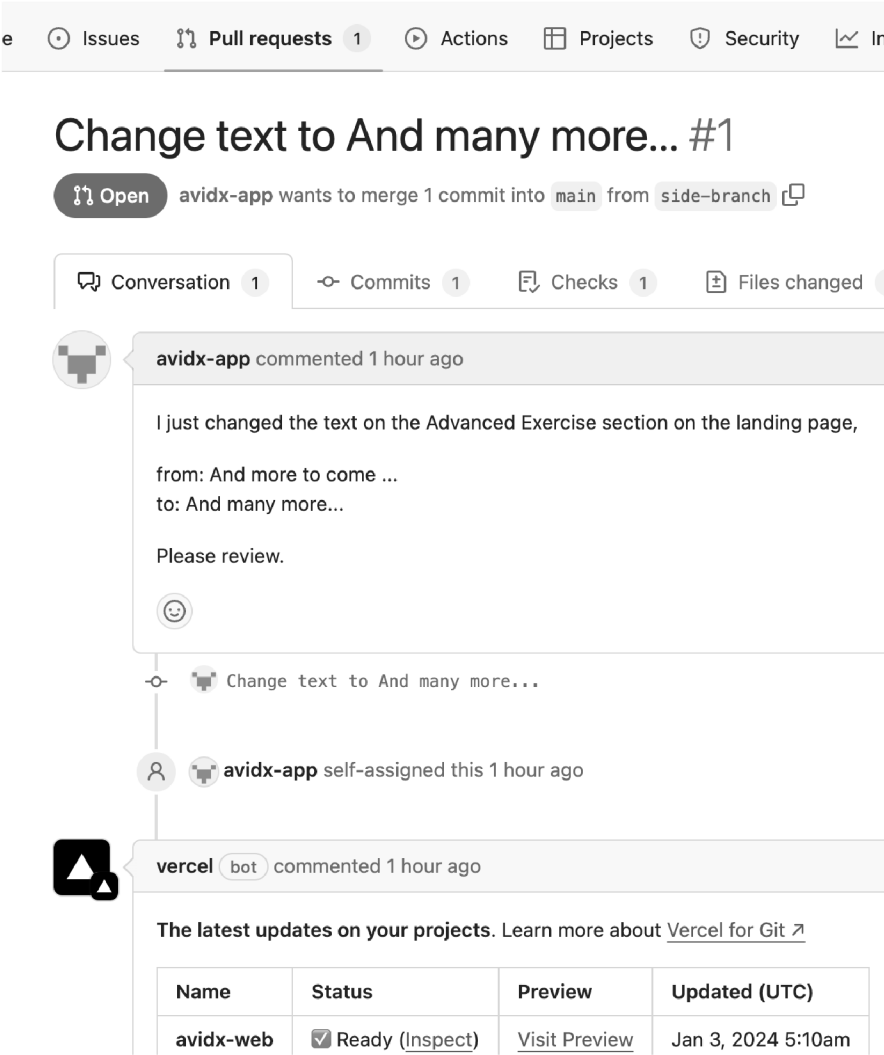
..		@@ -238,7 +238,7 @@	
238	238	<div class="coming-soon-feature-card">	
239	239	<figure class="coming-soon-feature-circle">	
240	240	</figure>	
241	-	<h2 class="coming-soon-feature-title">And more to come...</h2>	
	241	+<h2 class="coming-soon-feature-title">And many more...</h2>	
242	242	</div>	
243	243	</div>	
244	244	<ul class="footer-nav-container">	
....			
↓			

Screenshot 3.2. The diff isolates the one-line copy change on the side branch.

“Diff” is short for difference: a line-by-line comparison between two code versions. Here, the red minus line is removed and the green plus line is added; unchanged lines provide context. A diff focuses review on the change instead of showing the entire file.

