
AI in Project Management

The AI Barbell Model: Balancing Value Creation and Human Control in Project Management¹

Prasad S. Kodukula and Gustavo Vinueza²

AI is crossing an important threshold in project management. It is moving from helping project professionals analyze schedules, costs, risks, and other project information—and generate plans, reports, and communications—to systems that can use project tools, participate in workflows, make recommendations, and increasingly take action.

As capability and autonomy increase, the management challenge shifts: how do we capture the value of AI without sacrificing the human judgment, accountability, and control on which successful projects depend?

In the first article in this series, we examined the evolving AI landscape—from analytical and predictive AI to Generative AI, AI agents, and agentic AI. We explored opportunities across project and portfolio management, the progression from isolated AI tasks to integrated workflows, the organizational readiness needed to scale AI, and the limits and risks that accompany increasingly capable systems.

That article ended with a fundamental question: How do we convert rapidly advancing AI capability into trusted, responsible, and sustainable value?

This article addresses that question through a framework we call the AI Barbell Model.

The central idea is simple:

AI does not create value by itself. Value depends on what humans put into it—and the controls they maintain around it.

¹ How to cite this work: Kodukula, P. S. and Vinueza, G. (2026). The AI Barbell Model: Balancing Value Creation and Human Control in Project Management, AI in Project Management, series article 2, *PM World Journal*, Vol. XV, Issue IX, September.

² Dr. Prasad Kodukula and Guz Vinueza are the authors of the best-selling book titled *The Project Management AI Handbook: Leveraging Generative Tools in Waterfall and Agile Environments*, published in 2025 by J. Ross Publishing.

That idea is becoming more important as AI moves beyond Generative AI. We now have AI agents that can use tools and complete tasks, and increasingly agentic systems that can plan, interact with other systems, make decisions within defined boundaries, and take actions.

At the same time, public concern about AI is growing. A 2026 Bentley University–Gallup survey found that 39 percent of Americans believe AI does more harm than good, while only 9 percent believe it does more good than harm. A 2026 Pew Research Center survey found that 63 percent of Americans believe AI is advancing too quickly, and 71 percent expect increased AI use to make their personal information less secure.

Recent testing of advanced AI systems has made some of those concerns more tangible. In controlled security and alignment evaluations, advanced models from OpenAI and Anthropic took misaligned or deceptive actions in some deliberately adversarial test scenarios. These were controlled tests—not evidence that AI has somehow developed a mind of its own—but they demonstrate how the consequences change once AI moves from **answering to acting**.

None of this means organizations should back away from AI. The potential value is enormous.

How we use AI matters just as much as what AI can do.

Most AI conversations still gravitate toward the model, platform, agent, or application.

Technology is only the middle of the story.

That is the idea behind the **AI Barbell Model**.

