

AI Accountability Is Already PM Work

Why Project Managers Should Lead the Guardrails Conversation Now

Advisory Article

Ashley Essick, MBA, PMP

Abstract

AI adoption on project teams is outpacing formal governance, and the gap between what is approved and what is actually being used is widening. Shadow AI use, uninformed reliance on unvalidated outputs, and unclear ownership of AI-generated work are quiet risks already present on most teams. Project managers are uniquely positioned to address this as the function that has always owned work visibility, decision accountability, and output ownership. This article makes the case that AI governance is already PM work and offers a practical four-part framework: visibility, guardrails, checks, and accountability. Drawing from experience in highly regulated global clinical operations, the author examines how poor AI use patterns surface, what they cost teams when left unaddressed, and why PMs who begin these conversations now will be the ones leading enterprise AI governance as it matures.

Introduction

The application of AI is changing how teams work, but the ownership of the work is the key aspect of what should never change.

On most project teams, AI has been implemented to draft emails, summarize meetings, clean up trackers, build slide decks and to support with critical thinking tasks to drive to the next step. These avenues of AI workstreams typically first surface as low risk administrative task because these are the easiest to test on. However, as team members become more comfortable with the use of AI implementation quickly moves past the low-risk administrative functions. AI use can quickly move into communication, analysis, planning and decision support.

The silent challenge we see rising is not the use of AI itself but when teams begin using AI without clear boundaries, clear review, or little to no discussion about where it helps versus where it creates risk. Project managers are already strategically positioned to address this issue, not because we should monitor tools, but because oversight of application is already PM work. As a PM, we make work visible, clarify ownership and ensure accountability for the final output.

What Shadow AI Actually Means

The term shadow AI can often seem as an ambiguous amorphous term which feels unmanageable to many PMs. To simplify the matter, shadow AI is simply AI which is being used outside of approved workflows which are normally visible and have established governance but due to the unvetted application may open the workflow and team to unforeseen risk.

Use of shadow AI can surface as someone using an outside tool the team or company has not yet approved, or potentially they are utilizing an approved tool but it is unclear how often it is used, for what application and what level of review of output is occurring. In both cases, the issue is not the tool itself but rather that there is not a clear line of sight as to the application and use scenarios of the tool.

The scale of shadow AI use is becoming more visible as we study how AI systems are impacting the manner in which we work. A 2025 Gartner survey of 302 cybersecurity leaders found that 69% of organizations suspect or have evidence that employees are using prohibited generative AI tools (Gartner, 2025). The same research predicts that by 2030, more than 40% of enterprises will experience security or compliance incidents linked to unauthorized shadow AI. This is not a hypothetical problem for teams but a risk we must be actively addressing with our team members now.

One of the aspects which makes AI radically different from other tools we have implemented in the past is, AI can quickly produce an output which appears professional, complete, attributable, and well thought out. Still, a polished answer can be wrong, shallow or poorly reasoned. Without proper oversight and ownership of AI output the speed of the output becomes the vulnerability which creates risk. It is essential we define who owns the AI output and ensure verification and domain knowledge application before use or application of the output.

Why This Sits on the PM's Desk

I am a PM who operates in the sphere of large scale global clinical trials. In this environment, we are known for early adoption of new technology. However, I am seeing a more conservative approach when it comes to AI and for good reason. It is important for complex and high impact projects to implement AI which has clearly defined guardrails, validated secure confidentiality structures, and accurate attributable outputs.

These risks of AI use often surface as quiet failures due to shadow utilization.

There is a grey area on many teams around what can go into an internal AI tool versus what crosses a line. We are seeing most organizations do not have a defined pathway for the tracking of issues of this nature yet. Many times, we will see newer users of AI tools develop an uninformed reliance on the output because the structure is clear, confident and sounds credible. Unvalidated dependency can lead to misinformation being communicated to stakeholders,

decision being taken without domain experience being applied, and risk surfacing due to unverified application of the output.