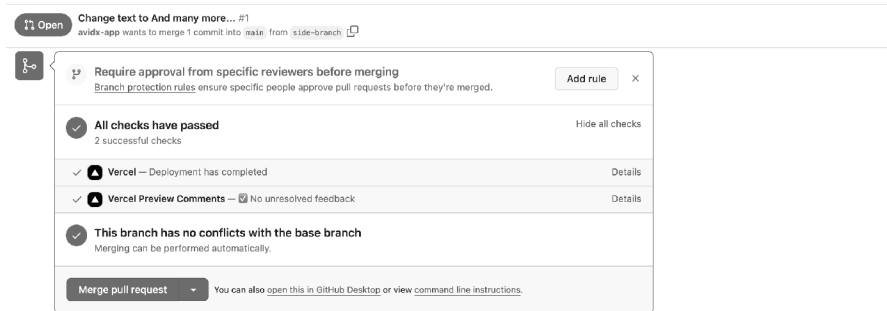
This matters when an AI coding agent writes the implementation. Its summary is not proof. The diff exposes touched files, unexpected scope, removed safeguards, and whether the code matches the request. Review it alongside the working product and tests.

Opening a pull request changed the status of the work from private implementation to a reviewable proposal. The pull request named the source and target branches, preserved the commit and diff, explained the intended copy change, and linked a deployable preview. A reviewer could therefore examine both the implementation and its rendered behavior before the work reached the main branch.



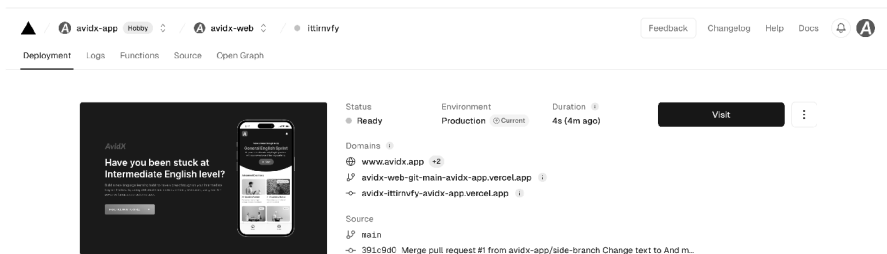
Screenshot 3.3. The open pull request brings the proposed change, branch relationship, commit, and preview into one review surface.

The automated state then turned green. In this repository, the named checks reported a completed Vercel deployment, no unresolved preview feedback, and no conflict with the base branch. That was useful evidence, but its boundary matters: it said those checks passed and GitHub could perform the merge automatically. It did not establish that the wording was a good product decision, that every user path was correct, or that the production page had changed.



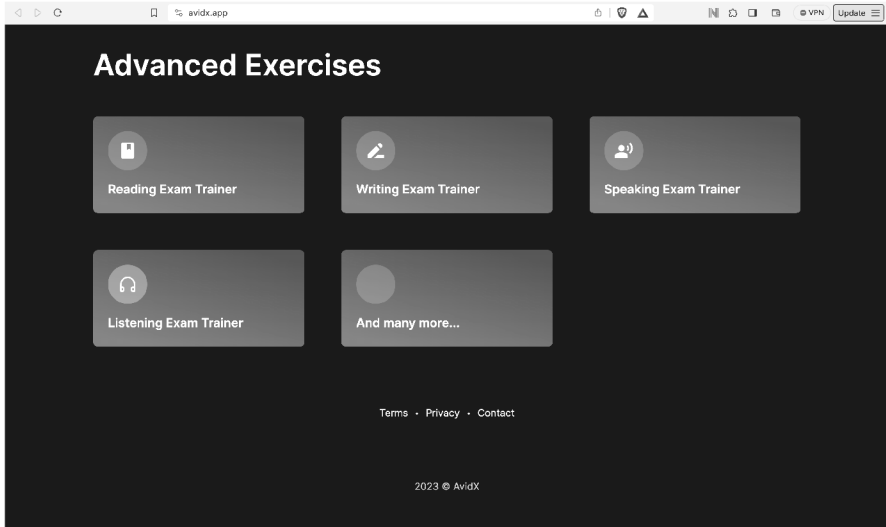
Screenshot 3.4. The “green” state shows that the named checks passed and the branch could be merged without a conflict.

After merge, the deployment system built from `main` and marked the production deployment ready. This was stronger evidence that the integrated version had travelled through the release mechanism, but deployment status still could not substitute for looking at the live result.



