

Principles of Machine Leadership

An integrated framework for leading in the age of AI

David Swanagon and Stephen McIntosh, The Machine Leadership Journal

Abstract: This paper provides an integrated coaching framework for the age of AI. The literature indicates that trust and skill deficiencies limit AI adoption. We introduce a coaching model that focuses on the key dimensions of Machine Autonomy, Trust, and AI Competencies, including an equation that helps coaches track performance. The significance of this model is that it provides coaches with a framework to engage leaders to better plan, manage, and sustain AI adoption.

Keywords: Business Coaching, Machine Leadership, AI adoption

Introduction

In the age of AI, executives will be required to lead machines, lead people that build machines, and lead organizations that adopt AI. The purpose of this paper is to introduce a model that coaches can use to drive AI adoption within organizations. The research has focused on the role that Machine Autonomy, Trust, and AI Competencies have on AI adoption. This includes analyzing factors that encourage equilibrium between these three dimensions, conditions that create divergence, and the implications for coaching.

The Machine Leadership Model provides a framework for coaches to optimize AI adoption by focusing on the balance between Machine Autonomy, Trust, and AI Competencies. Research indicates that many AI use cases have achieved near linear computational complexity for baseline operations, traditionally expressed as $O(n)$. This breakthrough in AI engineering has allowed highly complex models to scale in proportion to the size of their input. However, our findings suggests that AI models are influenced by the relationship between Machine Autonomy, Trust, and AI Competencies. In situations where $(A=T=C)$, there is no moderating impact on $O(n)$. However, any deviations from this creates a positive moderating balance. This results in higher AI adoption costs that increase based on the degree of imbalance. This level of technical complexity is important to coaching due to the unique challenges that leaders face integrating machines with the human workforce such as Responsible AI, data privacy, and employee skill programs.

This paper will provide a review of existing literature, followed by a description of the Machine Leadership Model, an introduction of the moderating equation, and coaching strategies for improving AI adoption. The significance of this model is that coaches can help AI engineers manage AI adoption in a manner that drives $(A=T=C)$. This approach optimizes organizational performance, further integrates the hybrid workforce, and lowers costs.

Methodology

To develop the Machine Leadership Model, the researchers conducted a mixed methods exploratory design. The first step consisted of a literature review that evaluated the theories and

best practices impacting AI adoption. This included analyzing the role that Human Autonomy, Trust, and AI competencies have on AI leadership. Furthermore, the team examined existing coaching models and their adaptations for AI leadership.

The researchers completed stakeholder interviews with fifteen AI engineers and corporate leaders. Separately, the team administered a 75-question survey to twenty participants. The survey was designed as a pilot instrument. The team investigated existing assessments and AI performance metrics such as Hogan's inventories, Korn Ferry's Leadership Architecture, the BBH Benchmark, and AGIEval. The researchers determined that unique questions were needed to directly address the factors influencing AI adoption. Survey data was collected from a variety of sources including: one AI Researcher, eleven AI Engineers, and eight Corporate leaders. Respondents were asked to provide their job role, industry experience, and geographic location. An ipsative 'forced choice' method was used for the survey questions. Furthermore, the researchers utilized the literature to identify a method for measuring baseline AI model complexity. From this, the team developed a moderating equation that accounts for the impact that Machine Autonomy, Trust, and AI Competencies have on AI adoption.

Regarding data analysis, the researchers mapped each respondent's Autonomy, Trust, and AI Competency level using their ipsative scoring from the survey. Separately, two coders evaluated the stakeholder interviews. The team determined that a fixed codebook and Kappa inter-coder reliability calculation was unsuitable for this study. Instead, the focus was on extensive collaboration. A comparative analysis was used to finalize key themes.

The moderating equation utilized a baseline complexity assumption, expressed as $O(n)$ for a standard feedforward neural network. The decision to use this neural network was to provide a simple approach for future studies. The calculation assumed an input layer with 4 nodes, a hidden layer with 8 nodes, a second hidden layer with 6 nodes, and an output layer with 2 nodes. The total operations for this baseline model were defined as: 40 (Layer 1) + 54 (Layer 2) + 14 Layer (3) = 108. The 108 represents the FLOPs (Floating-point operations) for a single forward pass, which Meng et al. (2024) describe as a standard measure of AI computational complexity. The researchers anchored the moderating equation against $O(n)$ to ensure that coaching strategies were aligned to AI industry practices for determining computational costs.

