



Algorithmic Governance and Trust Calibration in Project Management

Syed Osman Zulnorain

osmansyed015@gmail.com

Received: 08 Jun 2026; Received in revised form: 05 Jul 2026; Accepted: 10 Jul 2026; Available online: 15 Jul 2026

Abstract— Today, AI and predictive analytics can obsolete the primitive “take a stab at it” attitude towards project risk management. These technologies ensure that the businesses are aware of building projects which may be delayed or cost overruns in advance. Most of the existing research however focuses on enhancing the math’s and algorithms and has not studied the question of trust and usage of this data amongst real project managers (PMs) working in their day to day role. This paper does so. We combine the Technology Acceptance Model (TAM) with Sociotechnical Systems (STS) theory to describe the human aspect of the AI tools. We argue that there are two problems with an out-of-balance manager’s trust: either the manager appeals to the wrong data (automation rejection) or the manager hands over to the machine, believing that the machine is always in the right (automation bias). Both the errors negatively affect the successful completion of a project.

Keywords— Artificial Intelligence (AI); Project Risk Management; Trust Calibration; Algorithmic Governance; Sociotechnical Systems (STS); Technology Acceptance Model (TAM); Automation Bias; Explainable Artificial Intelligence (XAI).

I. INTRODUCTION

Project management is a fast changing part of our approaches. Project managers used to primarily guess, estimate and rely on their experience on where a project might face challenges using a manual Gantt chart and old Excel spreadsheets. Now it is different, however. Predictive software and machine learning models are increasingly frequently used in Project Management Offices (PMOs). These technologies are designed to analyse the data and warn the manager in advance about any activities that can be delayed, any areas where costs might go up and how resources can be managed.

But the real problem comes to light when these technologies are deployed in the businesses day to day operations. These algorithms are designed by data scientists, who think if they are set up correctly, the project manager will follow their advice. It’s not what people do in real life. I’ve regularly observed that project managers tend to stick to their own

outdated methods and disregard obvious software warnings. But others are doing just the opposite – they stop paying attention to what’s actually happening on the ground, because they are relying so heavily on that “management dashboard” within the software.

This disconnect between what the manager is actually doing, and what the computer recommends isn’t entirely a technical issue. Just a matter of trusting; that’s it! When a danger warning is displayed to a PM and the software, sans warning text, functions like a “black box”, giving only an alarming warning, the management becomes confused. They stop using the expensive software, lose confidence in the system and revert on the old practices! The Company’s huge capital investment goes to a loss. If the management only follows the machine blindly, however, it’s a pity that one piece of information is wrong or one datum is entered incorrectly in the system, the whole project will be jeopardized.

This human eye is totally missing from the existing world scholarly literature. All technical publications are simply intended to slightly enhance accuracy of the algorithms and coding. On the other hand, technology is talked about in such a broad way in management studies that it views a highly complicated AI risk tool as a new corporate email system. They have no clue what stress a PM is under when cash is a problem, deadlines are getting close and senior executives are looking for answers. A manager must always make a balance between his/her professional experience and data of a computer. In this paper, those pieces will be tied together to develop a straightforward model that companies can truly apply.

II. LITERATURE REVIEW

The Evolution of Algorithmic Risk Assessment in Project Environments

In the history of project management, risk has always been considered as a simple mathematical problem. Risk = Probability $\times$ Impact ($\$P \times \I) has been a formula used for a long time in traditional books including in the PMBOK guide. The computer's part in control of the system was still quite limited. It could only be used to store numbers typed by humans, and had the ability to graph these numbers in a colorful pie chart or bar graph. The machine had to be programmed.

But this setup has completely changed with the use of machine learning and predictive analytics in the PMO. AI systems do not just await inputs today! They are actively involved in activities. The track how fasts the devs are coding, how much people talk about on the team, and also check for updates on the supply chain and send their own alerts if something is going astray on the critical path. In other words, AI is no longer just a passive calculator but rather an entity that guides the management regarding their project.

Therefore, the position of Project Manager has been totally transformed. One of the things the management had to do was to spend time determining the hazards. These days, the computer does the job of identifying the dangers and the only thing that the manager is expected to do is to decide the validity of the computer's prediction. This introduces a whole new form of Mental stress that

have no remedies in the old Project Management training.