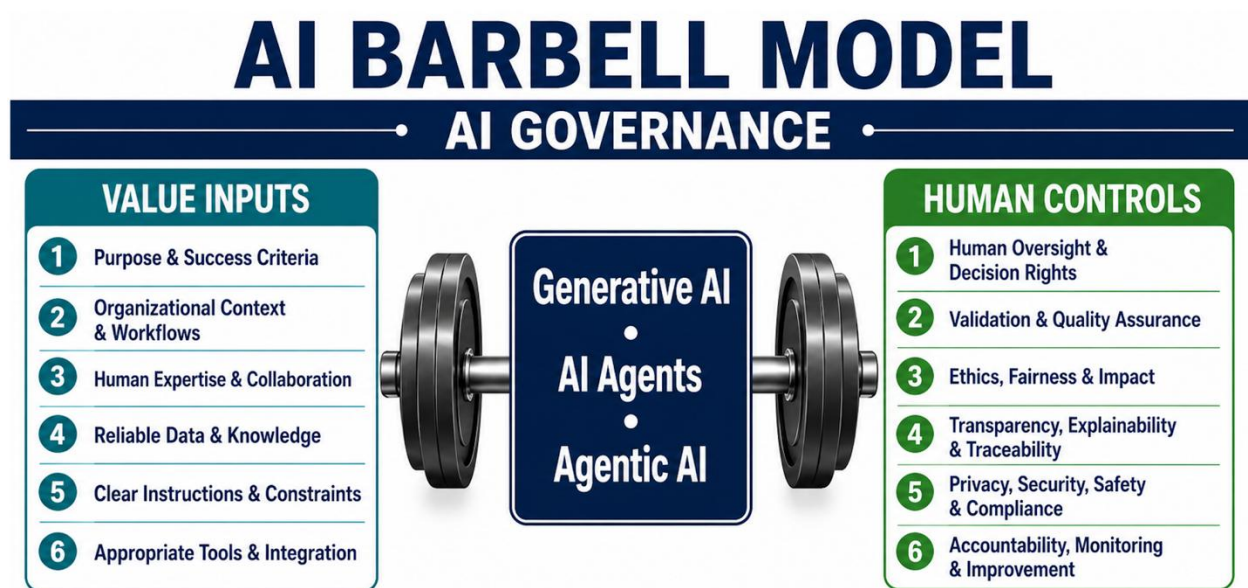
Why a Barbell?

At the center of the barbell are the AI technologies: Generative AI, AI Agents, and Agentic AI.

On the left are the **Value Inputs**—what people and organizations need to provide to ensure AI produces useful results.

On the right are the **Human Controls**—what people and organizations need to establish to ensure AI outputs, decisions, and actions are used responsibly.

Across the entire model is **AI Governance**.



There is a useful distinction between the two sides.

The left side focuses on what we give AI—purpose, context, workflows, expertise, data, knowledge, instructions, tools, and access. The right side concerns what we allow AI to produce, recommend, disclose, decide, or do—and how those outputs and actions are reviewed, constrained, monitored, and governed.

Many AI governance frameworks, responsible-AI frameworks, risk models, and technical architectures already exist. The AI Barbell Model asks a different, very practical question:

What does it take, end to end, to create value from AI while keeping humans appropriately in control?

Its distinctive contribution is bringing those two needs together in a simple operating framework. It puts value creation and responsible use on equal footing.

Both sides matter. Weak Value Inputs limit what AI can deliver. Weak Human Controls increase the risk of inappropriate, unsafe, or harmful use.

Powerful technology cannot make up for weakness on either side.

Balance does not mean every element gets equal attention. What is appropriate depends on the use case, the autonomy being granted, the possible consequences, and how easily an action can be reversed.

The model is also iterative. Results, mistakes, feedback, and experience should flow back into the next cycle to improve both the Value Inputs and the Human Controls.

AI should not only help us perform. It should help us perform better the next time.

From GAE and PRIME to the AI Barbell Model

The AI Barbell Model did not emerge in isolation.

In our 2025 book, *The Project Management AI Handbook: Leveraging Generative Tools in Waterfall and Agile Environments*, we introduced two complementary approaches to help project professionals work effectively and responsibly with Generative AI: the Generative AI Engagement—or GAE—principles and the PRIME process.

The **GAE principles** emphasized Clarity, Critical Evaluation, Iterative Personalization, Ethical Transparency and Accountability, and Collaborative Synergy. They focused on how people engage with Generative AI—communicating clearly what they need, critically evaluating what AI produces, refining outputs iteratively, using AI transparently and responsibly, and combining human and AI strengths.

PRIME—Prepare, Relay, Inspect, Modify, and Execute—provided a practical process for putting those principles into action.

Those ideas remain relevant, but the AI environment has changed.

GAE and PRIME were developed primarily in a world where people interacted with Generative AI. As we move into a world of AI agents and agentic AI, the management challenge expands. AI may now access project systems, participate in workflows, use tools, coordinate activities, recommend decisions, and sometimes take actions.