The moderating equation was then applied to the survey scores. This resulted in an AI adoption penalty based on the degree of imbalance, expressed as $M(A, T, C) > 0$. The calculated penalty is summarized in Figure 2. Finally, the adjusted scores were applied to a new coaching model specifically designed to optimize AI adoption. Existing coaching models such as GROW, OSCAR, and FUEL were reviewed for alignment. However, the researchers determined that a new framework was needed to specifically achieve $(A=T=C)$.

Theoretical Framework/Literature Review

The researchers focused on five themes to evaluate the factors influencing AI adoption. First, the team examined existing coaching models to identify best practices and understand how traditional frameworks have been adapted to AI. Then, a detailed analysis was conducted concerning the role

of Autonomy, Trust, and AI Competencies in driving organizational performance. Finally, the team investigated the process used to calculate computational complexity. Combined, the literature revealed that Autonomy, Trust, and AI Competencies are critical for AI adoption. Furthermore, that there is a standard method that companies use to calculate computational complexity, which can be applied to measure AI adoption. Lastly, that the interdisciplinary leadership of coaches can serve as a bridge between technical and non-technical teams seeking to address AI challenges.

Coaching Models Adapted for AI Leadership

The field of leadership coaching has well established models such as GROW, FUEL, and OSCAR for addressing employee development and team performance. Passmore et al. (2024) argued that coaches should adapt these models by becoming technology literate and embracing digital tools. This includes leveraging data analytics to assist with the Reality, Understand, and Situation stages of GROW, FUEL, and OSCAR, respectively. Likewise, utilizing AI tools such as Chatbots and digital platforms to scale services. AI tools can also assist with the development of action plans and accountability metrics. The researchers argue that AI cannot replace the coaching process. Instead, it should augment the models through robust data insights, automation, and process efficiencies. This “Human in the Loop” process requires coaches to develop new competencies that strengthen data literacy and generative AI knowledge.

A review of existing coaching models identified several themes that were relevant for AI. The seminal work by Sir John Whitmore (2002) provided the GROW framework. This approach emphasizes the need for robust goal setting, understanding reality, brainstorming options, and building action plans. This approach directly applies to AI Engineering tasks such as model selection, MLOps, and data privacy. Zenger and Stinnett (2010) added to this by introducing the FUEL model. This framework focuses on conducting quality coaching conversations. This is done by framing the purpose of the conversation, understanding the situational context, exploring potential solutions, and laying out action plans. This approach helps AI Engineers address domain specific challenges impacting models tied to Healthcare, Finance, Energy, and so forth.

The OSCAR model developed by Gilbert and Whittleworth (2009) compliments these concepts by helping leaders transition from their current state to a desired outcome. The model follows a structured process that includes identifying a future state, analyzing the situation, exploring choices, creating action plans, and reviewing results. This approach is useful in building robust technology roadmaps that help organizations integrate increasingly complex machines.

Methods for Calculating Computational Complexity

NIPS (2014) provided a comprehensive review of the core components that determine computational complexity for neural networks. This includes forward and backward passes and the reason these models follow a linear computational cost. Cormen et. al (2009) identified the expression $O(n)$ as a way of showing the efficiency and scalability of an algorithm. The notation indicates that runtime and memory usage grows as the input size (n) increases. The literature indicates that $O(n)$ is an effective way of showing the linear change in computational complexity as AI models gather more data, perform additional calculations, and integrate with other platforms. This method was used to support the Machine Leadership Model by establishing a baseline

computational cost for AI tools that is focused on model design. From this baseline, the AI tool is moderated by the Machine Autonomy, Trust, and AI Competencies within an organization.

Human Autonomy and Meaning in the Workplace:

According to Deci and Ryan (2000), Self-Determination Theory stipulates that autonomy and relatedness are critical factors in driving intrinsic motivation. This is important because employees that experience a sense of autonomy in their work are more likely to be committed despite experiencing external pressures. This was supported by Hackman and Oldham's (1976) Job Characteristics Model, which highlighted that jobs high in autonomy and meaning produce more effective work outcomes. Wrzesniewski et al. (2003) researched the role that autonomy has on developing a sense of meaning for employees and its positive effects on satisfaction. Separately, Pink (2009) noted that creative work which involves non-routine tasks requires a high degree of autonomy and purpose for employees to perform at high levels. These studies support the role of human autonomy as a key dimension in driving AI adoption using the Machine Leadership Model.