Screenshot 3.5. The production deployment is ready and traces back to `main` and the merged pull request.

The final check returned to the product surface. The live page now displayed “And many more...” Observing that state closed the narrow delivery loop for the copy change, while leaving the larger product question—whether the change improved anything users cared about—open.



Screenshot 3.6. The live AvidX landing page shows the deployed phrase, “And many more...”

The sequence creates an evidence chain: the baseline shows the starting behavior; the diff shows code scope; the pull request supplies review context; green checks supply bounded automated evidence; the preview supports pre-merge inspection; the merge integrates the change; the production record connects it to a live artifact; and the live check shows observed behavior. No single link proves that the change created user value.

A pull request is a product conversation

The AvidX example shows why a pull request is more than a code review mechanism. The diff shows what changed, while the surrounding description can connect that change to the ticket, customer problem, design, event plan, experiment, and release decision that gave it meaning. Together they give engineering, product, design, data, and QA one place to inspect whether implementation still matches intent. For PMs working with AI agents, this is especially important because generated work is exposed to the same human and automated challenge as any other contribution.

The conversation becomes weaker when one branch contains too many different decisions. Chapter 20 owns the competitor-ad analyst history; its durable lesson here is that passing CI cannot prove a diff is small and coherent enough for reliable review.

The PM should be able to explain:

- What problem the change solves.
- What behavior changed.
- What was intentionally left out.
- How the change was tested.
- Which risks remain.
- Whether analytics or events were added.
- Whether a feature flag is involved.
- Which environment should be used for review.

If the PM cannot explain the pull request, they probably do not understand the work well enough yet. This is not about pretending to be an engineer. It is about owning the product intent through delivery.

AI can write code, summarize a diff, draft the pull request description, and run a robot review against repository standards. Those accelerators do not answer the product questions the PM owns: whether the change solves the intended problem, behaves correctly, stays within scope, can be measured, and is safe to expose. An agent opening the pull request makes those judgments more urgent, not less necessary.

Continuous delivery means the path to production is practiced

Continuous integration checks whether code can be integrated safely. Continuous delivery extends that discipline by keeping the product in a releasable state and automating enough of the path that an approved change can be deployed reliably. Continuous **deployment** is the stricter practice in which every qualifying change is released automatically. Because teams use these terms differently, a PM should learn the local workflow rather than assume that every green build reaches users.

In many organizations, code that passes integration checks moves through a deployment pipeline into preview or production-like environments, where it can be tested again before release. The key idea is confidence. The team wants confidence that the change works not only on one developer's machine, but in environments that increasingly resemble production.

In the AvidX case, the pull-request preview supported inspection before merge, while the later production deployment was built from `main`. Keeping those records separate made it possible to say precisely what had been proposed, what had passed the repository's checks, and what had actually been released.

A typical path looks like this:

1. Development.
2. Continuous integration.
3. Deployment pipeline.
4. Continuous delivery.
5. User acceptance testing.
6. Production.

The names vary by company, but the stages should form an accumulating evidence path rather than a sequence of ceremonial gates. Local testing helps a contributor find obvious behavioral and implementation problems quickly. Staging reveals integration and environment-specific issues, UAT tests the experience against the intended user and business promise, and production monitoring shows what happens under real use. Treating those environments as interchangeable discards the reason the path has stages at all.

Localhost is where you learn fast

Localhost, or a local development environment, is a version of the product running on the contributor's machine. Its isolation and speed make it useful for changing code, debugging, and checking obvious

paths before asking anyone else to review. For PMs, the value is not the environment itself but the shortened feedback loop between describing a behavior and seeing whether it actually works.

If you ask an AI coding agent to change a registration flow, try it locally before sending it forward:

- Submit a valid email and complete the intended path.
- Submit an invalid email and inspect the error state.
- Check the password requirement and a duplicate username.
- Watch the loading state and read the error copy as a user would.

Localhost is not production, but it is where many preventable issues should be caught. It is also where PMs learn the product system by touching it.

Running the app locally moves the PM beyond screenshots and ticket summaries. They can reproduce the behavior, inspect an error, ask the agent to trace its cause, run a small query, and test an edge case while the context is still fresh. AI agents lower the mechanics of participating in that loop, while direct observation keeps the PM from delegating the judgment about what the product should do.