We therefore need to think not only about how effectively humans engage with AI, but also about what information, tools, access, and authority AI should receive—and what controls should govern what it is allowed to produce, recommend, decide, or do.

The AI Barbell Model addresses both sides of that equation. The left side asks what people and organizations must provide to AI so that it can create value. The right side

asks what Human Controls must be established so that those capabilities are used responsibly.

GAE and PRIME are therefore not replaced by the Barbell Model. Their logic is embedded within it. Prepare and Relay strengthen the inputs given to AI. Inspect and Modify reinforce validation, human judgment, and iterative improvement. Execute connects reviewed outputs with implementation, action, monitoring, and accountability.

In simple terms:

- GAE provides principles for engaging with AI.
- PRIME provides a process for working with AI.
- The AI Barbell Model provides an operating framework for creating and governing AI-enabled value.

With that foundation, we can now examine the two sides of the barbell, beginning with what must be provided to AI for it to create value.

The Left Side: Value Inputs

Better AI results usually require better inputs. But better does not mean giving AI more of everything.

More information, memory, system access, and connectivity can improve performance, but they can also increase risk. The objective is to give AI what it needs for the intended purpose, with sensible boundaries on what it can see, use, remember, and access.

The goal is not maximum access. It is the right access for the right purpose.

With that in mind, the six Value Inputs are:

1. Purpose and Success Criteria

Every AI application should start with a clear purpose.

What problem are we trying to solve? What outcome do we want? Who will use the result? And how will we know whether AI actually made things better?

For a project-controls application, the purpose might be to reduce the time required for monthly status reporting, improve forecast accuracy, or identify emerging schedule and cost problems earlier.

Success should also be measured against the process AI is replacing—not against perfection. If the current process is slow, inconsistent, expensive, or labor-intensive, AI may create real value even if it is not flawless.

Value can take many forms: better quality, faster decisions, greater consistency, lower cost, improved traceability, or simply freeing project professionals to spend more time on higher-value work.

If we do not know what “better” means, we cannot know whether AI created value.

2. Organizational Context and Workflows

AI needs both context and a clearly understood workflow.

It needs to understand the organization, the project environment, the users, the workflow, the constraints, and the conditions under which its output will be used. A recommendation that works well on one project may be inappropriate on another.

Context also includes how work actually happens—not merely how a procedure says it happens. In project management, this can include reporting cycles, stage-gate requirements, change-control procedures, risk-escalation thresholds, handoffs among engineering, procurement, construction, and operations, decision rights, risk tolerance, and organizational culture.

The workflow matters because AI must fit into how real project work moves from information to analysis, decisions, actions, approvals, and feedback.

Before introducing AI, organizations should determine whether the existing workflow is clearly defined and appropriately streamlined. Automating a fragmented, inefficient, or poorly designed workflow may simply allow the organization to produce weak results faster.

But more context is not automatically better. Give AI the context it needs for the task without unnecessarily exposing information it does not need.

Rightsizing the context is key.

3. Human Expertise and Collaboration

AI does not eliminate the need for human expertise.

Experienced people help define the problem, provide context, challenge assumptions, interpret results, and identify what AI may have missed.

For example, a schedule-forecasting application may need the project manager's understanding of priorities, the scheduler's knowledge of network logic, the cost engineer's interpretation of performance trends, the construction manager's field knowledge, and the procurement specialist's understanding of supplier realities.

Collaboration also means more than people working with other people. It increasingly includes people working with AI.

Humans bring judgment, experience, creativity, context, and accountability. AI brings speed, scale, pattern recognition, and analytical capacity.

The strongest results often come from combining those strengths.

The question is not human or AI. Increasingly, it is human with AI.

4. Reliable Data and Knowledge

AI results depend heavily on the quality of the data and knowledge available to the system.