The Importance of Trust

Dirks (2000) noted that trust is a fundamental component of leadership effectiveness. Teams that trust their leader's experience, temperament, and decision-making are more likely to perform at peak levels. This is supported by Simons (2002), which emphasizes the importance of leaders behaving in a consistent, predictable, and authentic manner. Mayer et al. (1995) noted the importance of character traits such as integrity, benevolence, and ability on cultivating trust. These concepts were further segmented by McAllister (1995) who defined classes of trust such as affect-based and cognition-based. The former focuses on the importance of building emotional bonds, while the latter is based on the leader's perceived competency level. The research also emphasized that affect and cognition-based trust must be balanced with cooperation. Zaheer et al. (1998) highlighted that teams which build high degrees of trust are more effective at engaging in interpersonal communication and cross-collaboration.

The literature indicated that Trust is an essential element of the Machine Leadership Model due to its influence on AI adoption, especially when balanced with autonomy. The researchers integrated these studies with the stakeholder interviews and survey results to establish a definition of trust that applies to the Machine Leadership Model. In terms of AI adoption, the team defines trust as the belief in the ability and willingness of a human or machine to meet your expectations.

Building AI Skills in Organizations

Manyika et al. (2017) highlighted the on-going and persistent AI skills gap across key domains such as machine learning, robotics, ethical AI, and data science. The lack of critical skills has fundamentally impacted the degree of AI adoption across industries. Van Laar et al. (2017) extended this concept by examining the complexity associated with defining the required skills for AI professionals. More so than most industries, AI requires engineers to have strong technical skills, problem solving, and interpersonal traits. Purdy and Daugherty (2017) added to this by noting the importance of integrating AI skills across business functions. This means technical and

non-technical professionals must have a fundamental understanding of AI technologies to lead a workforce comprised of AI agents and humans.

The literature emphasized the importance of learning agility. Wilson et al. (2017) examined the pace of AI innovation and noted the importance of upskilling and continuous learning on organizational performance. Companies need to continually reinvest in learning programs to respond to disruption and align new technologies with their enterprise strategy.

These studies indicated the importance of coaches in supporting the development of AI Competencies within an organization to ensure that Machine Autonomy is balanced with Trust. Furthermore, the interdisciplinary leadership of coaches can serve as a bridge between technical and non-technical teams seeking to address challenges with model design, utilization, and employee skills. This applies to the Machine Leadership Model as the decisions impacting Machine Autonomy, Trust, and AI Competencies involve multiple cross-functional stakeholders.

Summary of Major Findings

The research produced two themes. The first theme focused on computational complexity and moderation, while the second theme addressed the role of coaching in AI adoption. The first theme stipulates that AI adoption is influenced by the baseline complexity of the model, expressed as $O(n)$, which is moderated by the interaction between Machine Autonomy, Trust, and AI Competencies. The second theme suggests that coaching can play a significant role in balancing the moderating effect of Machine Autonomy, Trust, and AI Competencies, which will help organizations lower the cost of AI adoption.

The AI Innovation Frontier is the equilibrium where Machine Autonomy, Trust, and AI Competencies intersect. This is the place where AI tools are utilized in the most efficient manner. The frontier operates at various stages of model complexity. This allows organizations to utilize a standard framework for comparing novel inventions to large scale breakthroughs. The frontier seeks a positive correlation between Machine Autonomy and Trust, which means as autonomy increases, trust also increases at a rate moderated by AI Competencies.