AI is both a useful and important tool on our teams. It can help with speed and structure and first drafts, but it does not replace experience, expertise and human judgement. Verification through human-in-the-loop interaction must always be the key steps before an AI output is utilized. Newer users do not always understand the importance of this step, and it is the PMs responsibility to clarify this and hold team members accountable.

The PM's Role To Guide Accountability of Good AI Use

Those team members who apply AI to its best purpose are not those who use the most tools or prompt the fastest. It is those team members who know how collaborate with AI as a tool to support their workflows and understand when to override the output to extract the best quality.

The guidance I give my teams is straightforward; they are the domain experts, and it is their experience, judgement and personal understanding of the situation which makes the output of AI impactful or useless. It is through critical thinking and pushing back on AI outputs that we improve tool use and make the outputs work for us. Blind acceptance of every output AI provides is not proper use and is the equivalent to outsourcing team member expertise.

It is becoming more evident when poor use cases of AI outputs occur. I have seen on teams where work quality becomes detectibly uneven. Strong team members who indiscriminately use AI as a quick solution without validation can produce work which looks strong on the surface but falls in the moment it is pressure tested. Emails are frequently the most obvious area to spot this. When communications are not reviewed for natural language or original thought the message becomes flat and diminishes authority. Slide decks are another method of communication where these occurrences can lead to embarrassing guffaws and unintended miscommunication. The appearance of clean structure but with surface level content do nothing to support real work and decisions being done by team members daily.

One of the most impactful responsibilities I perform as a PM is to coach team members on best use practices of AI in one-on-ones. I inquire how each individual contributor is utilizing AI for high-impact deliverables. I share examples of where AI can help and where it can create risks. I also share deidentified examples of real-world lessons learned. It is important to me, as a PM, to create a safe space for my team members to test their approach with me so we can together build confidence in AI tool usage without creating unnecessary exposure. Many of these conversations have proven to be the most valuable in a work week. I remind my team members that rejecting an output from AI is not a sign that the tool has failed. It is a collaboration between them and the AI tool itself and by pushing back they are creating refined taste for useable output now and in future uses.

It is our role as PMs to ensure work is visible, keep decision paths clean and ensure important outputs are attributable and have human ownership. This means it is our responsibility to engage our team members and ask: What tools are being used? For what kind of work and when? What should never go into those tools? Where is review required? Who owns the final output? These are questions we already utilize when using validated software and other systems. Now is the time to expand the scope of these questions to our management of AI and how we guide our team with safe use and implementation of AI tools. AI is one more domain which will require governance and oversight by the PM.

When the Output Goes Wrong

One instance I have encountered of misuse was on a project I was overseeing. A functional lead working under a complex procedure needed to respond to a contracted vendor question. The team member utilized an enterprise AI tool to support assessment of the question and draft the vendor-facing response. They were acting in good faith, and took action quickly under operational pressure using a tool that was readily available. The failure came when the output was not fully accurate due to the AI tool not having the context which the functional lead had. However, because structure and answer appeared correct on the surface the unvalidated output was communicated to the vendor. The senior review team immediately identified the error and called the vendor to provide corrected guidance, revise the written guidance provided via email and also document the event as a lesson learned for the broader team.

In this instance, the tool being utilized was approved and the team member was trying to do quality work under a time constraint. However, this exemplifies how a generalized output was wrong due to insufficient context and lack of domain knowledge. If this was not caught it could have created inconsistent behavior with the vendor and an audit trail problem that would have been challenging to unwind. This is a case of shadow AI use where lack of visibility and governance can create unanticipated risk. This is why every AI output must be reviewed, verified and confirmed before use both internally and externally.

The Regulatory Landscape Is Already Moving

For teams operating in regulated industries the application of AI use cases is now a regulated governance question which requires PM oversight. The EU Artificial Intelligence Act (Regulation (EU) 2024/1689) is already in play, with a risk-based classification system that brings expectations to how we use documentation, transparency, and human oversight with the use of AI tools. It becomes fully applicable on August 2, 2026, with some provisions taking effect on different timelines (European Union, 2024). In those industries where regulation does not exist, this is an early signal of AI governance expectations which must be taken seriously. Strategic planning now will shape how our enterprises adapt AI and demonstrate accountability.