Theoretical Foundations of Trust in Automation

To grasp fully the process by which a project manager decides to pay heed to or ignore a software forecast, we must transcend the simple technology adoption models. The Technology Acceptance Model (TAM) was frequently used for research by corporate researchers for a long time. Based on this, employees will put into practice the use of software if they consider it easy to use and beneficial. So this doesn't work in a "time is money" project scenario – not even when everything goes perfect with a simple new digital time sheet or a new internal business chat application. The wrong decision on a massive project can cost millions in dollars, delay the launching of important projects or lead to the departure of the manager.

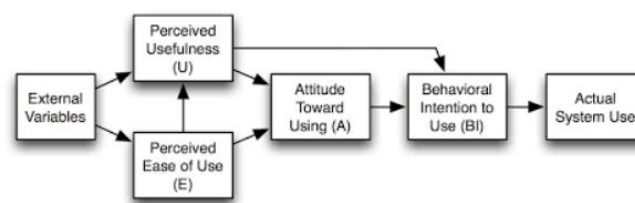


Fig.1: The classic Technology Acceptance Model framework (Davis, 1989) mapping behavioral intention layer. Source: Wikipedia

Asking someone if the software interface is easy to use or looks good won't be enough to anticipate how they will act under such strain that could destroy their career.

- The actual over time management of the peoples' faith in a given automated system needs to be examined. Based on a widely cited Trust research by Lee and See (2004), humans' confidence in robots is not an on/off switch. It consists of three different layers and it changes all the time:
- Part trusting attitude is innate, or dispositional trust, which means that people have a pre-existing level of trust in technology when they switch on their computers. It just depends on how that person was raised, and his or her personality and level of intrinsic interest or fear of new technologies.

- **Situational Trust:** This is a manager's trust in a tool at a particular time based on his or her environment. For example, a PM could be bogged down in the review of an AI alert when a project is progressing smoothly. But things can change just as fast because of stress if the entire project is a complete mess, and the due date is 2 hours away.
- **Learned Trust:** Trust based on life experiences. As the management use it week after week, they maintain a mental diary of the software's correctness and the stupid goofs that it contains. As time goes by, they come to understand how reliable the machine is.

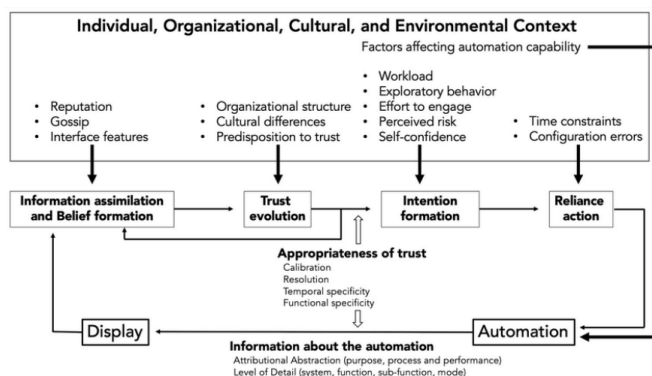


Fig.1: Conceptual model of human trust evolution and reliance actions in automated environments (Adapted from Lee & See, 2004).

Trust calibration is, in our theory, the closeness of an AI tool's true accuracy in real world and mental trust a manager has in an AI tool. It's a great combination when these two things go hand in hand. But, if you're still no match, there are two negative results. Management has become automatized and failed to critically assess the situation, followed the faulty readout and thus failed the project. Otherwise, they suffer from automation rejection, go into a defensive mode, ignore actually accurate information and work on a difficult project with their own outdated methods.

Sociotechnical Systems (STS) and the Black Box Dilemma

The project manager - AI clash goes beyond a psychological issue. There is a structural problem with the company. This can be explained referring to the Sociotechnical Systems (STS) theory. This idea is

based on the belief that a company can only perform its best when the social (human feelings, relationships and corporate politics) and technological (software, tools and data) sides are ideal and at the same time. Without taking into account how people on the ground run the show and take accountability, the system fails to hold if you pursue rapid technological advancements.