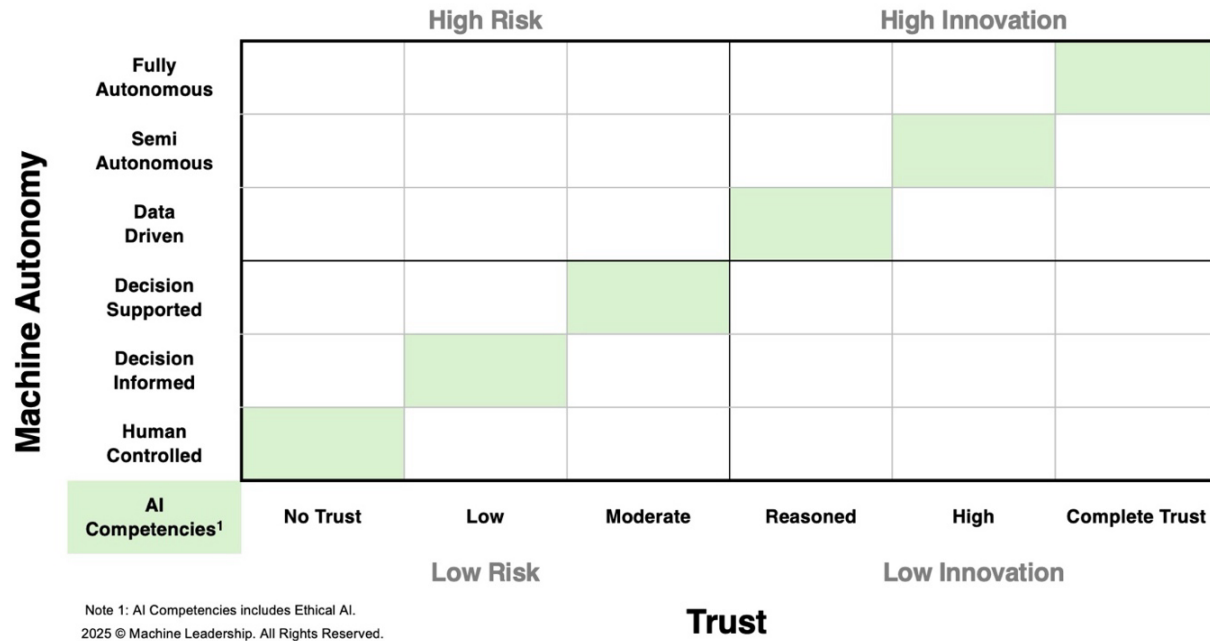
The moderating function for $O(n)$, which evaluates the equilibrium between Machine Autonomy, Trust, and AI Competencies is listed below.

$$\begin{aligned} \text{AI Adoption} &= O_{\text{influenced}}(n) = (1 + M(A, T, C)) * O(n) \\ \text{Moderating Function} &= M(A, T, C) = k * ((A-C)^2 + (T-C)^2 + (A-T)^2) \end{aligned}$$

The point where Machine Autonomy, Trust, and AI Competencies intersect is referred to as The AI Innovation Frontier. The reason is because there is no moderating influence over the general performance of an AI tool, expressed as $O(n)$. The moderating equation indicates that as AI tools scale, their adoption rate is influenced by the interaction between these variables. In situations where the moderating variables are at equilibrium, the formula is $(A=T=C)$. The moderator M quantifies the degree of deviation from this state of equilibrium.

Figure 1

The Machine Leadership Model



Deviations from the state of equilibrium create a positive moderating balance, expressed as $M(A,T,C) > 0$. This scenario incorporates a positive scaling constant k , alongside the sum of squared differences between the three elements. This is reflected using the following formula: $M(A,T,C) = k * ((A-C)^2 + (T-C)^2 + (A-T)^2)$. The k is a scaling constant that determines the degree of sensitivity towards $O(n)$ that results from deviations to the equilibrium. The first term $(A-C)^2$ measures the difference between Machine Autonomy and AI Competencies. The second term $(T-C)^2$ measures the difference between Trust and AI Competencies. The third term $(A-T)^2$ measures the difference between Machine Autonomy and Trust.

Proceeding to AI adoption, the formula $O_{influenced}(n) = (1 + M(A,T,C)) * O(n)$ can be analyzed as $O(n)$ represents the standard computational complexity of an AI model and stipulates the linear growth of a model as the input size increases. Cormen et. al (2009) identified this as a common measure showing the processing time and memory that change in proportion to the size of the inputs. When AI tools are out of balance, the baseline computational complexity $O(n)$ is moderated by a number greater than 0, expressed as $M(A,T,C) > 0$. This results in a penalty that increases proportionally based on the degree of imbalance.

The research team utilized a baseline complexity calculation for a standard feedforward neural network, expressed as $O(n)$. The total operations for this baseline model were defined as: 40 (Layer 1) + 54 (Layer 2) + 14 Layer (3) = 108 FLOPs. After this, respondent survey results were added to the moderating equation to determine if a computational penalty was needed. The

results are listed in Figure 2. The findings indicate that organizations with these model scores would incur a 200% penalty on their AI adoption programs.