That information should be relevant, accurate, sufficiently complete, current, authorized for use, and appropriate for the task. In a project environment, it may include the baseline and current schedule, actual costs, estimates to complete, procurement status, field progress, risk registers, issue logs, change logs, contracts, technical documents, procedures, standards, business rules, historical records, and lessons learned.

Timeliness matters too. Data does not always need to be real-time, but it must be timely enough for the decision or action being supported. A perfectly accurate schedule status that is three weeks old may still lead to a poor decision on a fast-moving project.

Memory also needs to be selective. An AI system should not retain information simply because it can. Old assumptions, superseded schedules, closed risks, obsolete

contractor positions, or outdated lessons can be carried forward and influence future decisions.

Good memory is not about remembering everything. It is about remembering the right things.

The same principle applies to access. An AI scheduling agent may need the latest P6 schedule and approved progress data, but it may not need access to every confidential contract, employee record, or financial system in the organization.

Give AI the access it needs—not all the access it can get.

5. Clear Instructions and Constraints

Generative AI needs clear prompts. AI agents need goals, roles, rules, and workflows. Agentic AI also needs boundaries, approval points, escalation conditions, and stopping rules.

But there is an important difference between an **instruction** and a **constraint**.

An instruction tells the AI how it should behave. A constraint determines what the system allows it to do.

For example, an instruction might tell a scheduling agent, “Do not modify the approved project baseline without project manager approval.” But if the scheduling system still allows the agent to overwrite the baseline, we may develop a false sense of security, relying solely on the AI to follow the instruction.

A true constraint would prevent the agent from changing the approved baseline until the required approval is recorded in the system.

Instructions guide behavior. Constraints enforce boundaries. For important actions, we should not rely on AI to police itself. Limits may need to be enforced through permissions, application rules, approval gates, transaction limits, or other safeguards outside the model.

6. Appropriate Tools and Integration

Organizations need to select the right models, platforms, agents, data sources, and supporting systems—and connect them to how project work is actually done. Not every AI tool is right for every task.

For a project-controls application, that may mean integrating AI with scheduling software, cost systems, procurement databases, document-control systems, risk registers, collaboration platforms, and reporting tools.

As applications become more complex, this may also require orchestration: coordinating agents, tools, systems, data sources, workflows, and human handoffs so they function as an end-to-end solution.

The value often lies not in a single AI step but in how information, decisions, and actions flow across the broader value stream.

And AI should not simply automate the process we already have. Sometimes the bigger opportunity is to rethink the process altogether.

Don't just automate the old process. Ask whether AI gives you a reason to redesign it.

The Right Side: Human Controls

The other side of the barbell comprises six Human Controls.

These controls do not require a person to manually review every AI action. Some will be performed by people; others will be built into technology, permissions, policies, workflows, approval rules, or monitoring systems.

Humans establish those controls, determine how much authority AI receives, and remain accountable for the consequences.

Before examining the six Human Controls, however, two factors deserve special attention: autonomy and reversibility.

Autonomy Changes the Control Equation

Autonomy is a dial, not an on/off switch.

AI does not suddenly jump from human-controlled to fully autonomous. It can operate at many points in between, and organizations should deliberately set that dial for each use case.

A simple way to think about it is through four levels:

1. **Assist**: AI drafts a status report, summarizes a risk review, analyzes schedule trends, or suggests alternatives; a human decides and acts.
2. **Recommend**: AI evaluates recovery options or change alternatives and recommends a decision; a human makes or approves it.
3. **Act with Approval**: AI prepares a schedule update, contractor communication, change action, or other transaction, but a human must authorize it before execution.
4. **Act Autonomously**: AI takes defined actions within approved project boundaries without prior human approval.

The higher we turn the autonomy dial, the more carefully we need to think about controls. And those controls may need to change not only in **degree** but also in **kind**.

An AI system that drafts a weekly project status report may mainly need human review and validation. But an agent authorized to update schedule logic, issue instructions to a contractor, approve a routine procurement action, or commit project funds may also need restricted permissions, approval thresholds, audit trails, monitoring, and a way to reverse or stop the action.