Staging is where the product gets closer to reality

Staging is an environment that closely resembles production. It is not for all users, but it should be similar enough to production that the team can catch issues that local testing may miss. Staging is useful for functional testing, regression testing, integration behavior, data behavior, and environment-specific issues.

A password reset feature is a useful example. A password reset may work locally, but staging can reveal problems with email providers, environment variables, authentication behavior, links, routing, or account state. Those are not always visible in a local setup. For PMs, staging is where product review becomes more realistic.

At that point, ask:

- Does the feature work with staging data?
- Does the flow behave correctly across roles or permissions?
- Does the email, notification, or external integration work?
- Does the UI still hold up with more realistic data?
- Does the event tracking fire in the expected place?
- Does the product still work after the code was changed to address review feedback?

Retesting after review changes is part of responsible delivery because a code-level correction can still break the intended feature behavior.

UAT is not a ceremony

User acceptance testing, or UAT, is where a designated group tests whether the product works as expected from a user and business perspective. In some teams, UAT happens in a sandbox environment. In others, it happens in a preview environment. The names differ, but the PM role is similar: validate functionality, usability, and fit against the intended requirement.

Consider a mobile banking example. Before a new banking feature reaches all users, designated testers might check account balance display, money transfer, bill payment, and related flows in a controlled environment.

Good UAT is more than “clicking around.” It asks whether the implemented behavior fulfils the product promise for the intended user segment, including whether the flow makes sense, the requirement omitted important states, the interface introduces unexpected usability problems, or realistic data and permissions expose gaps that the specification did not reveal.

For AI-assisted PMs, UAT is especially important because the agent may implement what was literally requested while missing what the product actually needed. The PM has to close that gap.

Production is where trust is spent

Production is the live environment where real users place trust in the product. By the time code reaches it, the team should have more than hope: the work should have passed appropriate checks, survived review, and been tested in the environments that expose its important dependencies. The release plan should also make monitoring and recovery explicit, because production is where incomplete assumptions become user-facing consequences.

Production asks questions that earlier environments cannot fully answer:

- Are real users adopting it?
- Are they dropping off?
- Is performance acceptable under real load?
- Are error rates increasing?
- Are support tickets changing?
- Are customers confused?
- Does the feature create the outcome we expected?

Delivery connects back to discovery in production, where a shipped product bet starts producing evidence.

Release controls are product controls

Delivery is not only about moving code from one environment to another. It is also about controlling who gets what, when, and how safely. This is where release terms become important for PMs.

A **staged rollout** releases a feature to a bounded group—a pilot customer, internal cohort, region, plan, or percentage of traffic—before access expands. A **feature flag** can make that exposure independently controllable, allowing the team to enable or disable behavior without deploying new code. Together, these controls let the team gather feedback, watch operational signals, and contain a defect while the affected population is still small.

Other controls solve different release problems. A **force update** moves users off an application version that can no longer be supported safely, while a **rollback** returns a release to a previous version when the current one causes serious harm. Rollback is not always the most precise response: if risky behavior is isolated behind a flag, disabling it may preserve unrelated improvements in the same deployment.

Versioning identifies which changes and compatibility assumptions belong to a release. **Release notes** explain those changes to users and internal teams, and **deprecation** manages the longer transition away from a feature, API, or behavior. All three affect support, migration, expectations, and trust; they are product work even when the underlying mechanism is engineering-owned.

Taken together, these controls determine how a change enters the market, how risk is staged, how users are informed, and how the team learns after release.

Feature flags change the PM's delivery imagination

Feature flags are worth special attention because they change how PMs can think about shipping. Without one, deployment and user exposure may happen together, turning release into a binary choice between hidden and live.

That was the competitor-ad analyst's release path. It predated the team's adoption of LaunchDarkly and repo guidance for feature flags, so production deployment enabled the feature for everyone. The absence of a staged cohort or independent disable control is one reason Chapter 20 treats release infrastructure as part of the product lesson, not as an engineering footnote.