Figure 2

AI Adoption Penalty - Summary of Respondent Survey Scores (Respondents = 20)

Baseline	Autonomy	Trust	Competencies	M(A,T,C)	Penalty
108 FLOPs	3	2	2	2	324 FLOPs

Using the survey respondent pool, the average score for Machine Autonomy was 3, Trust 2, and AI Competencies 2. The baseline complexity = 108 FLOPs. The new adjusted FLOPs incorporating this lack of equilibrium is 324, which represents a 200% increase in complexity. These are AI adoption costs that are incremental to the model's engineering design, which organizations incur to build trust and employee skills in new AI tools. For example, the penalty may require organizations to slow down deployment, incur higher training costs, or add Responsible AI monitoring platforms to achieve equilibrium.

$$M(3, 2, 1) = k * ((3-2)^2 + (2-2)^2 + (3-2)^2) = 2$$

$$O_{influenced}(n) = (1 + 2) * 108 = 324 \text{ FLOPs (+200\%)}$$

The findings suggest that coaches can directly impact the moderating function. For example, in situations where an AI tool has low Trust and high Autonomy ($T < A$), organizations may need to engage in significant monitoring to ensure issues such as Responsible AI and data privacy are mitigated. Coaching can reduce these issues by working with engineers and corporate stakeholders to improve organization trust. This could be achieved by using the Machine Leadership Model and coaching tactics listed in this paper.

Comparatively, when Trust is high and AI Competencies are low ($T > C$), additional protocols may be needed to manage data ingestion errors, input validation, and cyber security breaches. In this situation, engineering teams would need to ensure that employees developed baseline skills in threat prevention. An example would be a phishing email program. Coaches could support this program by developing action plans for teams and individual leaders to ensure that trust and skill programs were implemented. The AI adoption formula could be used to check if the moderating variables have reached equilibrium, and what additional steps were needed.

Machine Leadership Coaching Model

Current models such as GROW, OSCAR, and FUEL focus on coaching conversations that emphasize how the leaders focus on attaining goals, motivate and lead their teams, and navigate the complexity of change. These models are effective in the context of leaders leading people, but leading an AI machine adds additional complexity. To effectively coach AI leaders, a new coaching model is required that accounts for the Machine Autonomy and purpose of AI tools.

AI adoption is highly dependent upon the Trust created between the leader and the hybrid workforce. The research team developed an applied definition of Trust stipulating, "Trust is the

belief in the ability and willingness of a human or machine to meet your expectations.” In the hybrid workforce context, we can look at the variables of “ability” and “willingness” as being applicable to both humans and machines.

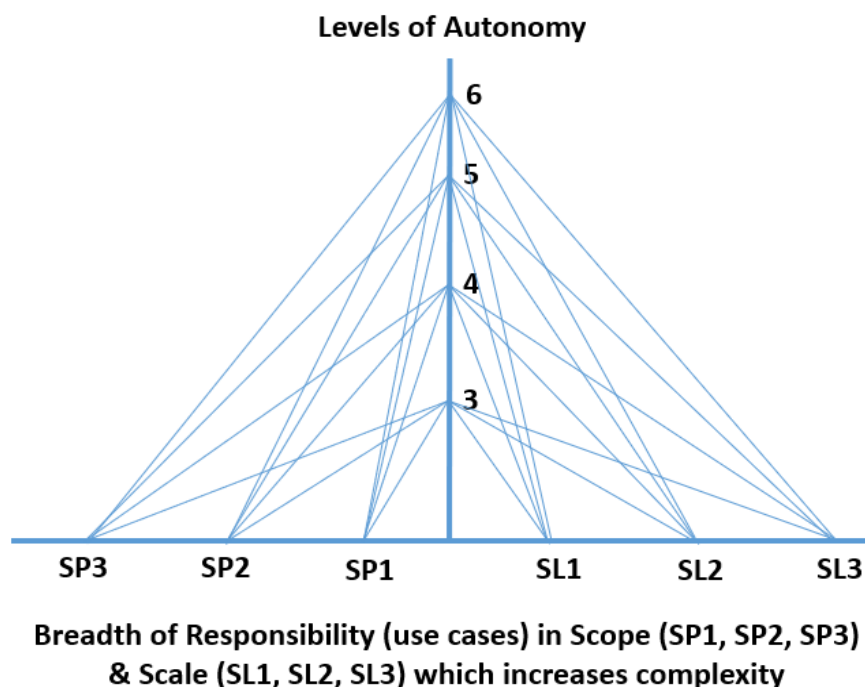
The term “ability” is defined as the breadth of responsibility given to successfully manage a use case at both the scope and scale expected. According to Deloitte (2024), autonomy is the power to act and make decisions independently. This concept can be extended to agentic AI. For these platforms, the goals are set by humans, but the agents determine how to execute them. Thus, each AI tool has a level of autonomy, which is described in Figure 1 by six levels. Levels three through six are coaching priorities because they indicate the AI tool is actively participating in the decision-making process from creating data driven criteria to being fully autonomous.