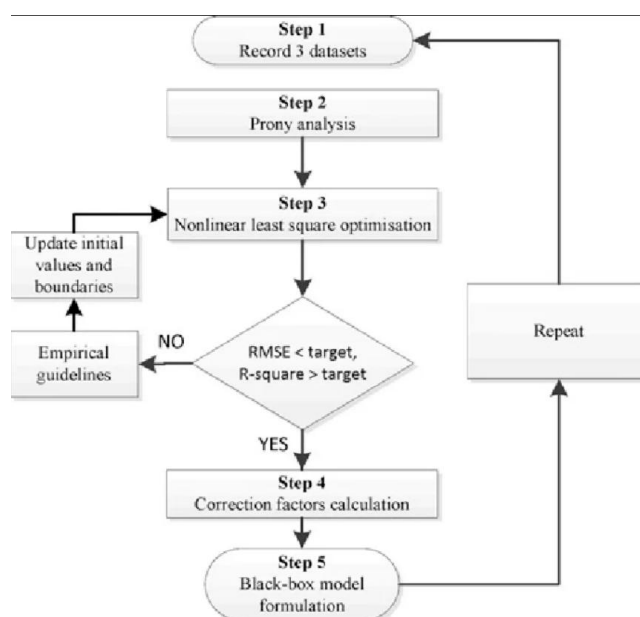
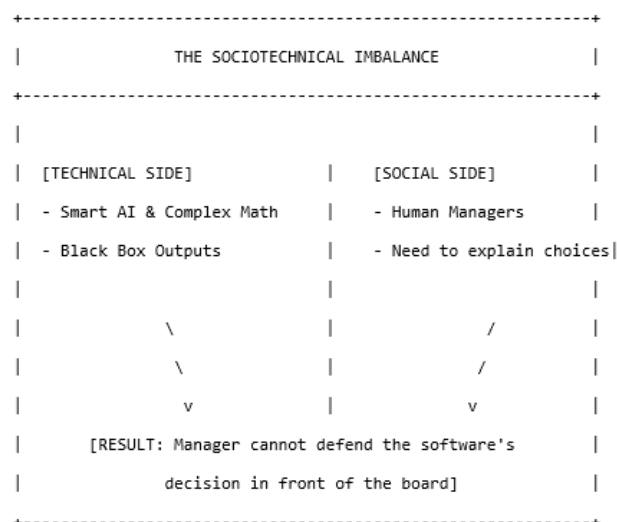


Fig.2: Methodological workflow of a non-linear optimization routine within data-driven black-box modeling architectures.

This is exactly what's happening in modern company workplaces. Sophisticated machine learning models, working on complex math's, are being bought by businesses to calculate project hazards. There is

however a large “black box” problem for these technologies. In data science, there is a elementary tradeoff: the more intelligent and precise a model is at predicting unpredictable delays on initiatives, the extra advanced its math will get. Because of the software’s complexity, it cannot give a simple, step-by-step explanation to humans of how and why it came up with its decision. All it does is display a warning score on the screen.

This means that there is a political problem for the project manager. There’s always a PM that’s on call to steering committees, clients and top executives. If a customer’s major launch is scheduled for an upcoming week and sudden computer trouble emerges, a manager can’t just inform the board of directors, “We’re postponing the project, since the computer screen turned red. “Why? The board will start its questioning by asking “Why?

When the manager answers, “uhuh thing goes out of the machine said so” then he looks ‘can’t handle it’ himself and is 100% liable if anything goes wrong. Since they can’t tell anyone else how the machine will work, managers will choose to ignore the program and to discredit the unsolicited advice given by the machine so as to protect their own jobs and reputations.

Detailed Explanation of the Conceptual Framework Blueprint

We explore the exact ‘mental clock’ in which a project manager operates when using an AI Tool, to make this easy to understand. The framework unpacks the whole process into three distinct steps which take place in order and which are logical, not to suggest a nebulous concept of trust:

[Step 1: The Setup] -----> [Step 2: The Mental Filter] -----> [Step 3: The Action]
(System, PM, & Project) (How Trust is Formed) (Bias vs. Good Work)

Step 1: Setting Up (Structural Antecedents)

This is the situation as soon as the manager steps in the software platform. Comprising of 3 types of components: project side (how disorganized/costly/messy the actual project is), human side (years of expertise and background of PM) and system side (how transparent/hidden the AI math is).