With a feature flag, those moments separate. Code can be merged while the experience remains hidden, then enabled for internal users, a pilot customer, or a small share of traffic before broader rollout. If the behavior degrades, the team can disable the exposed path quickly; the same mechanism can also support a controlled experiment when assignment and measurement are designed for that purpose.

Flags add complexity and need lifecycle management; a forgotten flag can become technical debt. For meaningful changes with adoption, usability, performance, or business risk, however, separating code merge from full exposure gives the team room to test, pilot, measure, and learn. That control is a particularly useful guardrail when an agent has compressed the time between specification and working code.

Delivery literacy makes specs better

Understanding CI/CD changes how a PM writes specs. A PM who understands delivery does not only write:

Build a new recommendation engine.

They ask:

- How will this be tested locally?
- What does staging need to support?
- Should this be behind a feature flag?
- Who should get access first?
- What event tracking is needed?
- What does success look like after release?
- What failure modes should we monitor?
- What should customer-facing release notes say?
- What happens if we need to turn it off?

The result is not a longer PRD for its own sake, but a sharper brief. The modern spec does not need to be a giant document that tries to settle every possible question before work begins. But it does need to name the objective, expected behavior, outputs, verification, and rollout shape clearly enough that the delivery system can carry it.

AI agents make a sharp release brief more important because they compress implementation time without removing the need to define readiness. Chapter 20 returns to the agent-specific workflow—local checks, robot review, human review, preview, flags, rollout, monitori-

ng, and GTM. The durable foundation here is simpler: no matter who or what produced the change, the team needs a visible, testable path from version control to production.

PM delivery checklist

Before a change moves toward release, a PM should be able to answer:

- What user or business problem does this change solve?
- What is the smallest useful version of the change?
- Which branch or pull request contains the work?
- Which automated checks have passed?
- Where has the feature been tested: local, staging, sandbox, preview, or production?
- What happy paths were tested?
- What failure paths were tested?
- Was the feature retested after review changes?
- Does this need a feature flag?
- Does this need staged rollout?
- Does this need event tracking?
- What should be monitored after release?
- Who needs release notes or internal communication?
- What is the rollback or disable plan?

The checklist is not bureaucracy. It is how the PM protects the quality of the product bet.

Exercise: turn a feature idea into a release path

Take one feature idea from your current product. Write the release path in plain language:

1. What is the objective?
2. What is the smallest buildable version?
3. What code change or product surface is likely involved?
4. How would you test it locally?

5. What should be verified in staging?
6. Who should test it in UAT or preview?
7. Should it be behind a feature flag?
8. Should rollout be staged?
9. What events or metrics should be tracked?
10. What would make you pause, roll back, or disable the feature?

If you cannot answer these questions yet, that is fine. That is the point of the exercise.

Delivery literacy is not about knowing every engineering detail upfront. It is about knowing which questions turn a feature idea into something the team can ship, observe, and learn from.

Evidence and source note

- ⓕ The brief callback to an oversized AI-assisted change reflects the author's delivery experience; Chapter 20 owns the full incident and outcome. The AvidX copy change and accompanying screenshots come from the author's side project and are reproduced for education, commentary, and analysis. The example documents one historical workflow; current GitHub, GitHub Desktop, and Vercel interfaces may differ.
- Ⓣ The registration, password-reset, mobile-banking, and recommendation release paths are illustrative rather than records of one team.
- Ⓢ Git terminology and continuous-integration principles are grounded in the Git reference and Martin Fowler's published treatment.
- Ⓥ CI providers, pull-request interfaces, preview environments, and feature-flag tooling vary by platform and version.
- Ⓟ The delivery checklist and staged-exposure guidance describe product participation. Release engineering, test architecture, incident response, and production authorization remain engineering-owned systems.

Part V: The AI-Augmented PM Operating Model

A validated problem is not yet a shipped outcome.

Between the two sit context, evidence, a bounded specification, a repository full of precedent, a prototype, implementation choices, automated checks, reviews, a preview, rollout controls, instrumentation, customer communication, and the learning that follows release. AI agents can help at almost every point. They can also accelerate ambiguity, reproduce bad code patterns, invent evidence, and produce work that looks complete before it is trustworthy.