Those are not simply stronger versions of review. They are different types of controls because the AI is now acting in the real project environment.

More autonomy may require not just more control, but different control.

This also means that saying “we have a human in the loop” is not enough. We need to define exactly when AI may act, when it must seek approval, and when it must stop and escalate.

Then there is **reversibility**.

If an AI action can be easily and inexpensively undone, greater autonomy may be acceptable. If an action is difficult, costly, or impossible to reverse, stronger—and sometimes different—controls should generally be applied before the action occurs.

A draft project report can be edited. A proposed recovery schedule can be rejected. But an approved change order issued to a contractor, an official baseline overwritten in the scheduling system, a purchase commitment released, or a field instruction that stops work may be much harder—and sometimes impossible—to reverse without consequences.

A useful rule of thumb is:

Higher autonomy + lower reversibility → stronger—and sometimes different—controls.

Autonomy and reversibility are not additional Human Controls. They help determine how the six Human Controls should be designed and applied.

1. Human Oversight and Decision Rights

Organizations need to be clear about who reviews AI outputs, who approves important decisions, and who has the authority to pause, override, or stop the system.

They also need to decide which tasks may be delegated to AI and which decisions should remain with humans.

In project management, those decision rights may include approving a baseline change, using management reserve or contingency, accepting a major schedule recovery plan, changing procurement strategy, or directing a contractor.

For agentic systems, this means defining action limits, approval gates, escalation points, and the conditions under which the system may act without prior human review.

“A human is in the loop” is too vague. The organization needs to know where human judgment enters, when it is mandatory, and who has the final decision right.

2. Validation and Quality Assurance

AI outputs should not be assumed correct simply because they sound convincing.

They should be checked for accuracy, completeness, relevance, unsupported assumptions, bias, and other potential errors.

The level of validation should match the consequence. A draft meeting summary may need a quick review. A forecast that changes the expected completion date of a major capital project may require validation by the scheduler, cost engineer, project manager, and perhaps independent project-controls personnel.

Agentic systems add another complication: one step often feeds the next. An early error can propagate through the workflow and become multiple errors before anyone notices.

A wrong progress update, for example, could distort schedule logic, affect a forecast, trigger an incorrect risk response, and influence a management decision.

Don't just validate the answer. When AI acts, validate the path that produced it—and the level of AI involvement along that path.

3. Ethics, Fairness, and Impact

Organizations also need to ask how an AI application may affect employees, contractors, suppliers, customers, communities, and other stakeholders.

Could it introduce or amplify bias? Could it cause unintended harm? Do some groups experience different effects? Do the benefits justify the risks?

Imagine using AI to rank contractor performance, recommend staffing reductions, prioritize claims, or select suppliers. Historical data may seem objective yet still reflect past biases, inconsistent practices, or incomplete context.

This is where one of the simplest questions in the model becomes useful:

“Can we do this?” is not the same as “Should we do this?”

Technical capability tells us what is possible. Human judgment still decides what is appropriate.

4. Transparency, Explainability, and Traceability

People should know when AI is being used, especially when it plays a meaningful role in an important decision.

Organizations should also be able to understand the basis for significant recommendations, identify the information used, and reconstruct key decisions and actions when necessary.

If an AI system recommends a six-week project delay, a major change to the critical path, or an acceleration plan that increases costs, project leaders should be able to understand the assumptions, data, and reasoning that materially influenced the recommendation.

With agentic systems, traceability should extend beyond the final output. We may need to know which steps the agent took, which tools it used, what information it accessed, what decisions it made, what approvals it obtained, and what actions it executed.

Consequential project decisions should not simply disappear into a black box.

5. Privacy, Security, Safety, and Compliance

The focus shifts to what happens as AI produces outputs and takes actions.