PMs who understand this direction now will be positioned to lead the conversations which shapes the management and application of AI use.

A Practical Guardrails Model

As PMs we are skilled at implementing manageable frameworks to govern the best use of AI tools. A simple model which works well and can be utilized in four parts:

Visibility. Know where AI is already being used. Ask directly. Make it a normal conversation, not an interrogation. Most teams are already using it somewhere, and you might be surprised what comes up once you start asking.

Guardrails. Set clear rules for what is acceptable and what is not. Confidential information, financial data, proprietary methods, and high-risk work should not be handled casually. The EU AI Act risk framework is a useful reference point here, even outside regulated industries (European Commission, n.d.). Teams need clear direction, not assumptions.

Checks. Require human review for every output. External communication, client-facing material, vendor direction, anything that could affect decisions, timelines, or financial assumptions. All of it should be reviewed before it leaves someone's desk.

Accountability. Keep human ownership explicit. A tool can help produce a draft. A person still owns the final answer.

This framework is functional, actionable and simple. It can be built through training, practical demonstration, and regular conversations about how AI is being utilized on a day-to-day basis. The goal is to make AI adoption safer, more consistent and genuinely useful within specified guardrails.

Do Not Wait for Perfect Governance

At this time, the appetite for AI adoption and implementation by teams is moving faster than formal governance procedure. Advanced users are experimenting with use cases which may not yet be vetted or have appropriate oversight. While leadership teams and PMOs are still evaluating what use cases are, how they want to incorporate AI into workstreams and then how to govern use.

PMs are uniquely situated at the intersection of operational expertise and AI implementation. It is the PM who should start the conversations and surface the need to implement guardrails where they are ambiguous or missing. Thus, providing a bridge between official policy and what teams are actually doing in practice. This helps the enterprise respond before small risks turn into

systemic ones. Structure and best practices should be part of how AI gets woven into daily work and through proactive management the PM has the capability of overseeing this.

PM Action Is Now

AI is likely already a highly integrated tool utilized by your team. Use cases will continue to grow and evolve. As PMs it is our responsibility to have oversight and create accountability for how these tools are utilized. Teams who learn how to use and implement AI in safe meaningful ways will have a genuine advantage. This advantage becomes more consequential when it comes from disciplined, visible, governed use where human-in-the-loop decision making is evident and ownership accountability is attributable to your team members.

As PMs, let us take action now by creating manageable frameworks and enabling safe usage of AI for our teams. By building this capability now and treating AI governance as a core project management task, we will be leading the next phase of how future teams work. This is a leadership opportunity for PMs to embrace. AI is the tool and PM governance is the hand which directs how the tool is implemented, use it well.

Ashley Essick, MBA, PMP, is a Global Project Manager leading multinational clinical trial programs across 23 countries. She has contributed operationally to three FDA-approved therapies and is the architect of an 18-solution enterprise AI portfolio designed for regulatory compliance and human-in-the-loop governance. The views expressed are the author's own.

References

European Commission. (n.d.). Artificial Intelligence Act: Regulatory framework overview. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, 12 July 2024. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

Gartner. (2025, November 19). Gartner identifies critical GenAI blind spots that CIOs must urgently address [Press release]. <https://www.gartner.com/en/newsroom/press-releases/2025-11-19-gartner-identifies-critical-genai-blind-spots-that-cios-must-urgently-address>

Microsoft. (n.d.). Microsoft Purview: Data leaks and shadow AI (overview).
<https://learn.microsoft.com/en-us/purview/deploymentmodels/depmod-data-leak-shadow-ai-intro>

AI Use Disclosure

Generative AI tools were used only for limited editorial assistance, including readability, grammar, and structural refinement. They were not used to generate the article's substantive content, analysis, examples, or conclusions. The author is solely responsible for the final manuscript, including all interpretations and factual accuracy.