Step 2: The Mental Filter (Mediating Trust States)

You are not immediately acted upon when you get a warning of risk from the AI tool. First it filters through the manager’s filters and filters.

In an open system when the manager is able to articulate their rationale, then the manager’s Learned Trust is enhanced.

Manager’s personal dispositional trust is a function of personal experience, age and the nature of the person.

If the endeavor becomes quite difficult or unpleasant then their Situational Trust is temporarily jolted/replaced by past experiences.

Step 3: The Action (Behavioral Outcomes)

These Mental Filters are what they will ultimately need to combine to create the real business result. If all this is well-balanced, the manager in charge achieves what is called Calibrated Reliance – taking full advantage of the machine and intervenes with human common sense when it is appropriate. When filters are distorted, however, the manager will either sink into the rut of Automation Rejection—whereby he/she turns off or disengages from the program altogether – and relies solely on intuition—or Automation Bias, where his/her attention wanders and he/she leaves the computer to make all the decisions.

Literature Synthesis and Identifying the Gap

When you review the currently available literature of scholars on the subject, you can notice that each group of scholars is divided into separate entities and do not communicate with each other.

A major part of the traditional project management studies are risk logs and simple arithmetic formulae. They think that managers are soulless machines that always take the right and logical decision. Technology adoption studies (e.g., TAM or UTAUT) overlook the degree of professional and career risk managers face in the corporate world when dealing with decisions concerning new technology and only focus on a software’s “funness”. Finally, although studies of automation trust have done a wonderful job of studying how people trust machines, almost all of the research has been done on how nuclear power plant engineers trust machines and how autopilot users trust aero planes.

What is missing in our current research is how – in an active way – a manager’s faith in an AI risk tool is undermined by the tensions and political pressures of

working within a project environment. This study is designed to fill this very target gap.

III. RESEARCH METHODOLOGY

Research Design and Philosophy

The research design that is used in this research is a quantitative based on a conceptual approach. Given that using AI for project risk governance is still a rather new concept in the corporate world, it would be illogical to dive straight into a large-scale quantitative survey replete with figures & data. Moving parts must be well defined before quantifying things with numbers can take place. Consequently, it is our approach to meticulously deconstruct established, tried-and-tested concepts from other disciplines (e.g., psychology, cognitive automation and organizational design) and adapt them for use in everyday life.

Variables are not arbitrary! Instead, we surveyed issues in user practice (why does a boss get really mad about a dashboard for example, or ignore an alarm displayed by a computer), and correlated them to known concepts in academic scholarship. This will help us build a stable and sound structure and make it easily replicable by future scholars via survey or business studies.

Literature Selection Strategy and Rules

To make sure that this system is based on solid foundations, we looked through prestigious academic sources such as Google Scholar and Scopus. Aside from research into the current state of things (up until 2026) showing how companies are actually using AI in the workplace, we looked for old school human psychology books.

In order to keep the paper contained and focused, and avoid a generic paper on computers, we established very strict measures on what to include and what not to include:

What we included: peer-reviewed academic publications specifically related to behaviors in project management, the influx of Explainable AI on business decisions and people's trust in automated systems.

What we didn't include: Articles that are purely technical in computer science, so there's only discussion about the math or the code or tweaking the hyperparams – never run it on real live human users.

In addition, biased whitepapers generated by AI vendors trying to sell their products and software ads were not included.

IV. RESULTS: THE CONCEPTUAL FRAMEWORK

These key elements need to be accurately specified if it is to be of service in future studies. Here we think of trust as a real connection (quantifiable) between actual maximalist behavior and software configuration. The model consists of three different parts: the inputs, the mental filters and the results.