Figure 3 illustrates how Scope (SP) and Scale (SL) interact to guide the decision on establishing the right autonomy level given to the AI Engineer or Machine. The combination of the two is the Breadth of Responsibility. When the Level of Autonomy is decided, then this coaching variable becomes the Degrees of Freedom that determines Ability. The ability of the engineer or machine to accomplish an increased scope or higher scale of work complexity leads to increased risk.

Therefore, the Breadth of Responsibility accounts for the progress through significant expansions in Scope as moving from SP1 to SP2 and SP3, while significant Scale increases are noted as SL1, SL2, and SL3. The Breadth of Responsibility can iterate across many levels until the underlying technology changes to the point a new Degrees of Freedom chart should be created.

Figure 3

Degrees of Freedom Chart - Levels of Autonomy by Breadth of Responsibility



Coaches working with AI Leaders will first ask the leaders to complete the Machine Leadership Model shown in Figure 1, while using the moderating equation in Figure 2 to determine the current AI adoption trends. Once an initial score is established, a discussion around the Breadth of Responsibility of team members (both humans and machines) and at what level of autonomy and skillset should be targeted. The objective being to achieve ($A=T=C$).

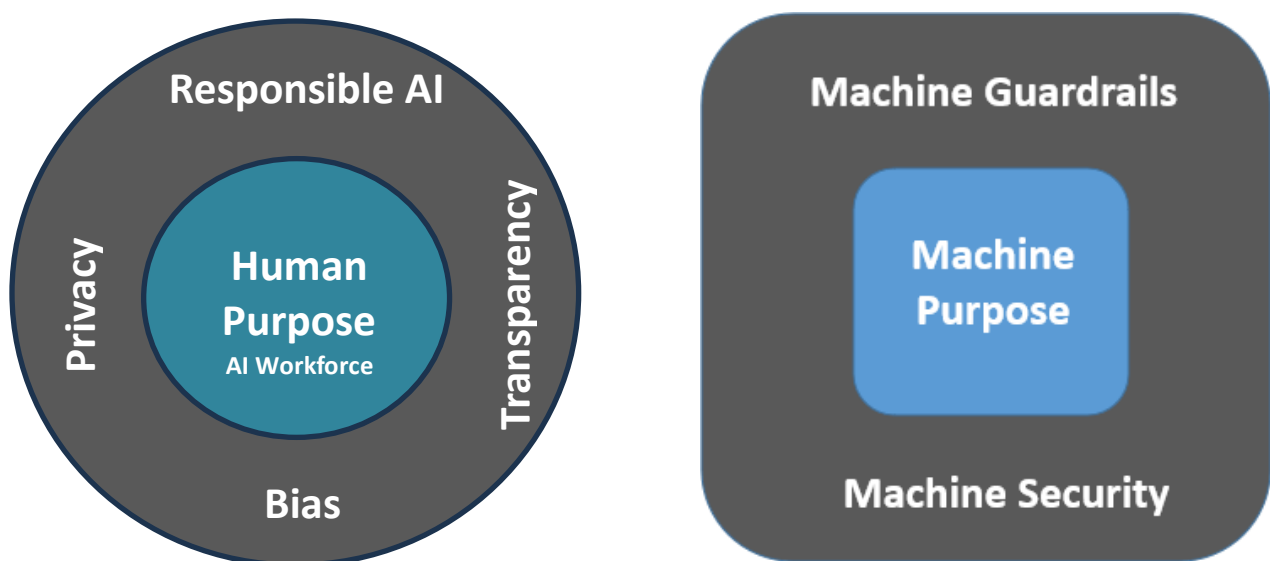
As the pace of AI innovation continues to scale, two more levels of autonomy will be added: Level 7 for Artificial General Intelligence (AGI) and Level 8 for Artificial Super Intelligence (ASI). These autonomy levels present a significant leap in how AI Leaders interact with their technology, since AGI and ASI will put AI Leaders and coaches at an asymmetrical information disadvantage concerning how the machines operate and make decisions.