The Introduction named the operating chain that holds the book together, and Chapter 1 turned it into a practical responsibility: follow a user need into expected behavior and verify what users actually received. Parts I through IV developed the floors needed to do that responsibly. Part V now returns to the same chain and follows it end to end, with AI agents participating in the work:

signal -> evidence -> decision -> specification -> system behavior -> review -> release -> measurement -> learning

A real Heatseeker feature will carry that chain through this part. Heatseeker AI already worked from customer and internal context. A customer request—and a broader product intuition—suggested that it also needed a bounded competitor-ad analyst: a delegated capability that could scan publicly observable competitor advertisements and compare their strategic angles with the user’s own market context and test evidence.

The purpose was not to critique ad creative or claim access to a competitor’s private performance. It was to help performance marketers and brand marketers interpret the market around them: which problems, jobs, positioning angles, and solution claims competitors appeared to be testing; where several players converged; and where meaningful whitespace might exist.

Chapter 16 establishes the working contract with an agent. Chapter 17 treats the repository as a product surface that communicates rules through code as well as documentation. Chapter 18 preserves evidence and provenance through discovery. Chapter 19 matches prototype fidelity to the learning question. Chapter 20 carries the same feature through the real delivery path.

The PM does not own every step. The PM does own the continuity of product intent: what problem is being solved, what evidence supports the decision, how the behavior will be verified, and what the team will learn when real users meet the result.

Chapter 16: AI Agents Are Teammates, Not Magic Assistants

The first contribution I ask a PM to make in an AI-ready repository does not use an agent.

The PM adds their name to a roster file, carries that one-line change through a branch and commit, opens a pull request, watches the checks run, responds to review, and merges only after approval.

The almost comically small scope is what makes the exercise work. The PM can experience the delivery system before an agent hides it: the branch isolates work, the diff makes change visible, checks enforce known rules, and another person decides whether the contribution belongs. Only after that path is understood do we delegate more complex work.

avidx-app / noted-main

Update ROSTER to include my name #79

Open

 raziibraham wants to merge 1 commit into main from team-os/update-roaster

team-os/ROSTER.md

+1 -1

		@@ -8,7 +8,7 @@ *(maintainer)* `@raz-ii` Product + engineering
8	8	*(code owner)* `@avidx-app` Organization / primary code owner
9	9	*(maintainer)* `@raz-ii` Product + engineering
10		- *(maintainer)* `@raziibraham` Product + engineering
	10	+ *(maintainer)* `@raziibraham` Product + engineering + chief everything officer
11	11	

All checks have passed

5 successful checks

format:check

— formatting is consistent

CI Pipeline

type:check

— no TypeScript errors

CI Pipeline

lint:check

— lint rules pass

CI Pipeline

test:coverage

— tests pass

CI Pipeline

post-checklist

— PR checklist acknowledged

PR Checklist

Merge pull request

Under team-os/, CODEOWNERS leaves the path without a required reviewer, so any teammate can merge. Code outside team-os/ still requires review from @avidx-app.

Screenshot 16.1. The first contribution, reconstructed from the author's Noted team-os/update-roaster branch: a one-line roster edit carried through a branch, a visible diff, passing checks, and a merge gate. Under team-os/, the repository's CODEOWNERS requires no named reviewer, so a teammate can merge; code outside team-os/ still requires review. The scope is deliberately tiny so the delivery path—not the change—is what the PM learns.

This sequence captures the teammate model better than the word *magic*. Tools such as Claude Code, Codex, and Cursor can extend a PM’s reach, but they still need a job, context, constraints, feedback, and review. They can misunderstand, optimize for the wrong thing, or continue a weak assumption with impressive speed.

A useful contribution does not require owning the whole system. It requires carrying the product objective into a technical surface, making a bounded change, testing the behavior, and escalating decisions whose risk or ownership belongs elsewhere.

Start with your own work

Most product managers meet generative AI through personal productivity. They summarize a document, rewrite a message, brainstorm names, or draft a spec. That is a reasonable beginning, but it is only the first useful stage.

In practice, I see responsible adoption move through four stages:

1. Try the tool.
2. Use it to solve a problem in your own work.
3. Use it to help solve a user's problem.
4. Build or adapt a more specialized system when the evidence justifies it.