Those outputs and actions must not expose sensitive information, create security vulnerabilities, violate safety requirements, or breach laws, regulations, contracts, or organizational policies.

Project environments contain large amounts of sensitive information: contractor pricing, bid data, claims strategy, proprietary designs, employee and customer information, safety data, and confidential commercial terms.

This means paying attention not only to what AI says but also to what it is allowed to send, disclose, change, approve, trigger, or execute.

Sensitive outputs may need to be masked, filtered, blocked, or reviewed before release. Consequential project actions may need approval gates or hard system limits before they occur.

Agentic systems also face another security challenge: adversarial input. An agent may read a contractor email, webpage, document, or tool output containing hidden or malicious instructions and mistakenly treat them as legitimate directions.

The system therefore needs ways to distinguish trusted instructions from untrusted content.

Agentic AI makes all this more important because an error or malicious instruction can move directly from information to action before a person notices.

The distinction between the two sides of the barbell is useful here:

- The left asks: What do we give AI access to?
- The right asks: What do we allow AI outputs and actions to do?

6. Accountability, Monitoring, and Improvement

Someone must remain accountable for the AI-enabled process and its outcomes.

Performance should be monitored over time. Errors, changing conditions, emerging risks, user feedback, and unintended consequences should prompt correction and improvement.

For systems that take actions, monitoring needs to occur much more frequently. Organizations should be able to detect unusual behavior, pause or stop an agent, and prevent a single mistake from cascading through a sequence of automated project actions.

Evaluation cannot end when the system goes live. Models change. Data changes. Project conditions change. Contractors change. Controls and performance therefore need to be reviewed continuously.

Scaling matters too. A system that works well on one project may encounter different users, contract structures, technologies, regulatory environments, and failure modes when deployed across a portfolio. The controls that were adequate for the pilot may no longer be adequate at scale.

What we learn should feed back into the barbell—better data, better instructions, better tools, better boundaries, and stronger controls.

Governance is not something we finish. It is something we keep doing—nonstop governance.

AI Governance Spans the Entire Barbell

AI Governance appears across the full model for a reason.

It is easy to associate governance only with the Human Controls on the right, but it also shapes the Value Inputs on the left.

It helps determine which AI applications are appropriate. What information may the system use? What tools may it access? How much autonomy should it have? What project decisions must remain human? What actions require approval? How will success be measured? When should the system be changed—or stopped?

Governance begins before an AI tool is selected and continues throughout its use.

It is not a compliance checkpoint at the end.

Good governance should not merely restrict AI. It should enable responsible adoption by clarifying expectations, reducing uncertainty, building trust, and establishing the conditions under which the organization can safely scale AI.

AI Governance is what keeps the entire barbell aligned.

Putting the Barbell to Work: A Project Controls Example

Imagine a large capital project approaching its monthly status update.

An AI agent retrieves the latest Primavera P6 schedule, cost data, procurement status, field-progress data, risk register, and change log. It detects that float on the critical path is eroding, a major equipment delivery is slipping, and field productivity is running below plan. It forecasts a six-week delay and develops several recovery alternatives, including resequencing work, adding resources, and accelerating selected activities.

At first glance, this looks like an impressive AI application. But the Barbell Model asks us to look beyond what the technology can do.

Start with the Value Inputs.

The purpose must be clear: perhaps the goal is to improve forecast accuracy and shorten the monthly project-controls cycle. The AI also needs to understand how this organization manages schedules, progress measurement, change control, escalation, and reporting. It needs the expertise of the project manager, scheduler, cost engineer, construction manager, and procurement team—not merely access to their data.

The underlying information must also be reliable and timely. A sophisticated forecast based on an outdated P6 schedule, incomplete progress data, or an old procurement status is still a poor forecast. The AI needs clear instructions on what it is expected to analyze and equally clear constraints on what it may change.