Following the Degrees of Freedom determination, Figure 4 illustrates major factors involved in ensuring the alignment of the human and machine purpose with the leader's purpose to determine the willingness to meet the leader's expectations. For example, AI agents operate with a purpose, whether it is a SLAM algorithm or LLM. They are programmed and designed to deliver a workflow. If a leader has a different purpose for the AI agent, then the misalignment of purpose will lead to the multi-agent system being modified or a new AI model being created to drive the new purpose.

Several factors should be considered when a coach helps a leader decide on a Machine. Does the purpose still align with the Degrees of Freedom of the machine such as the scale, scope of responsibilities, and level of autonomy? If the new scope of responsibilities has been added through a redesign of the multi-agent system, then the purpose of the machine may have expanded. This would necessitate a discussion regarding the possibility of installing new guardrails.

Figure 4

Alignment of Purpose determines Willingness



Alignment of purpose is a critical coaching activity for the hybrid workforce. For example, the coach needs to help the AI Leader consider the possibility of bias that an AI Engineer may program into the machine. If that bias is not supported by the leader or the company, then the AI Leader is not aligned on purpose with the AI Engineer. The ability may be strong, but additional steps are needed to build trust. The alignment of purpose determines the degree of trust, which the researchers defined as part of the findings as the willingness of the person or machine to meet the leader's expectations. A highly capable machine or AI Engineer with a misaligned purpose and high Degrees of Freedom can present a significant threat to a company.

The Machine Leadership Coaching Model supports Responsible AI by providing a framework to consider the expanding responsibilities AI is given and how much autonomy should match the responsibilities. Furthermore, it guides the conversation about whether the Degrees of Freedom given to the AI Engineer or machine aligns with the purpose. If a fully autonomous AI system with significant control over essential business processes is not willing to work within its intended purpose, then company leaders must make proactive changes to achieve (A=T=C).

Coaches will engage different audiences with a slightly varied approach on how to discuss (A=T=C) and apply the Machine Leadership Model. For the executive levels, coaches may focus on the need for careful consideration of autonomy levels as the pace of AI innovation scales. This also includes a discussion about the purpose of the AI tool and how to align it with the goals of the company. A key challenge for coaches will be to help organizations find the right equilibrium on the Machine Leadership Model to ensure technologies are fully adopted. Figure 5 provides examples of adoption tactics that can be employed based on the audience.

Figure 5

Adoption Tactics by Audience

Audience	Autonomy Tactics	Trust Tactics
Executives	Use Case Discussions Scenario Planning AI Council	Evaluation and Monitoring Practice Labs Reverse Mentoring
AI Workforce (AI Builders)	Bessemer Scale Review Risk Analysis and Testing Iterative Development	Data Validation AI Transparency Human-centered Criteria
General Workforce (Internal AI Users)	Pilot and Control Groups Quality Assurance Testing Systems Mapping	Positive Experiences Learning Communities AI Competency Building
Customers/ Clients (External AI Users)	Incentives Expanding Scope / Features Ethical AI Disclosures	Human-in-the-Loop AI Success Stories Customer Profiling / Memory

Conclusions

The Machine Leadership Model is an innovative framework that helps coaches drive AI adoption within organizations. The framework highlights the interconnected relationship between Machine Autonomy, Trust, and AI Competencies. The AI Innovation Frontier is the optimal state where the three dimensions intersect. This is the place where AI adoption is most efficient for an organization. In situations where $M(A,T,C) > 0$, the imbalance among the three variables moderates AI adoption. This means organizations are forced to incur additional costs to ensure AI tools are utilized in a responsible manner. Coaches can have a positive impact on this process by serving as critical bridges between AI machines, engineering teams, and corporate leaders. The coaching model incorporates AI tools while helping individuals and teams build action plans that optimize trust for AI adoption. For example, a coach may work with a physician to develop the trust and competencies needed to safely utilize a surgical robot.

The findings of the stakeholder interviews, respondent survey, and moderating equation suggest that AI adoption can be positively influenced by coaching. This was evidenced by the 200% penalty that was incurred from the baseline computational costs $O(n)$ that was observed in a standard feedforward neural network. The survey presented a wide distribution of scores for Machine Autonomy, Trust, and AI Competencies. These results indicate that AI adoption may be influenced by factors outside of engineering design. To reduce the impact of these moderating variables, coaching models should focus on the level of autonomy that is given to an AI Engineer or machine by scale and scope. As part of this, coaches and AI leaders must ensure that trust is being established by aligning individual purpose with the goals of the team.