Specification and Definition of Variables

If a researcher chooses to develop a survey or test this model in an organization, he must have a good grasp of exactly what these variables represent in real-life working situations:

- Algorithmic Opacity (\$O_A\$): How much of the workings of the AI tool is the project manager not able to see. If the opacity of a tool is low (transparent) a short description of the cautions is clearly visible. If it is set to high opacity (all black), then it displays a danger percentage, or a red color on screen, without displaying any mathematical processes.
- \$E_D\$ - PM Domain Experience is the domain name devoted to the human factor. It tracks the manager's amount of time, skills and experience in project paradigms (Agile or Waterfall), or with project crises.
- The environment variable is Project's Structural Complexity (\$C_P\$). Measures the extent of the project, number of teams involved, limited funding, rate of change of project goals and customer requirements etc.
- The main mental filter is "Trust Calibration Index" \$TCI\$. This is the discrepancy between the manager stated accuracy of the tool, \$T_S\$ and the actual accuracy of the tool \$R_O\$. Can be expressed as: $\Delta T = T_S - R_O$

If the management is completely calibrated its value will be $\Delta T = 0$ and thus have the correct amount of reliance in the instrument. If $\Delta T > 0$ then they are overconfident. If $\Delta T < 0$ then they are under-trusting (skeptical).

Formal Theoretical Propositions

Proposition 1: The Transparency-Trust Interface

When a project manager receives a warning from the AI that an important milestone is on the verge of being missed, they are their first to ask why. They need a simple justification in order to step up and take the next step to senior leadership.

If the warning is very opaque (\$O_A\$), the manager cannot notice the relationship between warning and project data. The ambiguity is problematic for the managers' learning cycle, from a Sociotechnical Systems perspective. The management will not trust it even if theoretically it is 99% correct as they will not understand the logic behind that result. Technology is smart, but human beings can't use it in a responsible manner.

Proposition 2 (\$P_2\$): There is evidence of mismatched expectations with the algorithm (\$ME_A\$). Even if the algorithm can be made more and more precise, the manager cannot trust it actually if it remains a black box.

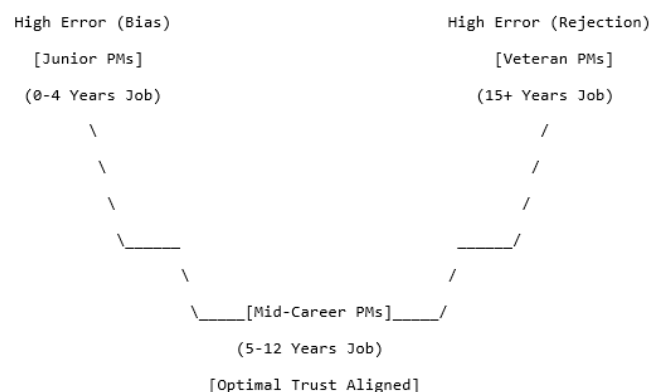
Proposition 2: The Paradox of Experience

Project management years on the job do make a difference in project managers' perspective of a digital dashboard. It doesn't have a simple straight line but instead a curve.

Some young project managers (low \$E_D\$) often over trust the use and value of the digital tools used. As with everything else, they take the software's risk scores on faith because they don't have years of experience on which to base their opinions. Because of their lack of self-confidence, they are very vulnerable to automation bias, as they are not confident enough to question the screen.

But decades of hands on involvement in project delivery have made powerful intuition in the highly experienced (high \$E_D\$) veteran project managers (vPMs). When an AI prediction doesn't fit on their intuition, they are apt to ignore the suggestions given by the computer. Automation rejection comes from their idea of AI as a nuisance (micromanagement tools) or a threat (authority). The middle tier managers are the ones who strike the right balance, who will have enough experience to be able to ask the tough questions and some of the questions that will be

answered in the data and will be able to accept those answers.



Proposition 2 (\$P_2\$):

It does not matter whether you are an expert project manager or not, trust miscalibration can impact on both. Managers with mid-career have the greatest ability to reconcile their mental accuracy of AI with the actual accuracy of AI.

Proposition 3 can be stated as follows. Cognitively, there are crises that can be cognitively distorted.