The stages are not a maturity contest. Not every workflow should end in training or fine-tuning a model. The point is that AI becomes more demanding as it moves closer to users.

When AI helps draft your meeting summary, you can catch an error before sharing it. When AI generates a product response, changes a code path, classifies customer feedback, or decides what a user sees, the quality bar is higher. Context, evaluation, privacy, failure handling, and human ownership become part of the product. PMs should learn responsible delegation on their own work before treating agent output as user-facing behavior.

Discovery jobs and delivery jobs

I divide my agent-assisted work into two categories.

Discovery jobs help the team understand, frame, and decide:

- Synthesize research findings.
- Mine call transcripts.
- Investigate product issues.
- Compare evidence across Notion, Slack, Linear, analytics, and the repo.

- Draft queries and inspect data.
- Generate alternative hypotheses.
- Map user and system flows.
- Create prototypes.

Delivery jobs help the team turn an approved direction into a working change:

- Load a Linear ticket and its linked context.
- Inspect the relevant code and conventions.
- Propose an implementation plan.
- Make a bounded change.
- Add event tracking.
- Run checks and tests.
- Review a diff.
- Prepare a pull request.
- Update documentation and release context.

The distinction matters because the jobs fail differently. A discovery agent can manufacture false clarity. It can combine weak sources into a persuasive summary, flatten contradictions, or turn an isolated customer comment into a theme.

A delivery agent can manufacture false completion. It can satisfy the wording of a ticket while missing the product intent, follow a legacy code pattern, or make a feature work locally without making it shippable. In both categories, the PM must define what evidence and verification are appropriate.

The PM owns the thread

An agent can continue a task. The PM must preserve the reasoning that justifies the task.

The thread includes:

- The user signal that started the work.
- The evidence that made the signal worth attention.

- The decision the team made.
- The behavior expected to change.
- The output the product should create.
- The constraints and ownership boundaries.
- What is intentionally out of scope.
- The signal that will show whether the change worked.
- The condition that would make the team stop.

If the context is weak, an agent continues weak thinking efficiently. “Improve onboarding” can become a beautiful flow that addresses the wrong drop-off. “Add tracking” can become an event with no useful properties. “Fix the bug” can become a patch to the symptom while the underlying state remains broken. The agent’s ability to proceed is not proof that the team is ready to proceed.

Delegation can be product behavior

The Heatseeker competitor-ad analyst makes the teammate metaphor architectural. The product has a main AI experience grounded in the user’s customer and internal context. Competitive-ad analysis is a bounded specialist job. Rather than forcing the main agent to contain every retrieval and analysis instruction, the product can delegate that job to a sub-agent with a narrower role, tool set, and output contract.

Users can enter through a quick action in a dashboard tile. The main agent can also recognize an appropriate intent in a user question and invoke the specialist. Those two paths create different product obligations:

- The explicit button must explain what information is needed and what will happen.
- Intent-based delegation must be predictable enough that users understand why competitor analysis appeared.
- Both paths must preserve the user’s question and scoped Heatseeker context.
- The sub-agent should receive only the context required for the task, not every fact available to the main agent.

- The result must return with sources and uncertainty labels that the main experience does not flatten away.

The minimum input is a competitor name. The product may also offer competitors already identified and stored during onboarding. That convenience creates a provenance question: did the user name this company, did Heatseeker infer it from public business context, and can the user correct the list before analysis?

Sub-agent architecture is therefore not only an implementation pattern. It shapes discoverability, consent, context boundaries, traceability, latency, failure, and the user's mental model of who—or what—did the work.

Prompting is executable product framing

Prompting is often described as a writing skill. For PMs, it is more useful to treat it as product framing in executable form.

A dependable prompt or context brief contains seven parts:

1. **Role:** What perspective should the agent take?
2. **Goal:** What decision, behavior, or outcome does the work support?
3. **Task:** What should the agent do now?
4. **Constraints:** What must it respect or avoid?
5. **Inputs:** Which evidence, files, data, examples, or prior decisions matter?
6. **Output:** What artifact or format should it produce?
7. **Verification:** How should the result be checked?