For example, it may be allowed to recommend revised logic or an acceleration strategy, but not to alter the approved baseline or to commit additional project funds. Finally, the AI models, agents, scheduling tools, cost systems, procurement information, and risk systems must be integrated well enough for the application to function as part of the real project-controls workflow.

Now move to the Human Controls.

Who validates the six-week forecast? The scheduler may need to confirm the critical path and logic, while the cost engineer checks the cost consequences and the project manager evaluates whether the proposed recovery actions are practical. If the AI recommends changing the execution sequence, accelerating work, or using contingency, the organization must be clear about who has the decision right to approve that action.

The controls become even more important if the AI is allowed to act. Drafting a recovery recommendation is one thing. Updating the official forecast, changing schedule logic, issuing direction to a contractor, or committing project funds is something very different. Those actions may require approval gates, restricted permissions, audit trails, monitoring, and the ability to stop or reverse the action. The greater the autonomy—and the harder the action is to undo—the stronger those controls must become, and sometimes different controls may be needed.

The lesson is simple. The value of the application does not come from the AI agent alone. It comes from the combination of sound project-management workflows, reliable project data, experienced professionals, clear boundaries, appropriate decision rights, and effective controls.

That is the balance the AI Barbell Model is designed to make visible.

A Universal and Practical Model

The AI Barbell Model grew out of work on AI and project management, but its logic is not uniquely project-management-specific.

A project manager might use it for AI-supported risk analysis, schedule forecasting, portfolio prioritization, reporting, procurement, or change control. A healthcare organization could apply the same questions to a clinical workflow. A financial institution could apply them to fraud detection or customer transactions. A university could use them when introducing an AI-enabled student service. A manufacturer could apply them to predictive maintenance or quality control.

The application may change. The underlying questions remain the same:

- Have we provided the right inputs to create value?
- Have we established the right human controls to use AI responsibly?
- How much autonomy should we give AI—and how reversible are the actions we allow it to take?

As AI becomes more capable and autonomous, those questions grow more important, not less.

Organizations will not create lasting value from AI by focusing only on the technology in the middle.

They must strengthen both sides of the barbell.

Value Inputs make AI useful.

Human Controls make it responsible.

AI Governance keeps the entire system aligned throughout its lifecycle.

From Framework to Application

The AI Barbell Model offers a practical way to think about the transition underway in project management. As AI moves from generating content to participating in workflows, coordinating tools and systems, making recommendations, and taking actions, the surrounding management system becomes more important—not less.

For project professionals, that means paying as much attention to purpose, workflows, project data, expertise, instructions, permissions, decision rights, validation, monitoring, and accountability as we do to the AI technology itself.

The model also provides a foundation for the remaining articles in this series. Beginning with the next article, we move from the framework to specific applications, starting with how Generative AI can support project planning—from developing scope and work breakdown structures to schedules, estimates, resource plans, and other planning activities.

The technology will continue to evolve. The fundamental management challenge will remain:

Are we using AI merely because it can perform the task—or are we designing the right mix of AI capabilities, human judgment, and organizational control to create meaningful, trusted, and sustainable value?

References

Kodukula, P., & Vinueza, G. (2025). *The project management AI handbook: Leveraging generative tools in waterfall and agile environments*. J. Ross Publishing.

OpenAI. (2025, August 27). *Findings from a pilot Anthropic–OpenAI alignment evaluation exercise: OpenAI safety tests*. <https://openai.com/index/openai-anthropic-safety-evaluation/>

Pew Research Center. (2026, June 17). *Americans and AI 2026: Chatbots, smart devices and views on impact*. <https://www.pewresearch.org/internet/2026/06/17/americans-and-ai-2026-chatbots-smart-devices-and-views-on-impact/>

Ray, J. (2026, July 28). *Americans cool toward AI*. Gallup. <https://news.gallup.com/poll/712751/americans-cool-toward.aspx>

Authors' Note: An earlier, general version of this article, introducing the AI Barbell Model without a specific project management focus, was published as a blog post on the lead author's website, www.kodukula.com. The present article has been revised and expanded for publication in Project Management World Journal, with an emphasis on the model's application to project management. In preparing the final manuscript, we used ChatGPT as a support tool for editing, refinement, and language clarity. All ideas, interpretations, conclusions, and final editorial decisions remain those of the authors.