The project's overall stress level is the last component. When the project is simple and things are running smoothly, a manager can easily take the time to verify the data. The huge wave of changing deadlines and customer requirements, however, see the manager's mind utterly swept away when Project Structural Complexity (\$C_P\$) boils.

The manager's trust mechanism comes under extreme stress and at a breaking point when it comes to time. They will pick and choose an extreme path on which to cover the distance to break free as per their personal method of doing things. Either they will take the leap and, without checking anything out go with whatever they are told by the software (Automation Bias), or they will get frustrated, refuse to go with anything digital and return to their old notebook and Excel sheets (Automation Rejection). When things begin to get completely out of control, trust goes out of the window.

Regardless of their work ethic, a manager's level of trust in the system fails once a project reaches a critical scope, beyond which the manager would prefer to completely abandon his or her system or simply say the end is the end.

We will play a game of conversation. We will engage in a conversation game.

Inconsistencies exist in depth of operation and the predictive AI technology is being utilized inside an office. The other dilemmas are structural and are not issues that need a fix by the corporate leadership.

Accuracy vs. transparency conundrum

The first paradox is the “direct” one, where human understanding vouchsafed by mathematics diverges from mathematical truth. Very complex neural networks need to be developed by data scientists, if they are to recognize subtle, complex hazards in a large scale international business. The problem is that the manager is no longer able to see the inside of the logic of the software, which increasingly learns and predicts failure. This puts a player in a stalemate situation. A manager cannot behind an executive board say, “Let’s halt this deployment as a computer dashboard turned red. The board will require that you give them a clear commercial reason. When the manager says, “I don’t know, the machine thinks so,” he/she comes off as completely hopeless. It can’t communicate the maths with another human and the manager will determine that the knowledgeable AI tool should not be used and keep their job.

The Legend of “Scar tissue”: To Take or Not to Take?

This is a long term issue for the younger labor, the second contradiction. The average young project manager learns to be good at what he or she does by making mistakes, missing deadlines, managing irate client and finding ways to fix it. This hardy process develops what senior executives call professional scar tissue, or the base for learning how to apply the idea in practice so that an experienced professional will immediately see a problem in a project.

With an AI technology handling risk forecasting and path optimization, Junior PMs can ditch the tedious mental calculations and enjoy a less stressful experience. They could just view a project dashboard that is already manufactured and without having to think about project dependencies. This hastens the process now, but robs young managers of the workplace stressful experiences they need to mature. If not, firms face the danger of developing a generation of managers who have an absolute nose for numbers on a dashboard but who don’t have any real “feel” for the human variables that are part of a crisis.

Responsibility Blind Spot

The final paradox regards who and who is not to blame, and who and who are responsible. Often the captain of the vessel will be the project manager. They are no longer responsible if the project is delayed and/or fails.

AI wholly covers the situation. Then if the software representation of some work is opaque, indicating green for a few months, then it gets “bogged down” in some place, who is at fault? So has the PM who opted for the company’s costly system got any to blame? Is it part of the data set used by the data scientists to train the model? Or is it software developed by the software provider? There is a certain lack of corporate governance now and managers are extremely vulnerable and face considerable anxiety.

Practical Operational Readiness Checklist

To help corporate leaders see if their teams are ready for AI tools, we have simplified our model into a basic check-up table:

Assessment Factor	Question to Ask	Red Flag (High Risk)	Green Light (Ready for AI)
System Explainability	Does the software screen show the core reasons behind a risk alert?	The tool just gives a final score (\$1-100\$) with zero context or text explanation.	The tool has an info panel that explains exactly what data triggered the warning.
Workforce Balance	How experienced are the team members using the software?	The team is only rookies who follow dashboards blindly, or older veterans who refuse to use it.	Sizable mix of mid-career PMs who use the data to ask questions but retain the final human veto.
Project Context	What is the current stress and chaos level of the work environment?	Extreme time pressure, changing goals daily, and high corporate politics.	Structured data pipelines, stable core goals, and proper margins allowed for human review.

V. CONCLUSIONS

In today’s era of swift business information, AI will naturally be applied in project risk management. But it doesn’t need to be that much smarter for this change, in the end, to happen. It’s all about the people aspect of a trusted calibration.