Consider two prompts. The first says:

Improve our onboarding.

The second says:

Act as a product engineer reviewing a proposed onboarding change.

Goal:

Reduce uncertainty about whether first-time users understand how to

```
create
their first useful document.
```

Task:

Read the linked product brief and inspect the existing zero-document flow.
Propose the smallest local prototype that lets us compare the current blank state with an intent-first scaffold.

Constraints:

- Do not connect new backend behavior.
- Use existing components and design tokens.
- Keep the prototype on a temporary route.
- Do not change the production entry point.
- Clearly label hard-coded and simulated behavior.

Inputs:

- product brief
- current flow files
- design rules
- existing document-creation pattern

Output:

- files likely to change
- prototype plan
- assumptions
- risks
- what the prototype can and cannot test
- questions before implementation

Verification:

Cite the files and rules used. Do not modify code until the plan is reviewed.

The second prompt does more than contain detail. It defines the job, the boundary, and the proof required before action. That is the same work a good product brief performs for a team.

The first contribution should teach the workflow

The opening roster exercise creates an important cost boundary:

- Do small, obvious work manually when delegation would be more expensive.
- Use the agent when it meaningfully reduces search, synthesis, implementation, or verification effort.
- Never use the agent merely to avoid understanding the workflow you are entering.

Technical participation grows from comprehension, not performance.

From short spec to portable context

Long PRDs still have a place in complex, risky, or cross-functional work. But many product changes need a smaller artifact that can travel.

A useful short spec includes:

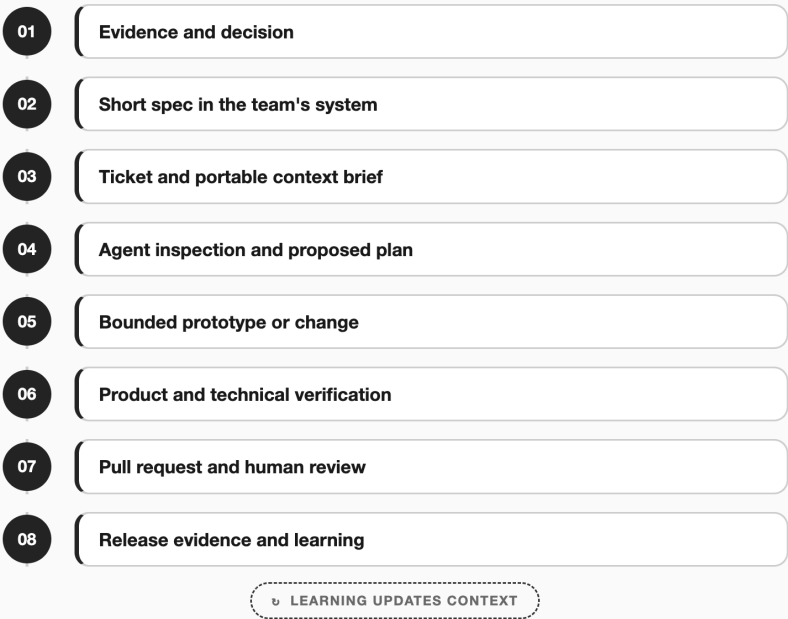
- Objective.
- User problem.
- Evidence.
- Expected behavior.
- Output.
- Constraints.
- Success signal.
- Guardrail or unacceptable outcome.
- Out of scope.
- Open questions.

That spec can begin in Notion, link to a Linear ticket, become input to a prototype, guide an implementation plan, inform event tracking, supply acceptance criteria, and anchor a pull request.

FIGURE 16.1

From evidence to verifiable work

Portable context keeps product intent intact across agents, artifacts, and review.



The useful unit of agent work is a verifiable contribution, not a prompt transcript.

F16.1

Figure 16.1. Agent context should preserve the evidence-to-decision thread and end in human-verifiable work, not an isolated prompt transcript.

The operating model continues.

Chapter 16 continues from portable context into repository rules, agent contracts, review, and verified delivery - then the book carries that work through evidence, prototyping, and shipping.

Read the full 20-chapter field guide

book.raziabraham.com