Additional Lines of Inquiry

This paper acknowledges the data limitations of the moderating equation in calculating aggregate AI adoption penalties for entire organizations. The approach would benefit from future studies that included a larger sample size of survey respondents, along with a broader portfolio of AI tools. The baseline complexity calculation used a simple feed forward neural network. Future studies should incorporate domain adapted AI tools to address industry needs such as healthcare robotics, finance LLMs, and energy deep learning algorithms. This will help determine the degree of sensitivity k that an AI tool has within a situational context. Separately, the coaching model should be incorporated into the analysis to determine the impact of the process along with the required time to change. This could include experimental design studies for specific technologies or longitudinal studies for entire industries.

Warrant Statement:

We warrant that given our paper, experiential session, demonstration or panel proposal is accepted, we will submit a formally written summary for inclusion in the conference proceedings. We agree that the summary will be typed and single-spaced and will respect the maximum number of words expected. We understand that if this summary is not submitted by July 28, 2025 our presentation will not be included as part of the Columbia Coaching Conference in New York City 2025. We also agree that formatting of the document according to conference specifications is our responsibility, and we understand that the document will be returned to us if it does not meet those specifications.

References

- Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). *Introduction to algorithms* (3rd ed.). The MIT Press.
- Deci, E. L., & Ryan, R. M. (2000). The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268.
- Deloitte. (2024) TMT Predictions 2025: The Gap Year.
- Dirks, K. T. (2000). Trust in leadership and team performance: Evidence from NCAA basketball. *Journal of Applied Psychology*, 85(6), 1004–1012.
- Gilbert, A., & Whittleworth, K. (2009). The OSCAR coaching model. Lulu.com.
- Hackman, J. R., & Oldham, G. R. (1976). Motivation through the design of work: Test of a theory. *Organizational Behavior and Human Performance*, 16(2), 250–279.
- Manyika, J., Lund, S., Chui, M., Bughin, J., Woetzel, J., Batra, P., Koenecke, E., & Sanghvi, S. (2017). *Harnessing automation for a future that works*. McKinsey & Company.
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734.
- McAllister, D. J. (1995). Affect- and cognition-based trust as foundations for interpersonal cooperation in work groups. *Academy of Management Journal*, 38(1), 24–59.
- Meng, X., Chen, W., Benbaki, R., & Mazumder, R. (2024). *FALCON: FLOP-Aware Combinatorial Optimization for Neural Network Pruning*. Proceedings of The 27th International Conference on Artificial Intelligence and Statistics.
- NIPS. (2014). On the Computational Efficiency of Training Neural Networks. *Advances in Neural Information Processing Systems*, 27. <http://papers.neurips.cc/paper/5267-on-the-computational-efficiency-of-training-neural-networks.pdf>
- Passmore, J., Diller, S. J., Isaacson, S., & Brantl, M. (Eds.). (2024). *The digital and AI coaches' handbook: The complete guide to the use of online, AI, and technology in coaching*. Routledge.
- Pink, D. H. (2009). *Drive: The surprising truth about what motivates us*. Riverhead Books.
- Purdy, M., & Daugherty, P. (2017). How AI boosts industry profits and innovation. *Technology Review*, 120(3), 20–29.
- Simons, T. (2002). Behavioral integrity: The perceived alignment between managers' words and

- deeds as a driver of employee's job satisfaction and retention. *Organization Science*, 13(1), 18–35.
- van Laar, E., van Deursen, A. J. A. M., van Dijk, J. A. G. M., & Huysman, M. (2017). The relation between 21st century skills and digital skills: A systematic literature review. *Computers in Human Behavior*, 72, 577–588.
- Whitmore, J. (2009). Coaching for performance: GROWing human potential and purpose—The principles and practice of coaching and leadership. Nicholas Brealey Publishing.
- Wilson, H. J., Daugherty, P. R., & Bianzino, N. (2017). The jobs that AI will create. *MIT Sloan Management Review*, 58(4), 14–16.
- Wrzesniewski, A., Dutton, J. E., & Debebe, G. (2003). Interpersonal sensemaking and the meaning of work. *Research in Organizational Behavior*, 25, 93–135.
- Zaheer, A., McEvily, B., & Perrone, V. (1998). Does trust matter? Exploring the effects of interorganizational and interpersonal trust on performance. *Organization Science*, 9(2), 141–159.
- Zenger, J. H., & Stinnett, K. (2010). The extraordinary coach: How the best leaders help others grow. McGraw-Hill.