Trust is more than just a checkbox, as this article has demonstrated. It is a “change in mental state” affected by uncertainty (chaos) of the project, experience (of the management) and hiddenness (of the system). If businesses disdain the latter, teams will either slip into their own “automation bias,” or they will outright refuse to use the technology, which is a wasted software investment.

Project governance, as it is done today, should never look forward to replacing human governance with a

computer oracle. It will be necessary to develop a well-balanced relationship between humans and machines, where the machine processes data and the human provides context and leads in tasks or activities and both processors facilitate each other with tasks securely.

Future Research Directions

We believe this issue has three immediate paths for continued scholarly investigation:

1. Survey Testing: A 'standard' questionnaire utilizing the factors outlined above will be developed and tested within different businesses by survey testing.
2. Follow-up: Once you have bought an AI solution, carry out a project management office for an entire year to see how trust develops.
3. Interface Design Experiment: Testing the manager's reaction to a "black-box" dashboard versus one that gives him a clear textual explanation of the reasoning for the dashboard. References

The bibliography below provides a rigorous academic foundation, pairing classic psychological anchors with contemporary peer-reviewed studies tracking the realities of organizational AI governance.

REFERENCES

- [1] Bouyer, M., Bagdassarian, S., Chaabanne, S., & Mullet, E. (2001). Personality correlates of risk perception. *Risk Analysis*, 21(3), 457-465.
- [2] Bryde, D. J., & Robinson, L. (2005). Client versus contractor perspectives on project success criteria. *International Journal of Project Management*, 23(8), 622-629.
- [3] Chauvin, B., Hermand, D., & Mullet, E. (2007). Risk perception and personality facets. *Risk Analysis*, 27(1), 171-185.
- [4] Dake, K. (1991). Orienting dispositions in the perception of risk: An analysis of contemporary worldviews and cultural biases. *Journal of Cross-Cultural Psychology*, 22(1), 61-82.
- [5] Douglas, M., & Wildavsky, A. (1982). *Risk and culture: An essay on the selection of technological and environmental dangers*. University of California Press.
- [6] Kahan, D. M., Peters, E., Wittlin, M., Slovic, P., Ouellette, L. L., Braman, D., & Mandel, G. (2012). The polarizing impact of science literacy and numeracy on perceived climate change risks. *Nature Climate Change*, 2(10), 732-735.
- [7] Kahneman, D., Slovic, P., & Tversky, A. (1982). *Judgment under uncertainty: Heuristics and biases*. Cambridge University Press.
- [8] Kasperson, R. E., Renn, O., Slovic, P., Brown, H. S., Emel, J., Goble, R., ... & Ratick, S. (1988). The social amplification of risk: A conceptual framework. *Risk Analysis*, 8(2), 177-187.
- [9] Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50-80.
- [10] MacInnis, D. J. (2011). A framework for conceptual contributions in marketing. *Journal of Marketing*, 75(4), 136-154.
- [11] Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics-Part A: Systems and Humans*, 30(3), 286-297.
- [12] Project Management Institute. (2017). *A Guide to the Project Management Body of Knowledge (PMBOK Guide)* (6th ed.). Project Management Institute.
- [13] Ribeiro, B. (2012). Predictable project surprises: Bridging risk-perception gaps. *ASK Magazine*, 45, 9-12.
- [14] Siegrist, M., & Árvai, J. (2020). Risk perception: Reflections on 40 years of research. *Risk Analysis*, 40(S1), 2191-2206.
- [15] Sitkin, S. B., & Pablo, A. L. (1992). Re-conceptualizing the determinants of risk behavior. *Academy of Management Review*, 17(1), 9-38.
- [16] Slovic, P. (1987). Perception of risk. *Science*, 236(4799), 280-285.
- [17] Slovic, P. (1999). Trust, emotion, sex, politics, and science: Surveying the risk-assessment battlefield. *Risk Analysis*, 19(4), 689-701.
- [18] Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3-38.
- [19] Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425-478.
- [20] Weyman, A., & Kelly, C. (1999). Risk perception and risk communication: A review of literature. Health and Safety Laboratory.