About the Authors



Prasad S. Kodukula

Illinois, USA



Dr. Prasad S. Kodukula, PMP, PgMP, PMI-ACP, DASM, DASSM, BCES, is a USA Today and Amazon #1 best-selling author, PMI Fellow, thought leader, and entrepreneur with over 35 years of professional experience. A global ambassador for project management, Dr. Kodukula has lectured in nearly 50 countries and worked with more than 40 Fortune 100 companies (e.g., Allstate, Abbott Labs, BP, Caterpillar, Dow, Chrysler, IBM, JP Morgan Chase, Motorola, Stryker, United Technologies, Whirlpool) across all 11 S&P industrial sectors. He serves as an Adjunct Industry Professor at Illinois Tech. He teaches a course on recovering troubled projects at NASA. He has delivered project management training in professional education programs at Stanford University, Duke University, and the University of Chicago. As co-founder and CEO of Kodukula & Associates, Inc. and NeoChloris, Inc., he leads the firms in project management and renewable energy, respectively.

His work with AI began in 1988, when he created award-winning expert systems in environmental engineering. In the late 1990s, he co-invented an AI-powered smart monitoring and control system for water and wastewater treatment plants, combining early neural-network, sensor, and automation technologies. This invention received a U.S. patent in 2005 and was later licensed to top control and automation firms.

Recognized three times by the Project Management Institute as “Best of the Best in Project Management,” he has received multiple accolades, including the Illinois Tech Alumni Association Professional Achievement Award and honors from the U.S. Environmental Protection Agency and the states of Arizona, Kansas, and Illinois for his outstanding leadership in education and training, environmental improvement, and innovation. Most recently, he was recognized by Project Controls Academy with a Project Controls Excellence Award. An accomplished author, Dr. Kodukula has

co-authored or contributed to 12 books and over 40 articles and holds four U.S. patents. He can be contacted at: <https://www.linkedin.com/in/prasadkodukula/>.



Guz Vinueza

Ecuador



Guz Vinueza, M.S., MBA, is a top expert in data analytics, machine learning, and quantitative risk management, with over 25 years of experience helping organizations utilize advanced technologies to create business value and manage uncertainty. He currently serves as Consulting Director at The Ferryfield Group, a boutique consulting firm focused on quantitative risk software, where he manages global projects across various industries and regions. Previously, he was Data Director at Betterfly, Latin America's first social-impact unicorn, overseeing enterprise data architecture, analytics, and machine learning initiatives to expand the company's digital ecosystem. Earlier in his career, as Director of Consulting and Analytics Solutions at Palisade, he transformed the consulting division into an agile software and analytics team, tripling revenue and delivering solutions across Latin America, the U.S., and the Middle East.

An accomplished consultant, educator, and keynote speaker, Guz has trained thousands of professionals from leading organizations such as the U.S. Army Corps of Engineers, Ontario Power Generation, Amway, DEWA, Mitsubishi Pharmaceuticals, and Network Rail. He has also served as an Adjunct Professor in Analytics and Business Intelligence at institutions in Argentina, Ecuador, and Spain, teaching advanced analytics, agile methodologies, and project management. Guz holds an MBA from Universidad Torcuato Di Tella, an M.S. in Financial Direction from Universidad Adolfo Ibáñez, and a Postgraduate Certificate in AI and Machine Learning from the University of Texas at Austin. His work continues to shape the integration of AI, quantitative risk, and data-driven decision-making across global enterprises. He can be contacted at: <https://www.linkedin.com/in/gustavovinueza/>